



МЕНЕДЖМЕНТ

УДК 65.011:004.8:340

DOI <https://doi.org/10.5281/zenodo.20713889>

ПРАВОВЕ РЕГУЛЮВАННЯ ВИКОРИСТАННЯ ШТУЧНОГО ІНТЕЛЕКТУ В СИСТЕМІ МЕНЕДЖМЕНТУ ОРГАНІЗАЦІЙ: УПРАВЛІНСЬКІ РИЗИКИ ТА КОМПЛАЄНС

Станіслав Салоїд,

к.е.н., доцент, доцент кафедри менеджменту підприємств,

Національний технічний університет України

«Київський політехнічний інститут імені Ігоря Сікорського»

e-mail: s_saloid@ukr.net

ORCID ID: <https://orcid.org/0000-0002-3294-2671>

Олександра Хлебінська

д.філ.з менеджменту, старший викладач кафедри менеджменту

підприємств,

Національний технічний університет України

"Київський політехнічний інститут імені Ігоря Сікорського"

e-mail: a.khlebynska@gmail.com

ORCID: <https://orcid.org/0000-0002-7977-0483>

Прийнято: 15.05.2026 | Опубліковано: 30.05.2026

Анотація: Стаття присвячена комплексному дослідженню правового регулювання штучного інтелекту (ШІ) у системі корпоративного менеджменту в умовах поглиблення євроінтеграційних процесів. Актуальність теми зумовлена стрімким поширенням ШІ-технологій у бізнес-практиці та суттєвим відставанням нормативно-правової бази від темпів цифрової трансформації підприємств. В основі



дослідження — системний аналіз Регламенту (ЄС) 2024/1689 (EU AI Act) як основоположного акта міжнародного регулювання у цій сфері, а також порівняльно-правовий аналіз його норм у зіставленні з чинним законодавством України та стратегічними ініціативами Міністерства цифрової трансформації. Детально розкрито чотирирівневу класифікацію ШІ-систем за ступенем ризику та проаналізовано її практичне значення для різних функціональних підрозділів організації — від управління персоналом і маркетингу до фінансового контролю і системи безпеки. У статті систематизовано три ключові групи управлінських ризиків, що виникають у процесі неконтрольованого впровадження ШІ: правові (порушення авторських прав та витік персональних даних), операційні та стратегічні (ефект «галюцинацій», алгоритмічна упередженість), а також етичні та репутаційні (ефект «чорної скриньки», зниження довіри персоналу). Розроблено та теоретично обґрунтовано покроковий інструментарій побудови системи корпоративного AI-комплаєнсу, що охоплює аудит ШІ-інструментів, оцінку ризиків, регламентування, технічний контроль і навчання персоналу. Запропоновано організаційну модель розподілу відповідальності, що передбачає введення посади AI Compliance Officer та Комітету з етики ШІ. Як практичний інструмент управління ризиками представлено матрицю комплаєнс-заходів для ключових бізнес-процесів. Окремо проаналізовано специфіку адаптації українського законодавства до вимог EU AI Act та основні бар'єри впровадження стандартів AI-комплаєнсу у вітчизняному бізнесі. Наукова новизна дослідження полягає у формуванні прикладної моделі корпоративного AI-комплаєнсу, адаптованої до умов євроінтеграції та специфіки діяльності українських підприємств, а також у теоретичному осмисленні «парадоксу людського контролю» у контексті автоматизаційного упередження менеджерів.



Ключові слова: штучний інтелект, управлінські ризики, корпоративний комплаєнс, цифровий менеджмент, цифрова трансформація.

