

DOI: 10.33663/1563-3349-2023-34-684-693

УДК 004.896

О. М. СТОЙКО,
доктор політичних наук*
ORCID: 0000-0002-1021-5270

ПОЛІТИКО-ПРАВОВІ НАСЛІДКИ ВИКОРИСТАННЯ ШТУЧНОГО ІНТЕЛЕКТУ

У статті розглянуто переваги та недоліки, пов'язані зі зростанням застосування штучного інтелекту, зокрема його вплив на політичну сферу. Обґрунтовано необхідність правового регулювання його використання та проаналізовано досвід інших держав у цій сфері, зокрема Канади та США. Особлива увага приділена законопроекту «Про штучний інтелект», що перебуває на розгляді Європарламенту, з огляду на євроінтеграційний курс України.

Ключові слова: штучний інтелект, алгоритми, дискримінація, дезінформація, демократія.

Stoiko Olena. Political and Legal Implications of the Use of Artificial Intelligence

The advantages and disadvantages associated with the growing use of artificial intelligence, in particular its impact on the political sphere, are considered. The author substantiates the need for legal regulation of its use and analyzes the experience of other countries in this area, in particular Canada and the United States. Particular attention is paid to the Artificial Intelligence Act, which is currently under consideration by the European Parliament in view of Ukraine's European integration course.

Key words: artificial intelligence, algorithms, discrimination, disinformation, democracy.

Вступ. З 2017 р. публікується Індекс готовності до штучного інтелекту (ШІ), згідно з яким у 2022 р. найбільш готовими до запровадження ШІ є США, Сінгапур, Велика Британія, Фінляндія та Канада¹. З-поміж колишніх країн соцтабору найкращий показник – у Естонії (19), Чехії (30), Польщі (34), Литві (38). Україна ввійшла до цього рейтингу лише у 2019 р. на 63-ю позицію, а зараз перебуває на 60-й, що свідчить про дедалі більшу актуальність правового регулювання ШІ. Автори Індексу розробили 39 індикаторів, що визначають готовність у трьох основних сферах: уряд, технологічний сектор та дані й інфраструктура. Так, уряд повинен мати стратегічне бачення того, як він розвиває та управляє ШІ, підкріплене відповідним правовим регулюванням і увагою до етичних проблем (управління та етика). Крім того, він

* **Stoiko Olena**, Doctor of Political Sciences

повинен мати потужний внутрішній цифровий потенціал, включно з навичками і практиками, які підтримують його адаптивність до нових технологій.

Огляд літератури. Проблематика впливу технологій на суспільно-політичні процеси має відносно тривалу історію, однак лише після 2010 р. почали з'являтися дослідження наслідків використання машинного навчання та штучного інтелекту в етичному, політичному та правовому аспектах. Такі питання розглядаються у працях Н. Бострома, Дж. Брайсон, Л. Девіллерс, У. Еггерса, М. Камінські, Т. Кемпбелла, М. Маро, Л. Шепарда та ін. Вітчизняні дослідники тему впливу ІІІ на гуманітарну сферу практично не досліджували.

Постановка проблеми дослідження. Дедалі частіше застосунки з елементами ІІІ використовуються не лише у технічній сфері, для підвищення ефективності послуг, що надаються приватним і публічним сектором, а й для прийняття рішень, що безпосередньо впливають на життя громадян: визначення кредитоспроможності, найму на роботу, доступу до соціальних благ (медицини, освіти тощо). Однак як будь-яке технологічне рішення ІІІ має як позитивні, так і негативні результати, які лише починають осмислюватися суспільствознавцями.

Метою статті є виявлення політико-правових наслідків застосування ІІІ та аналіз правових механізмів гарантування його безпечного використання на основі досвіду зарубіжних держав.

