

ОРГАНІЗАЦІЙНО-ПРАВОВІ ЗАСАДИ ВИКОРИСТАННЯ ШТУЧНОГО ІНТЕЛЕКТУ У СФЕРІ ЗАБЕЗПЕЧЕННЯ ДЕРЖАВНОЇ БЕЗПЕКИ

Проведено аналіз використання штучного інтелекту у сфері забезпечення державної безпеки, зокрема у правоохоронній сфері. На основі аналізу вивчених джерел, чинного національного і міжнародного законодавства, а також практичного досвіду визначено ризики застосування можливостей нейронних мереж. Запропоновано обґрунтовані рекомендації щодо вдосконалення нормативно-правової бази стосовно відповідальності, захисту персональних даних, кібербезпеки та етичного використання інструментарію штучного інтелекту.

Ключові слова: штучний інтелект, нейронні мережі, службово-бойова діяльність Національної гвардії України, державна безпека, законодавче забезпечення, кібербезпека.

Постановка проблеми. Сучасний світ характеризується безпрецедентним рівнем цифрової трансформації, що відкриває нові можливості для гармонійного розвитку суспільства і держави загалом. Нейронні мережі, інструменти штучного інтелекту (ШІ) стрімко інтегруються в усі сфери життя, зокрема й забезпечення державної безпеки, створюючи як безпрецедентні можливості, так і серйозні виклики й загрози. З одного боку, ШІ здатний значно поліпшити ефективність дій у запобіганні кіберзагрозам, терористичній діяльності і їх нейтралізації, з іншого – його використання потребує ретельного правового аналізу, аби уникнути різних зловживань, неконтрольованих дій і забезпечити етичне застосування його можливостей.

Актуальність дослідження зумовлена відсутністю чітко визначеної організаційно-правової бази для використання ШІ у сфері забезпечення державної безпеки як в Україні, так і у світі. Чинне законодавство не цілком урахує специфіку застосування ШІ, що призводить до правової невизначеності й обмежує ефективність його використання. Тому розроблення пропозицій з удосконалення організаційно-правових засад використання ШІ у сфері забезпечення державної безпеки вбачається важливим завданням, вирішення якого сприятиме підвищенню рівня національної безпеки й захисту прав людини.

Аналіз останніх досліджень і публікацій. Дослідження ґрунтується на теоретичних засадах, які викладено у працях А. Тюрінга, С. Шеннона, О. Кривецького та деяких інших науковців. Окремі аспекти використання штучного інтелекту стали предметом дослідження вчених-юристів: М. Карчевського, А. Петренка, І. Городиського, Н. Шишки. Вирішенню питань правового регулювання штучного інтелекту у сфері забезпечення безпеки й оборони приділяли увагу К. Ріпер, М. Watney, О. Радутний, Т. Каткова, О. Ястребов, Є. Харитонов та ін. Проте ж дискусійними залишаються питання щодо визначення поняття «штучний інтелект». Невирішеним є законодавче регулювання його застосування, а також актуалізується питання щодо стандартів формування відповідної правової моделі, яка забезпечить належне функціонування сфери державної безпеки і оборони не тільки в Україні, але й у світі.

Мета статті полягає в аналізі організаційно-правових передумов застосування інструментів штучного інтелекту і розробленні рекомендацій щодо створення ефективної та етично виваженої системи регулювання застосування штучного інтелекту у сфері забезпечення державної безпеки в Україні, зокрема й у службово-бойовому застосуванні підрозділами Національної гвардії України.

Виклад основного матеріалу. Ідея створення машин, здатних імітувати людський інтелект, виникла давно. Початкові концепції ідеї штучного інтелекту було закладено ще у середині ХХ ст. Витоки ШІ можна простежити у працях британського математика А. Тюрінга. У своїй статті «Обчислювальні машини та інтелект» (1950 р.) науковець задався питанням: а чи можуть машини мислити? [1]. У відповідь він розробив тест (Тест Тюрінга), що являв собою концептуальну спробу визначити, чи здатна машина спілкуватися таким чином, аби її не можна було відрізнити від людини. Ця праця А. Тюрінга стала підґрунтям для подальших досліджень у галузі штучного інтелекту, оскільки вперше публічно ставилися під сумнів межі можливостей машин.

Офіційним початком наукових досліджень у галузі ШІ вважається 1956 р., коли на семінарі у Дартмутському коледжі було запропоновано термін «штучний інтелект». Організатори заходу, зокрема Дж. Маккарті, М. Мінський, К. Шеннон, прагнули розробити програмне забезпечення, яке могло б навчатись і розв'язувати проблеми. Було репрезентовано 17-сторінковий документ під назвою «Дартмутська пропозиція». У документі висвітлювалися теми, які науковці вважали

фундаментальними для цієї галузі досліджень: нейронні мережі, теорія обчислюваності, творчість, оброблення мови. На думку авторів, досить детально описати будь-яку особливість розуму людини і передати цю інформацію машині, створеній для його імітації [2, 3].

Наступними десятиліттями штучний інтелект розвивався у напрямках автоматизації, машинного навчання та експертних систем. У 1960–1970-х роках дослідження зосереджувалися на створенні спеціалізованих експертних систем, які давали б рекомендації на основі введених даних. Відомі системи, як-от DENDRAL для хімії або MYCIN для медицини, продемонстрували можливість автоматизації складних інтелектуальних завдань. Однак ці системи були доволі вузькоспеціалізованими і не мали загальних здібностей адаптуватися до нових галузей.