**LEGAL REGULATION OF THE USE OF ARTIFICIAL
INTELLIGENCE IN ORGANIZATIONAL MANAGEMENT SYSTEMS:
MANAGERIAL RISKS AND COMPLIANCE**

Stanislav Saloid,

PhD in Economics, Associate Professor, Associate Professor of the Department
of Enterprise Management, National Technical University of Ukraine

"Ihor Sikorskyi Kyiv Polytechnic Institute"

e-mail: s_saloid@ukr.net

ORCID ID: <https://orcid.org/0000-0002-3294-2671>

Oleksandra Khlebynska,

PhD in Management, Senior Lecturer of the Department of Enterprise
Management, National Technical University of Ukraine

"Ihor Sikorskyi Kyiv Polytechnic Institute"

e-mail: a.khlebynska@gmail.com

ORCID: <https://orcid.org/0000-0002-7977-0483>

Abstract: *The article presents a comprehensive study of the legal regulation of artificial intelligence (AI) within corporate management systems in the context of deepening European integration. The relevance of the topic is driven by the rapid proliferation of AI technologies in business practice and the significant lag of the regulatory framework behind the pace of digital transformation at the enterprise level. The study is grounded in a systemic analysis of Regulation (EU) 2024/1689 (EU AI Act) as the foundational international regulatory instrument in this domain, complemented by a comparative legal analysis of its provisions against current Ukrainian legislation and the strategic initiatives of the Ministry of Digital Transformation of Ukraine.*



The article provides a detailed examination of the four-tier risk-based classification of AI systems established by the EU AI Act and analyses its practical implications for various functional divisions of an organisation — from human resources management and marketing to financial control and security systems. Three key groups of managerial risks arising from uncontrolled AI deployment are systematised: legal risks (intellectual property infringement and personal data leakage); operational and strategic risks (hallucination effects and algorithmic bias); and ethical and reputational risks (the black-box effect and erosion of employee trust).

A step-by-step framework for building a corporate AI compliance system is developed and theoretically substantiated, encompassing AI tool auditing, risk assessment, internal policy regulation, technical data protection controls, and personnel training. An organisational model for distributing AI governance responsibilities is proposed, introducing the roles of AI Compliance Officer and an AI Ethics Committee composed of cross-functional senior management. A compliance risk matrix is presented as a practical decision-support tool for key business processes including HR, marketing, finance, and security.

The article separately examines the specifics of Ukrainian legislative adaptation to EU AI Act requirements and identifies the principal barriers to AI compliance implementation in domestic business — namely the shortage of professionals combining IT and legal expertise, financial constraints under wartime conditions, and cultural resistance manifested as Shadow AI practices. The Estonian regulatory model is considered as a benchmark for state-supported compliance infrastructure.

The scientific novelty of the study lies in the formation of an applied corporate AI compliance model adapted to the conditions of European integration and the operational realities of Ukrainian enterprises, as well as in the theoretical conceptualisation of the "human oversight paradox" in the context of automation bias among managers. The article argues for the adoption of a human



augmentation model, wherein AI serves as an analytical assistant while final decision-making authority and legal accountability remain with the human manager.

Keywords: *artificial intelligence, managerial risks, corporate compliance, digital management, digital transformation.*

Сучасний менеджмент переживає глибоку трансформацію, зумовлену переходом до парадигми цифрового управління, в якій штучний інтелект поступово набуває статусу повноцінного суб'єкта прийняття рішень на основі аналізу масивів великих даних [1]. Ці зміни відбуваються в усіх функціональних сферах підприємства — від стратегічного планування і управління персоналом до операційного менеджменту та клієнтського сервісу. Попри незаперечну економічну привабливість технологій ШІ, темпи їх інтеграції суттєво випереджають еволюцію нормативно-правової бази, що створює системні ризики на рівні організацій [2].

Особливої актуальності ця проблематика набуває для українського бізнесу в умовах євроінтеграції. Вихід на ринки ЄС, залучення іноземних інвестицій та співпраця з європейськими партнерами вимагають відповідності нормам, передусім Регламенту (ЄС) 2024/1689 (EU AI Act), що встановлює уніфіковані правила застосування ШІ на всій території Євросоюзу та поширює свою дію екстериторіально. Відтак формування системи корпоративного AI-комплаєнсу, адаптованої до вимог як українського, так і міжнародного законодавства, перетворюється на критично важливе управлінське та науково-практичне завдання для менеджерів усіх рівнів.