Виклад основного матеріалу. Автоматизовані системи, що послуговуються ІІІ, допомагають спростити та удосконалити цілу низку процесів: від виробництва продуктів, надання послуг, передбачення метеорологічних явищ до пошуку оптимальних методів лікування хвороб. У звіті, підготовленому в Стенфордському університеті «Штучний інтелект і життя у 2030 році»², зазначено, що у державному секторі ІІІ та прогностична аналітика сприяли установам економніше та ефективніше використовувати ресурси. Ці технології мають значний потенціал для покращення охорони правопорядку та посилення безпеки громадян шляхом адміністрування величезних обсягів даних (зокрема завдяки технології розпізнавання облич), скорочення часу реагування та виявлення зловмисної активності, яка може мати спершу спорадичний характер. ІІІ також може бути корисним у виявленні нерегулярної фінансової діяльності, такої, як відмивання грошей, фінансування тероризму та шахрайства. ІІІ може допомогти контролювати дотримання податкового законодавства та звітності, класифікувати рахунки для прогнозування пробного балансу для податкових коригувань, підтримувати чат-боти для автоматизації відповідей платникам податків³.

У сфері правосуддя ІІІ може використовуватися не лише задля удосконалення адміністративних процедур, спрощуючи заповнення документів та удосконалюючи судові процедури, а й, наприклад, для визначення терміну покарання. При винесенні вироку правопорушнику судді можуть керуватися оцінкою ризику від даної особи, підготовленої ІІІ, з урахуванням його історії

правопорушень, умов життя, та визначатися, на який термін ця особа має бути ізольована від суспільства. Також існують програми на основі ШІ, які передбачають результати розгляду справ у Європейському суді з прав людини.

Використання ШІ актуалізує питання захисту прав інтелектуальної власності. У листопаді 2022 р. у США уже було подано позов до суду на технологічні компанії GitHub, Microsoft і OpenAI, оскільки створений ними інструмент GitHub Copilot, що автоматично генерує комп'ютерний код, по суті, є плагіатом роботи людей-розробників програмного забезпечення, що порушує їхні ліцензії. Постраждалими сторонами у цій справі, на думку позивача, є розробники, які працювали над проєктами з відкритим кодом, не даючи явного дозволу на використання результатів їх праці для навчання ШІ.

Не менш значущий вплив ШІ і на політичну сферу, особливо на сферу комунікації між громадянами та політиками і представниками органів державної влади, оскільки може застосовуватися як для генерування, так і виявлення сфальшованої аудіовізуальної інформації. Системи ШІ використовуються для поширення дезінформації в інтернеті, що перетворює їх на потенційну загрозу для демократії: від сфальшованих відео (deep fake) до онлайн-ботів, які маніпулюють публічним дискурсом, імітуючи консенсус і поширюючи сфабриковані новини⁴, існує небезпека, що системи ШІ підривають соціальну довіру. Технологія може використовуватися злочинцями, ідеологічними екстремістами або просто групами з особливими інтересами для маніпуляції людьми заради економічної чи політичної вигоди. Дезінформація є серйозною загрозою для суспільства і демократії загалом, оскільки вона ефективно змінює і маніпулює фактами, щоб створити соціальний зворотний зв'язок, який підриває будь-яке відчуття об'єктивної істини. Дебати про те, що є реальним, швидко переростають у дебати про те, хто має право вирішувати, що є реальним, що призводить до перегляду владних структур, які часто слугують укоріненним інтересам. Водночас різноманітні процедури фактчекінгу здійснюються із застосуванням ШІ, що зумовлює необхідність його введення у правові рамки.

ШІ також може справляти безпосередній вплив на один з ключових елементів демократії – вибори. З його допомогою можна максимально деталізувати профіль виборця, виокремити групу людей з нестійкими політичними уподобаннями і за допомогою таргетованої рекламної кампанії чи кастомізованого контенту в мережі запропонувати саме ту виборчу програму кандидата чи політичної партії, що найкраще відповідає їх потребам. В умовах демократії поразка на виборах є результатом неправильної оцінки настроїв виборців, а використання ШІ дасть змогу мінімізувати цю помилку. До того ж агрегація потреб виборців у реальному часі з широким використанням інструментів електронної демократії перетворить політичні партії на рудинти політичної системи.