Важливим проривом стало відновлення наукових досліджень у галузі нейронних мереж, зокрема завдяки розробленню методу зворотного поширення помилки для навчання багатошарових нейронних мереж [4]. Це відкрило нові можливості для оброблення великих обсягів даних і вдосконалення алгоритмів розпізнавання образів, що дало поштовх до подальшого розвитку інтелектуальних систем [5].

На початку 2000-х років розпочався стрімкий розвиток обчислювальних потужностей, що разом із доступом до великих масивів даних (Big Data) стимулювало нову хвилю наукових досліджень. Алгоритми машинного навчання і методи глибокого аналізу почали демонструвати значні результати у таких галузях, як розпізнавання зображень, голосу, ігрові симуляції та робототехніка.

Останні десятиліття – етап ключових наукових і технологічних проривів у розвитку штучного інтелекту. Одним із найпоширеніших його застосувань стала автоматизація рутинних завдань. Програмне забезпечення, оснащене ШІ, може виконувати повторювані дії набагато швидше й точніше за людей. Сьогодні нейронні мережі дають змогу проводити аналітичну роботу, виявляти шахрайські дії, оцінювати ризики, оптимізувати витрати тощо. Вони також здатні генерувати реалістичні зображення на основі текстових описів.

Сучасне розуміння штучного інтелекту має два аспекти. Перший – *ідеологічна* складова, що визначає вигляд, у якому має виконуватися той чи інших функціонал. Другий – це платформа, тобто *технічна* складова, що визначає, хто і яким набором інструментів керуватиме процесами ШІ. Людина створює розвиваючі алгоритми, нейронні мережі, а також робототехніку, яка є фізичною оболонкою штучного інтелекту. У зв'язку з цим гостро постає питання відповідального ставлення розробників, які повинні розуміти й передбачити певні обмеження у діяльності ШІ.

Наразі проблема невизначеності поняття «штучний інтелект» суттєво ускладнює розроблення адекватного правового забезпечення його застосування, особливо у такій важливій сфері, як державна безпека. Брак єдиного, загальноприйнятого визначення ШІ створює певні труднощі:

- а) нечіткість сфери регулювання спричиняє прогалини в законодавстві;
- б) ускладнюється класифікація різних систем ШІ за рівнем автономності, що додає труднощів й у розробленні диференційованих правових норм;
- в) правова невизначеність ускладнює питання щодо відповідальності за дії (або бездіяльність) систем штучного інтелекту, особливо у випадках завдання шкоди: хто несе відповідальність – розробник, власник чи користувач;
- г) різні підходи до термінології та визначення поняття «штучний інтелект» у різних юрисдикціях утруднюють міжнародне співробітництво у сфері регулювання штучного інтелекту.

Задля розв'язання зазначених проблем необхідно розробити чітке і збалансоване визначення ШІ, ураховуючи його різні аспекти й можливості. Таке визначення має закладатися в основу правових норм, що регулюють використання ШІ у сфері забезпечення державної безпеки з огляду на специфіку різних галузей застосування.

Розглянемо провідні напрями застосування штучного інтелекту у сфері забезпечення державної безпеки.

1. *Кібербезпека.* Із початком війни в Україні у геометричній прогресії зросли кібератаки й на національну інфраструктуру, і на приватні компанії [6]. Програмні засоби на основі ШІ спроможні ефективно виявляти й нейтралізовувати кібератаки шляхом аналізу величезних обсягів даних у реальному часі. Відомо, що кіберзлочинці навчилися використовувати ШІ для створення шкідливого коду в автоматичному режимі. Попри те, що технологічні гіганти намагаються запобігти використанню своїх мовних моделей для написання шкідливих програм-вірусів, хакери постійно винаходять нові способи обходження обмежень. Цілі платформи (HackedGPT, WormGPT) допомагають створювати шкідливе програмне забезпечення. Це робить динаміку кібератак дуже складною та непередбачуваною.

У перспективі автоматизовані кіберзагрози на базі штучного інтелекту можуть ставати набагато серйознішими. Так, група експертів із кібербезпеки «Нуас» створили «лабораторний» вірус «BlackMamba», який може динамічно змінювати свій код за допомогою ChatGPT. Під час експерименту цей «хробак-кейлоггер» успішно адаптувався під будь-які сценарії та уникав виявлення антивірусним програмним забезпеченням [7]. Важко уявити, наскільки подібні віруси можуть ускладнити роботу фахівців із кіберзахисту і до яких збитків може призвести їх поширення у приватних і державних організаціях.

2. *Штучний інтелект на варті правопорядку та протидії злочинності.* Запровадження штучного інтелекту в юридичних дослідженнях має потенціал для надання допомоги фахівцям галузі державної безпеки і перебудови всієї правоохоронної системи. Як приклад, алгоритми штучного інтелекту спрощують аналіз великих масивів розвідувальних даних, зокрема візуальну й текстову інформацію, що значно прискорює процес прийняття рішень і поліпшує точність прогнозів. Це дає змогу ефективніше реагувати на потенційні загрози і планувати запобіжні заходи. Системи відеоспостереження на основі штучного інтелекту значно розширюють можливості контролю публічних місць (вулиць), державного кордону, автоматично виявляючи підозрілу активність і потенційні порушення. Унікальні алгоритми розпізнавання обличчя людини, моторики тіла (за ходом та будовою) вже допомагають виявляти осіб, які вчинили злочини й перебувають у розшуку [8].