Наведені обставини свідчать про нагальну потребу в розробці прикладних механізмів мінімізації ризиків у сфері використання ШІ, що перебуває на перетині корпоративного управління, права та цифрових технологій.



Аналіз останніх досліджень і публікацій. Питання інтеграції ШІ в корпоративне управління знаходять відображення у зростаючому масиві наукових праць. А. Азензул, Н. Махуат та К. Мохліс [4] дослідили вплив штучного інтелекту на механізми корпоративного управління, зокрема щодо мінімізації ухилення від оподаткування та посилення наглядових функцій. К. С. Кертехян [5] провів бібліометричний аналіз проблематики алгоритмічного управління і встановив його суперечливий вплив на досвід і задоволеність працівників у сучасному середовищі. Правові аспекти застосування ШІ у вітчизняному контексті висвітлює С. Берназюк [6], аналізуючи практику Суду ЄС та її значення для формування національного законодавства.

Проблеми стратегічного управління підприємствами в умовах цифровізації та євроінтеграції досліджуються у працях В. С. Ніценка, Н. А. Герасимчук, Ю. В. Мазура та співавторів [7], де розглядаються трансформаційні процеси у маркетинговій діяльності підприємств. В. Білик, О. Фурсова та В. Кредісов [8] наголошують на вагомості ділової репутації як інструменту інтеграції до Єдиного ринку ЄС, що безпосередньо пов'язано з виконанням вимог AI-комплаєнсу. Водночас прикладні механізми побудови внутрішнього AI-комплаєнсу для українських підприємств залишаються недостатньо дослідженими і потребують системного наукового опрацювання.

Формулювання цілей статті (постановка завдання). Метою статті є розробка науково обґрунтованого практичного інструментарію побудови системи корпоративного AI-комплаєнсу для мінімізації управлінських ризиків, пов'язаних із впровадженням ШІ, та забезпечення відповідності вимогам національного та європейського законодавства.

Для досягнення зазначеної мети поставлено такі завдання: проаналізувати ключові норми EU AI Act та їх значення для управлінської



практики; систематизувати типи ризиків, що виникають у процесі некерованого застосування ШІ на підприємствах; розробити алгоритм впровадження системи корпоративного AI-комплаєнсу та матрицю мінімізації управлінських ризиків; визначити особливості та бар'єри адаптації українського бізнесу до вимог міжнародних стандартів у сфері ШІ.

Виклад основного матеріалу дослідження. Регламент EU AI Act характеризується широкою екстериторіальною сферою дії: його вимоги є обов'язковими для всіх суб'єктів господарювання, чиї продукти або послуги на основі ШІ чинять вплив на осіб, які перебувають на території ЄС [3, 12]. Це означає, що значна частина українських підприємств вже підпадає під дію нового регулювання — насамперед IT-компанії, що обслуговують клієнтів із ЄС, організації, які здійснюють рекрутинг для своїх філій в Євросоюзі, а також підприємства, що здійснюють збут товарів і послуг через алгоритмічні платформи. Нехтування цими вимогами загрожує накладенням штрафних санкцій на рівні до 35 мільйонів євро або 7% від сукупного річного обороту компанії [3].

Регламент передбачає поетапне введення в дію окремих положень. З лютого 2025 року набули чинності норми щодо заборони ШІ-систем «неприйняттого ризику», зокрема систем розпізнавання емоцій у робочому середовищі та соціального скорингу. З серпня 2025 року встановлено зобов'язання щодо прозорості для загального призначення ШІ-моделей (GPAI) — зокрема таких як ChatGPT чи Gemini, — включаючи вимогу інформування користувачів про взаємодію з ШІ [10]. Наступний етап — серпень 2026 року — передбачає введення повного режиму контролю для систем високого ризику у сферах фінансового скорингу та автоматизованого відбору персоналу [3]. Попри можливе перенесення окремих дедлайнів для найскладніших систем на 2027–2028 роки [12], базові вимоги щодо прозорості та захисту персональних даних є чинними вже нині.