Оскільки системи штучного інтелекту виявляються кориснішими в реальному світі, вони розширюють сферу свого застосування, що призводить до

зростання ризиків зловживань. Зі збільшенням можливостей систем ШІ та їх глибшою інтеграцією в суспільство наслідки втрати ефективного контролю над ними викликають дедалі більшу стурбованість⁵. Однією з найбільш нагальних небезпек ШІ є техносолюцизм – погляд, згідно з яким ШІ можна розглядати як панацею, хоча насправді він є лише інструментом⁶. Чим більшого успіху досягає ШІ, тим більшою стає спокуса його застосування для розв’язання усіх суспільних проблем у суспільстві. Але технологія часто створює більші проблеми в процесі вирішення менших. Наприклад, системи, які впорядковують і автоматизують надання соціальних послуг, можуть швидко стати негнучкими і відмовити в доступі до них окремим групам осіб (наприклад, іммігрантам). Коли постає вибір між алгоритмами та людьми, доволі поширеною є думка, що алгоритми пропонують менш упереджені рішення. Однак на практиці автоматизоване прийняття рішень часто може слугувати для відтворення, поглиблення і навіть посилення тих самих упереджень, і навіть посилювати ті упередження, які б хотілося виправити. Технології не є панацеєю від усіх бід і можуть створювати петлі зворотного зв’язку, які посилюють дискримінацію. Алгоритми рекомендацій, наприклад списки відтворення музики чи відео, навчені визначати пріоритетність найбільш «релевантних» елементів на основі аналізу вподобань інших користувачів. Однак метод відбору цих користувачів може сприяти збору упередженої інформації, яка стає ще більш суб’єктивною після обробки алгоритмами і робить користувачів ще більш упередженими, надаючи їм відфільтровану інформацію⁷.

Автоматизоване прийняття рішень може призвести до спотворення результатів, які повторюють і посилюють наявні упередження⁸. Потенційна небезпека полягає в тому, що громадськість сприймає висновки, отримані ШІ, як вірогідні. Такий детерміністський підхід до прийняття рішень машиною може мати трагічні наслідки як у кримінальній, так і в медичній сферах. Коли набори даних непропорційно представляють менш впливових членів суспільства, ймовірним результатом є кричуща дискримінація.

Оскільки ШІ стає здатним аналізувати дедалі більше чинників, які можуть корелювати з ризиком, який становить підсудний для оточення, суди й суспільство загалом можуть прийняти алгоритмічну ймовірність за факт. Тобто алгоритмічна оцінка ризику особи для суспільства може бути інтерпретована іншими як майже вірогідна – результат, що вводить в оману, про що застерігали розробники інструменту. Попри те, що можна побудувати статистично керовану систему ШІ, яка б повідомляла про ступінь вірогідності разом з кожним прогнозом, немає жодної гарантії, що люди, які користуються цими прогнозами, будуть використовувати їх розумно. Прийняття ймовірності за вірогідність означає, що минуле завжди диктуватиме майбутнє. Існує ореол нейтральності та неупередженості, пов’язаний з прийняттям рішень ШІ, у результаті чого ці системи приймаються як об’єктивні, навіть якщо вони можуть бути результатом упереджених історичних рішень або

навіть відвертої дискримінації. Без прозорості щодо даних або алгоритмів ШІ, які їх інтерпретують, громадськість може залишитися в невіданні щодо того, як приймаються рішення, що мають суттєвий вплив на їхнє життя. Не маючи достатньої інформації для подання судового позову, люди можуть втратити доступ як до належної правової процедури, так і до відшкодування шкоди, якщо вони вважають, що були помилково оцінені системами ШІ.

Усвідомлення небезпеки безконтрольного використання ШІ змусило низку країн вдатися до пошуку правових інструментів мінімізації негативних наслідків його використання. Канада прагне позиціонувати себе як лідера в галузі штучного інтелекту, зокрема, завдяки Пан-канадській стратегії розвитку штучного інтелекту^{9*}. Наразі на розгляді парламенту перебуває законопроект С-27 під назвою «Закон про імплементацію Цифрової хартії»¹⁰, покликаний посилити законодавство Канади про захист приватного життя в приватному секторі, створити нові правила для відповідальної розробки та впровадження штучного інтелекту (ШІ), а також сприяти подальшому просуванню імплементації Цифрової хартії Канади. Пропонований Закон про штучний інтелект та дані запровадить нові правила для зміцнення довіри канадців до розробки та розгортання систем штучного інтелекту, зокрема захист канадців шляхом забезпечення того, щоб високоефективні системи ШІ розроблялися паралельно з виявленням, оцінкою та заходами з мінімізації ризиків заподіяння шкоди та упередженості. Законопроект також містить чіткі визначення заборон та покарань щодо використання даних, отриманих незаконним шляхом, для розробки ШІ, а також у випадках, коли бездумне впровадження ШІ завдає серйозної шкоди та за наявності умислу завдати значних економічних збитків шляхом його впровадження.