3. *Штучний інтелект на варті протидії дезінформації.* Сучасні IT-технології на базі штучного інтелекту вже виробляють алгоритми, які сприяють виявленню та нейтралізації дезінформаційних кампаній, аналізуючи тексти, зображення і відео в соціальних мережах та інших джерелах. Окремі дослідження свідчать про ефективність цього методу [9, 10].

4. *Штучний інтелект на варті безпеки й оборони.* Підтвердженням важливості використання ШІ для забезпечення національної безпеки й оборони є результати досліджень Науково-технічної організації НАТО, що визначають із них найбільш суттєві для розвитку технологій на найближчі два десятиліття [11]. Ключовими технологіями вважаються, зокрема, такі: BigData, нейронні мережі, автономні транспортні засоби, космос, гіперзвукові літальні апарати, квантові технології, біотехнології, нові матеріали тощо [12]. Такий підхід дає поштовх до переоснащення об'єктів промислового комплексу, кардинальної зміни підходів виробництва зброї і техніки військового призначення.

Збройні сили окремих розвинутих країн світу оснащують свої бойові системи та зброю штучним інтелектом, використовуючи його на повітряних, морських і наземних платформах. Застосування ШІ на бойових платформах сприяло появі ефективних систем ведення війни, повністю незалежних від впливу людини. За цих обставин відбувається підвищення продуктивності платформ, збільшення ефективності та взаємодії між бойовими системами, а також спрощення їх обслуговування.

Роботизовані бойові платформи вперше з'явилися у військах західних армій на стадії дослідної експлуатації на початку 2000 років. Із початком російсько-української війни у 2014 р. і, особливо після повномасштабного російського вторгнення в Україну 2022 р., вони почали активно використовуватися на полі бою [13]. Наразі є такі українські розробки роботів наземних бойових платформ: «Шабля» (наземна платформа, що за своєю сутністю є роботизованою кулеметною туреллю), «Ironclad», «Рись», «Миротворець», «Скорпіон-2, 3» та ін. Ці бойові платформи являють собою самохідні машини, які керуються дистанційно, спроможні самостійно оцінювати обстановку навколо за допомогою відеокамер та програмного забезпечення на базі ШІ і виконувати різні бойові завдання.

Логістика. Штучний інтелект уже відіграє вирішальну роль у логістиці та підтриманні мереж постачання, оскільки визначення альтернативних безпечних маршрутів під час перевезення військових вантажів і бійців є невід'ємним складником реалізації успішної військової операції. Інтеграція ШІ у транспорт і логістичні ланцюжки зменшує витрати на допоміжний персонал, викорінює корупцію серед постачальників, мінімізує фактор людських помилок. У транспортній авіації та на флоті ШІ також уможливує швидке виявлення поломок компонентів у техніці, що безпосередньо запобігає їх аваріям чи катастрофам.

Виявлення і розпізнавання цілей. У складних бойових умовах штучний інтелект підвищує точність розпізнавання невідомих цілей. Це дає змогу силам оборони/наступу досягати повного розуміння у потенційному полі діяльності завдяки глибокому аналізу документів, звітів, карт, новин та іншої важливої інформації, що допомагає у визначенні потенційної цілі. Системи розпізнавання цілей на базі ШІ мають кілька особливих переваг: прогнозування поведінки ворога; оцінювання ландшафту місцевості й усіх невідомих об'єктів, розміщених на території; визначення наслідків для довкілля від застосування тієї чи іншої зброї. Так, на початку 2020 р. компанія «Рейтер» оголосила про розгортання системи розвідки, спостереження й цілевказання «Істар» на літаках військово-повітряних

сил Великої Британії «Сентініел». Ця система заснована на ШІ й забезпечує знаходження наземних та морських об'єктів і спостереження за їх пересуванням [14].

Допомога бойовим медикам. Роботизовані наземні платформи, оснащені ШІ, використовуються для надання невідкладної допомоги пораненим бійцям у зоні активних бойових дій для подальшої їх евакуації з небезпечної зони. Такі платформи швидко діагностують пораненого бійця, визначають рівень його ураження, рекомендують методи лікування, що допомагає бойовим медикам у польових умовах приймати правильні рішення у стресових ситуаціях [15, 16]. Слід зазначити, що незалежно від розташування лінії зіткнення ШІ має доступ до бази даних лікарень і медичних книжок військових, що дає змогу алгоритмам виявляти й сортувати проблеми зі здоров'ям у військових.

Моніторинг потенційних загроз і ситуаційний аналіз. Ці операції використовуються для оброблення інформації задля підтримання широкого спектра військових дій. Штучний інтелект допомагає швидко обробляти великі обсяги даних одночасно з кількох серверів, підсумовує проаналізовану інформацію та пропонує готові висновки. Такий аналіз дає змогу військовим визначати закономірності, узагальнювати висновки і в деяких випадках вносити свої корективи. Наприклад, завдяки безпілотним системам, оснащеним ШІ, військові можуть контролювати великі території протягом тривалого часу, оцінювати скупчення ворога, ворожої техніки, його переміщення та швидко передавати інформацію.