Українське законодавство у сфері ШІ перебуває у стані активного формування. Міністерство цифрової трансформації реалізує підхід «знизу вгору» (Bottom-up), що передбачає поступову, а не примусову адаптацію регуляторного середовища з метою уникнення бар'єрів для інновацій [9]. Стратегічний вектор розвитку визначено в Білій книзі з регулювання ШІ, що фіксує курс на гармонізацію з нормами ЄС. Для безпечного тестування розробок без ризику санкцій запроваджено механізм правової «пісочниці» (regulatory sandbox), а для захисту прав людини інтегровано методологію HUDERIA. Питання військових і оборонних застосувань ШІ виведено за межі загального цивільного регулювання [13].

Практичне значення для управлінців має чотирирівнева класифікація ШІ-систем, що є ключовою нормативною категорією EU AI Act і визначає обсяг вимог до кожного типу систем [3]. Деталізовану характеристику рівнів ризику та їх зв'язок із конкретними управлінськими завданнями представлено у Таблиці 1.

Таблиця 1

Класифікація ШІ-систем за рівнями ризику в менеджменті

Категорія ризику	Сфера в менеджменті	Приклади інструментів	Ключові вимоги
Неприйнятний	Нагляд контроль та	Розпізнавання емоцій, соціальний скоринг працівників	Повна заборона з лютого 2025 р. Штрафи до 35 млн євро.
Високий	Рекрутинг, фінансовий оцінювальний скринінг	Алгоритмічний скринінг резюме, ШІ-оцінка КРІ, кредитний скоринг	Оцінка відповідності, обов'язковий людський контроль, логування.
Обмежений	Маркетинг	Чат-боти підтримки клієнтів, генератори текстів	Обов'язкове маркування контенту (інформування користувача).
Мінімальний	Адміністрування	Спам-фільтри, ШІ-оптимізація розкладу	Спеціальних юридичних вимог немає.

Джерело: сформовано авторами на основі [3, 12]

Нераціональне та безсистемне використання ШІ в управлінській практиці генерує три взаємопов'язані групи ризиків [3, 14, 15]. Правові ризики включають: мимовільне порушення авторських прав при генерації



комерційного контенту або програмного коду [10]; витік конфіденційної інформації та персональних даних у разі використання персоналом публічних ШІ-сервісів, що прямо порушує вимоги вітчизняного законодавства та Регламенту ЄС 2016/679 (GDPR) [3]. Операційні та стратегічні ризики пов'язані з ефектом «галюцинацій» — генерацією ШІ правдоподібних, проте недостовірних фактів або некоректних кількісних показників [14] — а також з явищем алгоритмічної упередженості (AI bias): якщо систему навчено на статистично нерепрезентативних або застарілих вибірках, вона відтворює і навіть підсилює дискримінаційні патерни у процесах рекрутингу чи оцінювання [15]. Нарешті, етичні та репутаційні ризики зумовлені ефектом «чорної скриньки» (Black Box): непрозорість алгоритмічних рішень щодо преміювання або звільнення породжує технострес серед персоналу, руйнує організаційну довіру та знижує залученість [16].

