

МІЖНАРОДНИЙ ДОСВІД РЕГУЛЮВАННЯ ШТУЧНОГО ІНТЕЛЕКТУ У СФЕРІ БЕЗПЕКИ

Моргунова Тетяна Іванівна

*к.т.н., доцент, доцент кафедри
кримінального аналізу та інформаційних технологій
Одеського державного університету внутрішніх справ
<https://orcid.org/0000-0002-3512-2425>*

М'ясоєдова Анастасія Григорівна

*здобувач першого (бакалаврського)
рівня вищої освіти спеціальності 081 «Право»
Одеського державного університету внутрішніх справ*

Розвиток цифрової економіки супроводжується інтенсивним впровадженням передових цифрових технологій, зокрема штучного інтелекту та аналітики великих даних (Big Data), що спричиняє глибокі трансформаційні зміни в різних галузях [1].

Сучасний етап розвитку цифрових технологій супроводжується стрімким впровадженням систем штучного інтелекту (ШІ) у сферу безпеки, що охоплює як державні, так і приватні структури.

Штучний інтелект не лише моделює людське мислення та аналізує великі обсяги даних, а й активно впроваджується в різні сфери діяльності. Зокрема, він забезпечує ідентифікацію ризиків і шахрайських схем, оптимізацію портфелів інвестицій та автоматизацію процесів прийняття рішень на основі машинного навчання [2].

Використання різних інструментів на базі штучного інтелекту, включно з машинним навчанням, розпізнаванням обличчя чи аналітикою прогнозів, вже давно перестало бути чимось суто експериментальним. Це поступово змінює підходи до безпеки – і громадської, і національної, і навіть кіберпростору. Деякі рішення ухвалюються швидше, ніж раніше, інколи навіть занадто швидко, і не завжди зрозуміло, як саме вони були сформовані. Тут і виникають певні сумніви, бо поряд із зручністю з'являються ризики – від порушення приватності до не зовсім прозорих алгоритмів, які працюють ніби «всередині чорної скриньки». Власне, через це тема регулювання штучного інтелекту на міжнародному рівні почала звучати значно частіше, ніж ще кілька років тому.

Якщо подивитись на різні країни, стає помітно, що єдиного підходу тут так і не сформувалося. Кожен рухається у свій спосіб, іноді

навіть досить суперечливо. Це залежить і від економічних інтересів, і від політичних традицій, і від того, наскільки взагалі країна готова до таких технологій. Хоча в цілому проглядається спільна логіка – знайти якийсь баланс між розвитком і безпекою, щоб не загальмувати інновації, але й не втратити контроль над ситуацією.

Європейський варіант виглядає більш формалізованим, якщо так можна сказати. Тут робиться акцент на правах людини, і це відчувається майже у всіх документах. Системи ШІ оцінюються через призму ризиків, і залежно від цього до них застосовуються різні вимоги. З одного боку, це створює певну чіткість, але з іншого – іноді виглядає надто складно і бюрократично. Особливо коли мова заходить про такі речі, як прозорість алгоритмів або відповідальність розробників. У сфері безпеки це означає обережне ставлення до масового спостереження, хоча на практиці все не завжди так однозначно.

У Сполучених Штатах підхід інший, більш гнучкий, навіть трохи хаотичний на перший погляд. Там більше покладаються на саморегулювання і галузеві стандарти, що дає змогу швидше впроваджувати нові рішення. У сфері безпеки технології ШІ використовуються досить активно, зокрема правоохоронними структурами. Але водночас виникають дискусії про те, де проходить межа між ефективністю і втручанням у приватне життя. І ці дискусії часом заходять у глухий кут.

Китайська модель виглядає ще інакше. Там держава відіграє ключову роль, фактично задає напрям розвитку і контролює його. Системи відеоспостереження, соціальні рейтинги – усе це вже стало частиною повсякденності. Такий підхід дозволяє досить ефективно управляти процесами, хоча питання приватності тут відходить на другий план. І це викликає неоднозначну реакцію, особливо з боку інших країн.

Не можна обійти увагою і міжнародні організації, які намагаються хоча б частково узгодити підходи. Наприклад, Організація економічного співробітництва та розвитку пропонує певні принципи відповідального використання ШІ – прозорість, підзвітність, безпека. Подібні ініціативи з'являються і в рамках Організації Об'єднаних Націй, де більше уваги приділяється етичним аспектам. Хоча іноді ці документи виглядають більше як рекомендації, ніж реальні механізми впливу.

Окремо варто згадати військову сферу, де штучний інтелект починає відігравати все помітнішу роль. Автономні системи викликають багато суперечок, особливо через питання контролю. Ідея збереження людського фактора звучить логічно, але реалізувати її не так просто. Міжнародні обговорення тривають, але чіткої позиції досі немає, і це, мабуть, теж показник того, наскільки складною є ця тема. Іноді здається, що технології рухаються швидше, ніж здатність суспільства домовитися про правила їх використання.

Аналізуючи різні підходи до регулювання штучного інтелекту, доцільно узагальнити їх ключові характеристики, що відображено у табл. 1.