Симуляція і навчання – ще одна вкрай корисна сфера застосування інструментів ШІ, яка поєднує системну інженерію, розроблення програмного забезпечення та інформатику для створення комп'ютеризованих моделей. Дедалі більше країн інвестують у програми для моделювання ситуацій і навчання військових на основі сучасного досвіду. Відпрацювання на тренажерах різноманітних ситуацій дає змогу мінімізувати втрати на полі бою, урахувати людські помилки (прорахунки) і сприяє розробленню заходів безпеки [17].

Як бачимо, штучний інтелект – дієвий інструмент, здатний значно підвищити рівень багатьох складових державної безпеки України. Однак його застосування пов'язане з певними ризиками, зокрема з можливістю зловживань і залежністю від високих технологій. Тому важливим завданням вбачається розроблення і впровадження ефективних стратегій управління ризиками, створення етичних норм використання штучного інтелекту та забезпечення на його основі систем кібербезпеки. Використання ШІ у правоохоронній діяльності актуалізує питання дотримання прав людини, можливості упереджених підходів і дискримінації. Автоматизація рішень на основі ШІ небезпечна через втрату людського контролю та прийняття необ'єктивних рішень. Слід також зазначити, що широке впровадження інструментів ШІ сприяє автоматизації праці й може призвести до зростання безробіття.

Тому паралельно з розвитком технологій необхідно приділяти максимальну увагу розробленню ефективних механізмів правового регулювання та механізмів контролю задля мінімізації ризиків і максимізації позитивного впливу ШІ на суспільство й державу. Забезпечення балансу між швидким розвитком технологій і забезпеченням безпеки є ключовим завданням для сучасного суспільства.

Аналіз вивчених джерел, вітчизняного й зарубіжного законодавства дає змогу дійти висновку, що наразі нема єдиного підходу до визначення штучного інтелекту та особливостей нормативно-правового регулювання його застосування. Проте вже зроблено деякі кроки на шляху розв'язання зазначених проблем як на національному, так і міжнародному рівнях. Протягом останніх років багатьма державами розроблено довгострокові національні стратегії розвитку штучного інтелекту і здійснено певні заходи щодо їх упровадження. Стратегії розвитку ШІ різних країн світу можна умовно поділити на такі основні групи [18].

1. Країни (США, Велика Британія, Німеччина, Саудівська Аравія тощо), яких відрізняє реалістичне ставлення до формування стратегій ШІ, глибокий аналіз не тільки стану сфери застосування ШІ у країні, але й дійсних потреб її розвитку. Стратегії цих країн мають фундаментальний характер і відображають як загальні світові проблеми впровадження ШІ, так і конкретні плани цифровізації багатьох галузей національної економіки й різних сфер суспільних відносин.

2. Група країн, які характеризують ґрунтовний і прагматичний підхід до цілей та етапів їх досягнення з урахуванням дійсних потреб держави та формування окремих унікальних завдань і цілей розвитку ШІ: Велике князівство Люксембург, Малайзія, Литва.

3. Країни, стратегії яких виконано у формалізованому вигляді, визначають базові цілі розвитку країни в напрямі впровадження технологій із ШІ у певних сферах суспільної життєдіяльності: Австралія, Республіка Австрія, Королівство Іспанія, Катар, Португалія, Республіка Ізраїль, Швейцарська Конфедерація тощо [18].

В Україні прийнято документи, які визначають стратегію використання технологій штучного інтелекту у сфері національної безпеки й оборони [19–24]. Серед основних можна виділити такі.

1. Стратегія національної безпеки України, згідно з якою розробляються системи озброєнь на основі технологій у сфері ШІ, створюються нові матеріали, робототехніка та автономні безпілотні апарати.

2. Стратегія інформаційної безпеки і Стратегія кібербезпеки України, якими передбачено запобігання гібридним загрозам у вигляді дезінформації та інформаційних операцій, захист персональних даних тощо.

3. Стратегія розвитку оборонно-промислового комплексу України, у якій прямо зазначається про широке застосування новітніх технологій на виробництві, зокрема із застосуванням ШІ.

Кабінет Міністрів України затвердив 12 травня 2021 р. «План заходів з реалізації Концепції розвитку штучного інтелекту в Україні на 2021–2024 рр.», згідно з яким на зазначений період передбачалося розробити низку заходів і законодавчих ініціатив, зокрема: запровадження правового регулювання з питань формування державної політики в галузі ШІ; запровадження державної підтримки використання технологій ШІ у пріоритетних галузях економіки; упровадження технологій ШІ у національну систему кібербезпеки для проведення аналізу і класифікації загроз та вибору стратегії їх стримування й запобігання їх виникненню; визначення пріоритетних напрямів і основних завдань розвитку технологій ШІ у документах оборонного планування тощо [25].

Аналіз чинного законодавства України свідчить про наявність певних ініціатив щодо регулювання застосування штучного інтелекту у сфері державної безпеки й оборони, які ґрунтуються на нагальних потребах. Проте бракує єдиного стратегічного документа, наприклад, Національної стратегії розвитку ШІ у сфері забезпечення національної безпеки і оборони. Саме це гальмує системне впровадження цих технологій. Така ситуація стимулює активне обговорення та наукові дослідження з приводу перспектив використання ШІ в оборонному комплексі України. Актуалізуються питання визначення перспективних напрямів застосування ШІ, аналізу його впливу на оборонно-промисловий комплекс та оцінювання досвіду інших країн щодо цього.