У США у жовтні 2022 р. оприлюднено проект Білля про права в галузі штучного інтелекту¹¹, покликаний регулювати розробку та використання автоматизованих систем, своєрідне керівництво для суспільства, яке захищає всіх людей від цих загроз та використовує технології для зміцнення демократії. У ньому визначено основні вимоги до ШІ:

1. Безпечні та ефективні системи. Керівництво організацій має брати на себе відповідальність і забезпечувати підзвітність за використання ШІ.

2. Захист від алгоритмічної дискримінації. Використання ШІ не повинно сприяти різному ставленню до осіб на основі їхньої раси, кольору шкіри, етнічної належності, статі, релігії, віку, національного походження, інвалідності, ветеранського статусу, генетичної інформації або будь-якої іншої класифікації, захищеної законодавством.

3. Конфіденційність даних. Особисті дані громадян повинні бути захищені від зловживань, а громадянам має бути надано більше прав, щоб впливати на те, як використовуються дані про них. На першому місці мають бути поси-

* Pan-Canadian Artificial Intelligence Strategy (CIFAR, 2017). URL: <https://www.cifar.ca/ai/pan-canadian-artificial-intelligence-strategy>.

лення захисту і обмеження для даних і висновків, пов'язаних із чутливими сферами, включаючи охорону здоров'я, роботу, освіту, кримінальне правосуддя і фінанси, а також для даних, що стосуються молоді.

У чутливих сферах особисті дані повинні використовуватися лише для виконання необхідних функцій, супроводжуватися етичною перевіркою. Громадяни повинні бути вільні від безконтрольного спостереження; технології спостереження повинні підлягати посиленому нагляду, що включає принаймні попередню оцінку їхньої потенційної шкоди та обмеження сфери застосування для захисту приватності та громадянських свобод. Безперервне спостереження та моніторинг не повинні використовуватися у сфері освіти, роботи, житла або в інших контекстах, де застосування таких технологій спостереження може обмежити права, можливості або доступ.

4. Повідомлення та пояснення. Громадяни мають право отримати пояснення того, як використовується будь-яка автоматизована система ШІ, і розуміти, як і чому вона отримує конкретний результат. Автоматизовані системи повинні надавати пояснення, які є технічно обґрунтованими, значущими та корисними для користувачів, а також для операторів.

5. Наявність альтернативного способу комунікації з людиною, перегляд рішень та запасні варіанти. Громадяни повинні мати можливість відмовитися від участі в системі й доступ до людини, яка може швидко розглянути та розв'язати їх проблеми. Громадянам слід гарантувати доступ до своєчасного розгляду та виправлення ситуації людиною, якщо автоматизована система виходить з ладу, видає помилку або якщо особа має намір оскаржити або опротестувати її вплив. Можливість комунікації з людиною та запасні варіанти повинні бути доступними, справедливими, ефективними, підтримуватися, супроводжуватися відповідним навчанням операторів і не повинні бути надто обтяжливими для громадян.

Найближче до запровадження базових стандартів регулювання ШІ наблизився Євросоюз. Його законодавці виходили з усвідомлення, що розвиток штучного інтелекту не лише обіцяє безліч переваг, а й створює нові ризики для користувачів. Відповідно до широкого застосування продуктів з ШІ слід вжити обов'язкових запобіжних заходів. Перший варіант закону про штучний інтелект¹² був опублікований Єврокомісією у квітні 2021 р. Запропонований закон об'єднує чинні правила та норми ЄС щодо ШІ в один документ та дає визначення системи ШІ як «програмне забезпечення, розроблене з використанням одного або декількох методів і підходів, перелічених у Додатку I (машинне навчання, логіка, підходи, засновані на знаннях, або статистичні підходи), і здатне для певного набору цілей, визначених людиною, генерувати такі результати, як контент, прогнози, рекомендації або рішення, що впливають на середовище, з яким вони взаємодіють».