Таблиця 1

Порівняльна характеристика моделей регулювання штучного інтелекту у сфері безпеки

Країна/регіон	Основний підхід	Рівень державного контролю	Особливості у сфері безпеки
ЄС	Ризик-орієнтований, нормативний	Високий	Обмеження масового спостереження, захист прав людини
США	Гнучкий, ринковий	Середній	Активне використання ШІ правоохоронними органами
Китай	Централізований, державний	Дуже високий	Масштабні системи цифрового контролю
ОЕСР / ООН	Рекомендаційний	Низький	Формування етичних стандартів та принципів

Якщо дивитись на це трохи простіше, то видно, що жодна модель не є ідеальною сама по собі. У кожної є свої сильні сторони, але й свої слабкі місця, і вони якимось природно впливають із того, що саме для держави важливіше в конкретний момент. Європейський варіант більше «про правила» і захист прав, інколи навіть занадто обережний, через що розвиток технологій може трохи гальмуватися. Американський підхід виглядає більш живим, швидшим, але при цьому не завжди встигає за власними наслідками,

тому питання контролю постійно повертається в дискусіях. Китайська модель, навпаки, показує ефективність у сенсі організації й безпеки, хоча ціна цього іноді здається занадто високою, особливо якщо говорити про свободи.

Для України ця тема виглядає трохи інакше, бо контекст все ж складніший, і безпековий фактор тут відчувається значно сильніше. Просто взяти одну з моделей і скопіювати її навряд чи спрацює. Швидше, мова йде про певне поєднання різних підходів, навіть якщо це звучить не дуже чітко. Логічно виглядає варіант, де європейська ідея з оцінкою ризиків якось поєднується з американською гнучкістю, плюс додається більш стратегічне бачення, яке часто згадують у контексті Китаю. Хоча на практиці це, напевно, буде виглядати зовсім не так акуратно, як у теорії.

При цьому ключові речі залишаються доволі очевидними, навіть трохи банальними: безпека, права людей і розвиток технологій. Інша справа, що між ними постійно виникає напруга, і баланс тримати складніше, ніж це зазвичай описують у підручниках. Іноді здається, що рішення приймаються вже «по ситуації», а не за якоюсь чіткою моделлю, і це, мабуть, теж частина реальності.

Важливо враховувати, що сучасний міжнародний досвід регулювання штучного інтелекту у сфері безпеки поступово переходить від фрагментарних рішень до формування комплексних підходів, які поєднують правові, технологічні та інституційні механізми. У цьому контексті дедалі більшого значення набуває гармонізація нормативно-правових вимог між країнами, зокрема щодо використання ШІ в правоохоронній діяльності, кібербезпеці та оборонному секторі. Практика свідчить, що ефективне регулювання неможливе без встановлення чітких процедур оцінки ризиків, визначення меж допустимого застосування автономних систем і забезпечення прозорості алгоритмів. Крім того, міжнародна координація дозволяє уникнути регуляторних прогалин, які можуть бути використані для зловживань, і водночас сприяє формуванню спільного безпекового простору, в якому інновації розвиваються без критичних загроз для суспільства.

Таким чином, міжнародний досвід регулювання штучного інтелекту у сфері безпеки демонструє різноманітність підходів і необхідність їх адаптації до національних умов. Формування ефективної системи регулювання потребує комплексного врахування технологічних, правових та етичних аспектів, а також активної

участі держави, бізнесу та громадянського суспільства. У перспективі подальший розвиток міжнародного співробітництва у цій сфері може сприяти створенню узгоджених стандартів, які забезпечать безпечне та відповідальне використання штучного інтелекту на глобальному рівні.

Список використаних джерел:

1. Ying J., Wan Y. The development trend of artificial intelligence in the big data environment. *2022 3rd International Conference on Electronic Communication and Artificial Intelligence (IWEC AI)*. Zhuhai, China, 2022. <https://doi.org/10.1109/IWEC AI55315.2022.00064>

2. Цегельник Н.І., Гладков Д.А. Big Data та штучний інтелект в обліку інвестицій: вплив на економічну безпеку підприємств. *Актуальні питання економічних наук*. 2025. № 9. DOI: <https://doi.org/10.5281/zenodo.15108352>

**ЗАСТОСУВАННЯ ТЕХНОЛОГІЙ ШТУЧНОГО
ІНТЕЛЕКТУ В КРИМІНАЛЬНОМУ АНАЛІЗІ
ТА КІБЕРБЕЗПЕЦІ У КОНТЕКСТІ
ОПЕРАТИВНО-РОЗШУКОВОЇ ДІЯЛЬНОСТІ:
АНАЛІТИЧНІ, ТЕХНОЛОГІЧНІ ТА ПРАВОВІ АСПЕКТИ**

Поляков Євген Валентинович

*доцент кафедри оперативно-розшукової діяльності
навчально-наукового інституту підготовки фахівців
для підрозділів кримінальної поліції Національної поліції
України, кандидат юридичних наук, доцент*

Мотрюк Максим Дмитрович

*курсант 201 взводу навчально-наукового інституту
підготовки фахівців для підрозділів кримінальної поліції
Національної поліції України*

Оперативно-розшукова діяльність (ОРД) є важливим інструментом протидії злочинності, що забезпечує отримання, обробку та аналіз інформації з метою попередження, виявлення та розкриття кримінальних правопорушень. У сучасних умовах цифровізації суспільства та зростання кіберзагроз особливого значення набуває використання технологій штучного інтелекту (ШІ) у