США, Європейський Союз (ЄС) і Китай активно розробляють власні стратегії і нормативно-правову базу для регулювання ШІ, водночас як інші держави можуть не мати достатньо ресурсів або технологічної готовності для розроблення та імплементації ефективних регуляторних механізмів. Це створює підґрунтя для геополітичних та економічних розбіжностей у ставленні до технологій, зокрема щодо того, як саме застосовувати етичні принципи, як захищати приватність даних і яким чином запобігти зловживанню ШІ у сферах безпеки чи інтелектуальної власності.

Європейський Союз можна вважати лідером на міжнародному рівні в аспекті регулювання ШІ. У 2021 р. ЄС представив проєкт «Акт про ШІ», що передбачає створення єдиного правового поля для застосування ШІ в межах Європи [26]. Цей документ пропонує системний підхід до ризиків, пов'язаних із використанням ШІ, зокрема визначає різні категорії ризиків. Так, європейська нормативна база передбачає більш суворий контроль за високоризиковим застосуванням ШІ (системи для автоматизованих рішень у судочинстві або застосування ШІ в робототехніці), тоді як для сфер менш ризикових пропонуються ліберальніші підходи.

США також активно працюють над створенням національної стратегії в галузі ШІ [27]. Проте їхній підхід зосереджено більше на сприянні інноваціям, розвитку технологій і підтримці приватного сектору, що робить США однією з найбільш інноваційних країн у галузі. Однак менша увага до регулювання й контролю у майбутньому може призвести до певних проблем з етикою та приватністю. Закон про Національну ініціативу штучного інтелекту від 2020 р. (National Artificial Intelligence Initiative Act of 2020) заохочує та субсидує інноваційні зусилля зі штучного інтелекту у ключових федеральних агентствах США. Проєкт білля про права на штучний інтелект, представлений у 2022 р., містить набір керівних, необов'язкових принципів для безпечного й надійного розвитку штучного інтелекту [28]. Порівняно з Європейським Союзом США менш орієнтовані на регуляцію використання ШІ, хоча деякі штати (наприклад Каліфорнія) вже почали впроваджувати власні закони, спрямовані на захист приватності і безпеки даних.

Китай – ще один важливий гравець на міжнародній арені, особливо в контексті розвитку ШІ. Він активно інвестує в дослідження та розробки в цій галузі, маючи амбітні плани до 2030 р. стати світовим лідером у використанні штучного інтелекту. Водночас Китай має інший підхід до регулювання (Адміністрація кіберпростору Китаю / Cyberspace Administration of China – CAC), орієнтуючись на сильний державний контроль і виявляючи меншу увагу до прав людини та етичних питань порівняно з Європою чи США. Це створює виклики для міжнародної співпраці, оскільки розбіжності у підходах можуть спричинити труднощі щодо глобального управління технологіями та

їхнього впливу на суспільство. Нові китайські правила стосуватимуться лише тих сервісів генеративного ШІ, які доступні широкому загалу, а не тих, що розробляються, наприклад, у науково-дослідних установах. Одна з вимог – сервіси генеративного ШІ змушені будуть отримати ліцензію на роботу. Якщо постачальник послуг генеративного ШІ виявляє «незаконний» контент, він повинен ужити заходів для припинення генерації такого контенту, вдосконалити алгоритм, а потім повідомити про це відповідний орган.

ООН та інші міжнародні організації також займаються питанням регулювання штучного інтелекту на глобальному рівні. Так, створена ООН спеціальна комісія з етики в галузі ШІ займається дослідженням етичних, соціальних і правових аспектів використання цих технологій. У 2021 р. була розроблена й опублікована «Глобальна етика використання штучного інтелекту», яка покликана допомогти урядам та іншим зацікавленим сторонам у визначенні принципів для справедливого використання ШІ. Вона охоплює питання етики, прозорості, підзвітності та захисту прав людини. У цьому документі, зокрема, зазначено, що структури Організації Об'єднаних Націй повинні забезпечити, аби системи штучного інтелекту не переважали над свободою та автономією людини у прийнятті рішень, а також гарантувати нагляд з боку людини. Усі етапи життєвого циклу системи штучного інтелекту мають передбачати людино-центричні практики проєктування та залишати значущі можливості для прийняття рішень людиною [30].

Здійснюється також активна співпраця між національними урядами і приватними технологічними компаніями. Наприклад, такі великі технологічні корпорації, як Google, Microsoft, IBM, беруть активну участь у розробленні політик щодо регулювання ШІ [31, 32]. Вони створюють власні етичні кодекси і працюють над удосконаленням технологій для забезпечення їх безпечного використання. Однак співробітництво зазнає труднощів через недостатній рівень прозорості у процесі розроблення політик, а також суперечність між прагненням компаній розвивати технології для комерційних цілей і необхідністю забезпечувати їх безпеку та етичність.

Глобальні ініціативи щодо регулювання ШІ досі ще на стадії розвитку. Це одна з головних причин, чому важливо створювати міжнародні платформи для обміну досвідом і розроблення спільних стандартів. Технічні стандарти, які враховують як інноваційний розвиток, так і потребу у безпечному й етичному застосуванні, можуть стати підґрунтям для ефективного міжнародного співробітництва. Проблеми міжнародного співробітництва у регулюванні ШІ також охоплюють питання глобальної конкуренції та ризиків, пов'язаних із технологічною гонкою. Країни можуть мати різні пріоритети в розробленні ШІ – від підвищення національної безпеки до комерційних та економічних інтересів. Це може призвести до нерівного доступу до технологій і впливу ШІ на різні країни та регіони.