Ризики використання ШІ виділено у чотири категорії: неприйнятні, висококоризиковані, з обмеженням і мінімальним ризиком.

I – неприйнятні ризики, такі як використання ІІІ в соціальному оцінюванні урядами, як це робиться в Китаї, або іграшки з голосовою підтримкою, що заохочують небезпечну поведінку.

II – високий ризик: технології ІІІ, що використовуються в одному з секторів, перелічених у Додатку III, зокрема: на критично важливих об'єктах інфраструктури (наприклад, транспорті); у сферах, які можуть загрожувати життю і здоров'ю громадян, їх освітній або професійній підготовці (доступ до освіти та кар'єрне зростання особи, наприклад, підрахунок балів на іспитах); як компоненти безпеки продуктів (наприклад, у роботизованій хірургії); у ході працевлаштування, управління працівниками та наданні доступу до самозайнятості (програмне забезпечення для сортування резюме для процедур найму на роботу); при наданні основних приватних і державних послуг (наприклад, кредитний скоринг, який позбавляє громадян можливості отримати кредит); у роботі правоохоронних органів, які можуть втручатися у фундаментальні права людей (наприклад, оцінка достовірності доказів); при управлінні міграцією, наданням притулку та прикордонним контролем (перевірка автентичності проїзних документів); при здійсненні правосуддя та демократичних процедур (наприклад, застосування закону до конкретного набору фактів);

До систем ІІІ з високим ступенем ризику будуть застосовуватися суворі зобов'язання, перш ніж вони зможуть бути випущені на ринок: адекватні системи оцінки та пом'якшення ризиків; висока якість наборів даних, що надходять у систему, для мінімізації ризиків і дискримінаційних результатів; реєстрація діяльності для забезпечення відстежуваності результатів; детальна документація, що надає всю необхідну інформацію про систему та її призначення органам влади, щоб вони могли оцінити її відповідність; чітке та адекватне інформування користувачів; належні заходи людського нагляду для мінімізації ризиків; високий рівень надійності, безпеки та точності.

При цьому усі системи віддаленої біометричної ідентифікації вважаються високоризикованими й до них висуваються суворі вимоги. Використання віддаленої біометричної ідентифікації в публічно доступних місцях для правоохоронних цілей у принципі заборонено. Винятки суворо визначені й регламентовані, наприклад, коли це необхідно для пошуку зниклої дитини, для запобігання конкретної та неминучої терористичної загрози або для виявлення, визначення місцезнаходження, ідентифікації або переслідування злочинця або підозрюваного в скоєнні серйозного кримінального злочину. Таке використання підлягає санкціонуванню судовим або іншим незалежним органом і має відповідні обмеження за часом, географічним охопленням і базами даних, в яких здійснюється пошук.

Однак категоризація систем, що підпадають під цей перелік випадків використання, не буде автоматичною, оскільки вони повинні «створювати ризик заподіяння шкоди здоров'ю, безпеці або основоположним правам фізичних осіб у спосіб, що спричиняє юридичні наслідки для них або має

еквівалентний значний ефект». Після цього постачальники ШІ можуть звернутися до захищеної системи тестування – так званої «пісочниці», для встановлення, чи потрапляє їхня система в категорію високого ризику. Якщо вони не вважатимуть, що це не так, вони повинні будуть подати обґрунтовану заяву до компетентного національного органу, щоб бути звільненими від відповідних зобов'язань. Якщо система буде використовуватися в більш ніж одній державі – члені ЄС, заявка буде направлена до Ради з питань штучного інтелекту – органу ЄС, створення якого депутати Європарламенту обговорюють з метою впорядкування правозастосування на європейському рівні. Основним завданням Ради з ШІ буде забезпечення послідовного застосування та виконання закону в державах – членах ЄС, налагодження гнучкої та оперативної взаємодії із зацікавленими сторонами.

Законодавці також хочуть, щоб під час тестування систем ШІ з високим ступенем ризику розробники враховували не лише наслідки їх використання, а й обґрунтовано передбачувані зловживання, а також будь-який негативний вплив на вразливі групи населення (насамперед на дітей). Що стосується наборів даних, на основі яких працюють алгоритми, то розробники ШІ нестимуть відповідальність за весь життєвий цикл системи й повинні будуть вживати заходів з управління даними та управління ризиками, пов'язаними з практикою збору даних, включаючи перевірку законності джерела даних. Крім того, розробники ШІ повинні враховувати, чи можуть такі набори даних призвести до упередженості, що негативно вплине на здоров'я, безпеку або фундаментальні права людини (наприклад, призвести до незаконної дискримінації). Також до уваги братиметься контекст і цільове призначення системи.

ШІ – обмежений ризик, використання у додатках з конкретними зобов'язаннями щодо прозорості. При використанні таких систем ШІ, як чат-боти, користувачі повинні знати, що вони взаємодіють з машиною, щоб мати змогу ухвалити обґрунтоване рішення про продовження або припинення взаємодії.

VI – мінімальний ризик, наприклад, відеоігри зі ШІ або спам-фільтри. До цієї групи належить більшість систем ШІ, що наразі застосовуються в Євросоюзі.

Парламент має проголосувати за проєкт закону про штучний інтелект до кінця березня 2023 р. Очікується, що після цього голосування у квітні розпочнеться обговорення між державами-членами, Парламентом та Комісією (так званий трилог). Якщо цей графік буде дотримано, остаточний варіант закону про ШІ має бути ухвалений до кінця 2023 р.

Загалом саме Європейський Союз, можливо, зробив найсміливіший крок, запропонувавши закон про штучний інтелект, який не тільки обіцяє стати першою в історії правовою базою для управління ШІ, а й потенційно може бути прийнятий як майбутній глобальний стандарт так само, як це зробив Загальний регламент про захист даних.

Незалежно від остаточної версії закону про ШІ його дотримання буде обов'язковим для будь-якої компанії, що надає послуги з використанням

штучного інтелекту жителям ЄС. Деякі платформи з міжнародною клієнтською базою можуть прийняти стандарти ШІ для всіх користувачів, тоді як платформи з більш ізольованими і локалізованими алгоритмами можуть відмовитися від такого кроку.

Після ухвалення Закон про ШІ стане першим горизонтальним законодавчим актом в ЄС, який регулюватиме системи ШІ, запроваджуючи правила безпечного та надійного розміщення на ринку ЄС продуктів із компонентом ШІ. Очікується, що екстериторіальна сфера застосування закону (тобто його дія поширюватиметься на постачальників і користувачів за межами ЄС, якщо продукція, вироблена системою, використовується в ЄС) та надзвичайно високі штрафи – до 30 млн євро або до 6% від загального світового річного обороту компанії за попередній фінансовий рік – визначатимуть регуляторні вимоги за межами ЄС, як це було у випадку з Європейським загальним регламентом про захист даних.

Висновки. Україна дещо відстала від розвинених держав у правовому регулюванні застосування ШІ. Так, Концепція розвитку ШІ в Україні¹³ була схвалена лише у 2020 р. і визначила ШІ як «організовану сукупність інформаційних технологій, із застосуванням якої можливо виконувати складні комплексні завдання шляхом використання системи наукових методів досліджень і алгоритмів обробки інформації, отриманої або самостійно створеної під час роботи, а також створювати та використовувати власні бази знань, моделі прийняття рішень, алгоритми роботи з інформацією та визначати способи досягнення поставлених завдань». У Концепції виділено 9 галузей застосування ШІ та визначено зміст заходів, зокрема, у публічному управлінні слід сформувати перелік адміністративних послуг, рішення за якими приймаються автоматично, а у сфері правового регулювання потрібно «привести принципи використання ШІ в українському законодавстві до європейських норм». Врахування європейського досвіду та української специфіки у національному законодавстві про застосування цифрових технологій має полегшити як адаптацію до входження в європейський правовий простір, так і сприяти розвитку технологічного сектора в країні.