Інша важлива проблема, яку породжує брак чіткої дефініції, полягає в тому, що законодавчі органи не можуть створювати відповідну інфраструктуру для моніторингу й оцінювання розвитку технологій ШІ. Через швидкий розвиток технологій ШІ актуалізується питання визначення авторських прав на об'єкти, створені штучним інтелектом. Хто є автором такого твору – людина, яка створила алгоритм, чи сам ШІ? Уже були випадки, коли суди в різних країнах розглядали питання авторських прав на твори, створені ШІ. У США, наприклад, авторські права не можуть бути надані ШІ, оскільки він не є фізичною особою [33].

Питання відповідальності за дії штучного інтелекту також становить одну з найбільш складних і важливих правових проблем у контексті сучасних технологій. Завдяки своїй здатності до автономного прийняття рішень ШІ значно ускладнює традиційне розуміння відповідальності, яке у праві зазвичай покладається на конкретну фізичну або юридичну особу. Технології, здатні до самонавчання та автономної взаємодії з оточенням, перед правовими системами ставлять питання, на яке складно відповісти в рамках існуючих норм.

У звичайних ситуаціях, коли технології не мають високого рівня автономії, відповідальність за їх дії зазвичай покладається на їхніх творців – розробників або власників. Коли ж припускається помилки система (наприклад робот), відповідальність покладається на компанію, що його виготовила, або на власника. Однак щодо більш автономних систем, які здатні самі приймати рішення і здійснювати дії без втручання людини, традиційне трактування відповідальності не є таким очевидним. Зокрема, коли йдеться про ситуації, в яких ШІ завдає шкоди. Так, у разі аварії з участю автономного автомобіля важливо визначити, хто відповідальний за наслідки: розробник програмного забезпечення, власник системи чи безпосередньо ШІ.

Один із можливих підходів полягає в тому, що відповідальність за дії ШІ має залишатися на розробниках і власниках технології, навіть якщо система діє автономно. Це означає, що розробники, які створили алгоритм або модель ШІ, повинні нести юридичну відповідальність за те, як ці системи працюють. До власників технології можуть також застосовуватися санкції, якщо вони не забезпечили

належного рівня контролю за роботою системи. У такому разі це можна розглядати як відповідальність за недогляд або неналежне використання технології. Важливими вбачається також тестування і перевірка систем на наявність помилок перед їх надходженням для використання. І це як обов'язок має покладатися на розробників і власників.

З іншого боку, у випадку із системами, здатними до самонавчання (наприклад, нейронні мережі, які адаптуються до нових даних), складно точно передбачити, як система поводитиметься в тій чи іншій ситуації. Це створює додаткові труднощі для правового регулювання, оскільки помилку може спричинити не дефект алгоритму, а неочікувані результати самонавчання системи. Оскільки це явище ще не набуло чітких правових визначень, правова система має пристосуватися до нових викликів, які створюють такі технології. Одна з ідей, що обговорюється у колах юристів і технологів, полягає в тому, щоб на певних етапах використання ШІ обмежити відповідальність розробників у випадках неправильного застосування системи користувачами [34, 35].

Одним із цікавих аспектів є пропозиції щодо можливої відповідальності самого ШІ [18, 36]. На сьогодні жодна правова система не розглядає ШІ як юридичну особу, здатну нести відповідальність за свої дії. Однак із розвитком автономних систем і можливістю самостійного прийняття рішень, імовірно, у майбутньому постане необхідність створення нових правових категорій для регулювання цього аспекту. Це може бути розвиток концепції «цифрових осіб» або створення нових інститутів права, здатних працювати в контексті дії технологій, що взаємодіють зі світом на автономному рівні [37, 38, 39].

Отже, питання організаційно-правового регулювання дії ШІ відкриває цілу низку юридичних, етичних і технічних проблем, які вимагають чітких відповідей від правових систем по всьому світу. Розв'язання зазначених проблем потребує комплексного підходу, що має враховувати як інтереси розвитку технологій, так і захист прав осіб, що можуть постраждати від неконтрольованих дій автономних систем.

Висновки

Штучний інтелект має значний вплив на національну безпеку, оскільки може бути використаний як у позитивному, так і в негативному контекстах для зміцнення або ослаблення безпеки держави. У військових, розвідувальних, правоохоронних та інших державних структурах штучний інтелект уже активно застосовується для поліпшення ефективності управління, аналізу даних і виявлення загроз. З одного боку, це відкриває нові можливості, а з іншого – ставить держави й суспільства перед новими викликами, пов'язаними з етикою, контролем і потенційними ризиками.

Один із найбільш очевидних напрямів застосування штучного інтелекту у сфері державної безпеки – це оборона. Штучний інтелект допомагає модернізувати оборонні технології, зокрема у розробленні таких автономних систем, як безпілотні літальні апарати (БПЛА), роботизовані бойові одиниці, а також інтелектуальні системи для управління військовими операціями. Вони здатні швидко обробляти величезні обсяги інформації і приймати рішення у реальному часі, що дає змогу здійснювати більш точні та ефективні операції в умовах складних бойових ситуацій. Штучний інтелект також застосовується в контексті кібербезпеки для виявлення і нейтралізації кібератак, а також для прогнозування можливих загроз і запобігання ним на етапі формування.

Використання штучного інтелекту в контексті розвідки та збирання інформації є ще одним із важливих складників. Алгоритми машинного навчання і аналізу великих обсягів даних дають змогу розпізнавати патерни, які можуть указувати на підготовку до терористичних актів, військових конфліктів чи інші загрози.

Отже, штучний інтелект має величезний потенціал для зміцнення державної безпеки, але водночас виникають і нові виклики, які потребують обережного й відповідального підходу до регулювання, етики та захисту прав людини. Використання цієї технології має бути збалансованим і забезпечувати прозорість у процесах, де вона може впливати на громадянські свободи і стабільність держави.

Подальшими напрямами наукових досліджень вбачається вирішення питань правового забезпечення функціонування штучного інтелекту з урахуванням вимог національної безпеки. Науковому аналізу мають бути піддані етичні правила і законодавчі ініціативи, які б давали змогу ефективно використовувати штучний інтелект, зберігаючи при цьому баланс між безпекою і правами громадян. Це передбачає встановлення чітких стандартів щодо моніторингу й контролю, захисту даних, а також збереження демократичних принципів у процесах, де штучний інтелект має значний вплив на управління державою.

Перелік джерел посилання

1. A. Turing. Computing Machinery and Intelligence. *Mind*, Volume LIX, Issue 236, October 1950. Pp. 433–460. DOI: <https://doi.org/10.1093/mind/LIX.236.433>.
2. Джон Маккарті — «батько» штучного інтелекту та хмарних обчислень. *GIGACLOUD.UA*. URL: <https://surl.lu/zmcsep> (дата звернення: 03.11.2024).
3. The Dartmouth Artificial Intelligence Conference: The next 50 years. *Dartmouth News Press Release*, July 6, 2006. URL: <https://surl.li/lnntkd> (accessed: 02.11.2024).
4. Fetzer J. H. What is Artificial Intelligence? Its Scope and Limits. Studies in Cognitive Systems. *The Springer*. 1990. Vol 4 : Dordrecht. DOI: https://doi.org/10.1007/978-94-009-1900-6_1.
5. Hunt E. B. Artificial Intelligence. London, New York : Academic press, inc., 1975. 361 p.
6. Держкомспецзв'язку повідомив про збільшення кількості кібератак в Україні за минулий рік. *Слово і Діло*. URL: <https://surl.li/slffso> (дата звернення: 20.10.2024).
7. Платформа розвідки та розслідування кіберзагроз. *HYAS Insight*. URL: <https://surl.lu/dioexw> (дата звернення: 14.11.2024).
8. AI People, Vehicle, Face Detection. *i-PRO Expert platform*. URL: <https://surl.lu/rxbdiip> (дата звернення: 02.11.2024).
9. Штучний інтелект і дезінформація: викриття цифрової пропаганди. *Національний проєкт з медіа грамотності «Фільтр»*. URL: <https://surl.li/nzsoaf> (дата звернення: 02.11.2024).
10. Петрів О. Дезінформація та штучний інтелект: (не)видима загроза сучасності. *Центр демократії та верховенства права*. URL: <https://surl.li/hwjke> (дата звернення: 15.11.2024).
11. Баловсяк Н. В. Кібератаки з використанням штучного інтелекту, за оцінкою НАТО, є критичною загрозою. *Інтернет Свобода*. URL: <https://surl.li/mshnlo> (дата звернення: 02.02.2025).
12. Хаустова В. Є., Решетняк О. І., Хаустов М. М., Зінченко В. А. Напрямки розвитку технологій штучного інтелекту в забезпеченні обороноздатності країни. *Бізнесінформ*. 2022. № 3. С. 17–26.
13. Наземні бойові платформи: новий гравець на полі бою. *Militaryni*. URL: <https://surl.lu/rjcxgo> (дата звернення: 12.02.2025).
14. Пацурія Н. Впровадження технологій штучного інтелекту у забезпечення національної безпеки та обороноздатності України: проблеми та перспективи повоєнного періоду. *Координата*. URL: <https://surl.li/ynkqck> (дата звернення: 12.02.2025).
15. D. Quick. Battlefield Extraction-Assist Robot to ferry wounded to safety. *New Atlas*. 2023. Vol. 2. URL: <https://surl.li/kuqfoq> (accessed: 12.02.2025).
16. В Україні випробували роботизований комплекс THeMIS для евакуації поранених. *Militaryni*. URL : <https://surl.li/lylylo> (дата звернення: 12.02.2025).
17. Rademacher, T. Artificial Intelligence and Law Enforcement. In: Wischmeyer, T., Rademacher, T. (eds) *Regulating Artificial Intelligence*. *The Springer*. 2020. Vol. 5. DOI : https://doi.org/10.1007/978-3-030-32361-5_10.
18. Костенко О. В. Аналіз національних стратегій розвитку штучного інтелекту. *Інформація і право*. 2022. № 2 (41). С. 79–76.
19. Про рішення Ради національної безпеки і оборони України «Про Стратегію забезпечення державної безпеки» : Указ Президента України від 16.02.2022 р. № 56/2022. URL: <https://surl.lu/mpuima> (дата звернення: 13.02.2025).
20. Про рішення Ради національної безпеки і оборони України «Про Стратегію національної безпеки України» : Указ Президента України від 14.09.2020 р. № 392/2020. URL: <https://zakon.rada.gov.ua/laws/show/392/2020#Text> (дата звернення: 13.02.2025).
21. Про рішення Ради національної безпеки і оборони України «Про Стратегію інформаційної безпеки» : Указ Президента України від 28.12.2021 р. № 685/2021. URL: <https://surl.lu/pbsxaa> (дата звернення: 13.02.2025).
22. Про рішення Ради національної безпеки і оборони України «Про Стратегію кібербезпеки України» : Указ Президента України від 26.08.2021 р. № 447/2021. URL: <https://surl.li/cc/ryvkmt> (дата звернення: 13.02.2025).
23. Про рішення Ради національної безпеки і оборони України «Про Стратегію воєнної безпеки України» : Указ Президента України від 25.03.2021 р. № 121/2021. URL: <https://surl.li/kkilrt> (дата звернення: 13.02.2025).
24. Про рішення Ради національної безпеки і оборони України «Стратегія розвитку оборонно-промислового комплексу України» : Указ Президента України від 20.08.2021 р. № 372/2021. URL: <https://surl.li/ikvfcg> (дата звернення: 13.02.2025).