1. Oxford Insights Government AI Readiness Index 2022. URL: <https://www.oxfordinsights.com/government-ai-readiness-index-2022> 2. Artificial Intelligence and Life in 2030. URL: https://ai100.stanford.edu/sites/g/files/sbiybj18871/files/media/file/ai100report10032016fnl_singles.pdf 3. How Tax is leveraging AI – Including machine learning – In 2019. URL: <https://www.pwc.com/gx/en/tax/publications/assets/how-tax-leveraging-ai-machine-learning-2019.pdf> 4. Buchanan B., Lohn A., Musser M., Sedova K. Truth, Lies, and Automation: How Language Models Could Change Disinformation. URL: <https://cset.georgetown.edu/publication/truth-lies-and-automation/> 5. Christian B. The Alignment Problem: Machine Learning and Human Values. N.Y.: W. W. Norton & Company, 2020. 496 p. 6. Philanthropy's Techno-Solutionism Problem6 Democracy and Civic Life: What Is the Long Game for Philanthropy? URL: <https://knightfoundation.org/philanthropys-techno-solutionism-problem/> 7. Noble S.U. Algorithms of Oppression: How

Search Engines Reinforce Racism. N.Y.: NYU Press, 2018. 248 p. **8.** Gathering Strength, Gathering Storms: The One Hundred Year Study on Artificial Intelligence (AI100) 2021 Study Panel Report. URL: https://ai100.stanford.edu/sites/g/files/sbiybj18871/files/media/file/AI100Report_MT_10.pdf **9.** Pan-Canadian Artificial Intelligence Strategy (CIFAR, 2017). URL: <https://www.cifar.ca/ai/pan-canadian-artificial-intelligence-strategy> **10.** Digital Charter Implementation Act, 2022. URL: <https://www.parl.ca/legisinfo/en/bill/44-1/c-27> **11.** Blueprint for an AI Bill of Rights: Making automated systems work for the American people. URL: <https://www.whitehouse.gov/ostp/ai-bill-of-rights/> **12.** Proposal for Artificial Intelligence Act and amending certain legislative acts. COM/2021/206 final. URL: <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX:52021PC0206> **13.** Про схвалення Концепції розвитку штучного інтелекту в Україні: Розпорядження Кабінету Міністрів України від 2 грудня 2020 р. № 1556-п. / Pro skhvalennia Kontseptsii rozvytku shtuchnoho intelektu v Ukraini; Rozporiadzhennia Kabinetu Ministriv Ukrainy vid 2 hrudnia 2020 r. URL: <https://zakon.rada.gov.ua/laws/show/1556-2020-%D1%80#Text>

Stoiko Olena. Political and Legal Implications of the Use of Artificial Intelligence

Increasingly, applications with AI elements are being used not only in the technical field, to improve the efficiency of services provided by the private and public sectors, but also to make decisions that directly affect the lives of citizens. However, like any technological solution, AI has both positive and negative results, which are only beginning to be understood by social scientists.

The purpose of the article is to identify the political and legal consequences of AI application and to analyze the legal mechanisms for ensuring its safe use based on the experience of foreign countries.

As AI systems prove to be increasingly useful in the real world, they expand their scope of application, which leads to an increase in the risks of abuse. The consequences of losing effective control over them are of growing concern. Automated decision-making can lead to distorted results that repeat and reinforce existing biases. There is an aura of neutrality and impartiality associated with AI decision-making, resulting in these systems being accepted as objective, even though they may be the result of biased historical decisions or even outright discrimination. Without transparency about the data or the AI algorithms that interpret it, the public may be left in the dark about how decisions that have a significant impact on their lives are made.

Awareness of the dangers of uncontrolled AI use has led a number of countries to seek legal instruments to minimize the negative consequences of its use. The European Union is the closest to introducing basic standards for AI regulation. A draft of Artificial Intelligence Act was published in 2021 and classifies the risks of using AI into four categories: unacceptable, high-risk, limited, and minimal. Once adopted, the AI Act will be the first horizontal legislative act in the EU to regulate AI systems, introducing rules for the safe and secure placement of AI-enabled products on the EU market. Taking into account the European experience and Ukrainian specifics in domestic legislation on the use of digital technologies should facilitate both adaptation to the European legal space and promote the development of the technology sector in the country.

Key words: artificial intelligence, algorithms, discrimination, disinformation, democracy.