25. Про затвердження плану заходів з реалізації Концепції розвитку штучного інтелекту в Україні на 2021–2024 роки : розпорядження Кабінету Міністрів України від 12.05.2021 р. № 438-р. URL: <https://surl.li/ikvfcb> (дата звернення: 13.02.2025).
26. The EU Act of AI: Directive (EU) 2020/1828 (Artificial Intelligence Act). URL: [EUAIAct.com](https://european-council.europa.eu/media/eu-act-ai) (accessed: 12.02.2025).
27. National Artificial Intelligence Initiative Act of 2020. *US Congress summaries*. URL: <https://surl.li/txefrv> (accessed: 12.02.2025).
28. Blueprint for an AI Bill of Rights. *The White House website*. URL: <https://surl.li/rquezm> (accessed: 12.02.2025).
29. Баловсяк Н. У Китаї затверджені правила регулювання штучного інтелекту. *Інтернет Свобода*. URL: <https://surl.li/snhcig> (дата звернення: 12.02.2025).
30. UNESCO's Ethics of AI Recommendation. URL: <https://surl.li/jonfrw> (дата звернення: 12.02.2025).
31. Making AI helpful for everyone. *Google inc.* URL: <https://ai.google/> (accessed: 12.02.2025).
32. Artificial Intelligence (AI) Solutions. *IBM inc.* URL: <https://surl.li/sdsvqp> (accessed: 12.02.2025).
33. Арсенів М. В. США вирішили, що штучний інтелект не може мати авторських прав на свої продукти. *Судово-юридична газета*. URL: <https://surl.li/fazoyd> (дата звернення: 12.02.2025).
34. Четверта промислова революція: зміна напрямів міжнародних інвестиційних потоків : монографія / А. І. Крисоватий та ін.; за наук. ред. А. І. Крисоватого та О. М. Сохацької. Тернопіль : Осадца Ю. В., 2018. 480 с. URL: <https://surl.li/lfcuhy> (дата звернення: 12.02.2025).
35. Шишка Н. В. Штучний інтелект в українському правосудді: правові передумови запровадження. *Юридичний науковий електронний журнал*. DOI: <https://doi.org/10.32782/2524-0374/2021-3/35>.
36. Радутний О. Е. Суб'єктність штучного інтелекту у кримінальному праві. *Право України*. 2018. № 1. С. 123–136.
37. White paper. On Artificial Intelligence – A European approach to excellence and trust. European Commission. Brussels, 19.02.2020 COM(2020) 65 final. URL: <https://surl.li/hxgqhm> (accessed: 12.02.2025).
38. Павленко Ж. О., Водорезова С. Р. Поняття електронної особи в цифровій реальності. *Вісник НЮУ імені Ярослава Мудрого. Філософія, філософія права, політологія, соціологія*. 2021. Т. 3. № 50. С. 59–70.
39. Колодін Д. О., Байталюк Д. Р. Щодо питання цивільно-правової відповідальності за шкоду, завдану роботизованими механізмами зі штучним інтелектом (роботами). *Часопис цивілістики*. 2019. № 33. С. 87–91.

Стаття надійшла до редакції 02.03.2025

UDC 342.9, 351, 004.8

V. Korniienko, V. Hmyria, M. Styshuk

ORGANIZATIONAL AND LEGAL PRINCIPLES OF USING ARTIFICIAL INTELLIGENCE IN THE SPHERE OF PROVIDING STATE SECURITY

The article analyzes the use of artificial intelligence in the field of ensuring state security, including in law enforcement activities. The work contains an analysis of existing legislation, identifies the risks of using artificial intelligence, and offers substantiated recommendations for improving the regulatory framework, in particular, regarding liability, personal data protection, cybersecurity, and ethical use of artificial intelligence tools.

Keywords: artificial intelligence, neural networks, service and combat activities of the National Guard of Ukraine, state security, legislative support, cybersecurity.

Корнієнко Василь Володимирович – кандидат юридичних наук, доцент, старший викладач кафедри правових дисциплін, Національна академія Національної гвардії України
<http://orcid.org/0000-0002-7682-1281>

Гмиря Вікторія Петрівна – кандидат економічних наук, доцент, провідний науковий співробітник науково-організаційного відділу, Державний науково-дослідний інститут випробувань і сертифікації озброєння та військової техніки
<https://orcid.org/0000-0003-3070-0158>

Стишук Михайло Сергійович – курсант групи 283 командно-штабного факультету Національної академії Національної гвардії України
<https://orcid.org/0009-0000-7459-2078>