Для нейтралізації зазначених ризиків менеджменту необхідно сформувати цілісну систему внутрішнього регулювання — корпоративний AI-комплаєнс, базованим документом якого є Корпоративна політика використання ШІ (Corporate AI Policy) [17]. На підставі аналізу кращих практик запропоновано п'ятикроковий алгоритм її впровадження. Першим кроком є комплексний аудит із метою інвентаризації всіх ШІ-інструментів організації, включно із так званим «тіньовим ШІ» (Shadow AI) — несанкціонованими сервісами, що використовуються персоналом поза корпоративною інфраструктурою [19]. Другий крок — оцінка ризиків кожного виявленого інструменту відповідно до класифікації EU AI Act та стандарту ISO/IEC 42001 [12, 11]. На третьому кроці, регламентуванні, визначаються категорії інформації, заборонені для завантаження у зовнішні ШІ-сервіси, та встановлюються обов'язкові процедури верифікації результатів людиною [18]. Четвертий крок — технічний контроль — передбачає налаштування систем захисту від втрати даних (Data Loss



Prevention, DLP), що унеможливляють несанкціоновану передачу конфіденційних матеріалів на зовнішні сервери [18]. Завершальним, п'ятим кроком є систематичне навчання та інструктаж персоналу щодо чинних корпоративних правил і законодавчих вимог.

Ефективність системи AI-комплаєнсу забезпечується відповідним розподілом організаційної відповідальності [11]. Ключовою у цій структурі є роль AI Compliance Officer — менеджера з моніторингу змін законодавства, ведення реєстру ШІ-інструментів та проведення їх первинної оцінки на відповідність нормативним вимогам [19]. У великих організаціях, що застосовують системи високого ризику, рекомендується утворення міжфункціонального Комітету з етики ШІ у складі представників топменеджменту (технічного, юридичного та HR-директорів), який здійснює стратегічний нагляд і санкціонує впровадження ризикованих алгоритмів [19].

Практичний інструмент для прийняття рішень у сфері AI-комплаєнсу представлено у Таблиці 2 — матриці мінімізації управлінських ризиків, що встановлює взаємозв'язок між бізнес-процесами, ідентифікованими ризиками та відповідними превентивними заходами.

Таблиця 2

Матриця мінімізації управлінських ризиків AI-комплаєнсу

Бізнес-процес	ШІ-система	Ідентифікований ризик	Превентивний захід
HR	Скринінг та ранжування резюме	Алгоритмічна дискримінація, непрозорість відмов	Людський контроль (Human-in-the-loop), аудит вибірок, надання права на апеляцію
Маркетинг	Генерація контенту через LLM	Порушення авторських прав, поширення дезінформації	Використання корпоративних API, маркування контенту, обов'язковий факт-чекінг
Фінанси	Предиктивні моделі скорингу	Фінансові втрати через «галюцинації» моделей	Застосування інструментів Explainable AI (XAI), дублювання класичними методами
Безпека	Системи розпізнавання облич	Незаконний збір біометрії, порушення приватності	Локальне зберігання даних, письмова згода працівників, заборона оцінки емоцій

Джерело: розроблено авторами



Проведене дослідження дає підстави для таких висновків. По-перше, безсистемне впровадження ШІ в організаційну практику генерує комплекс правових, операційних та репутаційних ризиків, які можуть спричинити фінансові збитки, юридичну відповідальність і деструктивні зміни в організаційній культурі. По-друге, екстериторіальна дія EU AI Act у поєднанні з євроінтеграційним курсом України перетворює побудову системи AI-комплаєнсу на стратегічний пріоритет для вітчизняного менеджменту вже у короткостроковій перспективі. По-третє, компанії, що першими впровадять прозорі та верифіковані стандарти використання ШІ відповідно до вимог ISO/IEC 42001, отримають конкурентну перевагу у вигляді підвищеної довіри з боку міжнародних партнерів та інвесторів.

Для практичної імплементації запропонованих підходів рекомендується здійснити такі першочергові кроки: провести повний аудит ШІ-інструментів, що використовуються в організації; розробити та затвердити Корпоративну політику використання ШІ (Corporate AI Policy); ввести посаду або функцію AI Compliance Officer; налаштувати DLP-системи захисту корпоративних даних від несанкціонованого розголошення через зовнішні ШІ-сервіси.

Перспективи подальших наукових розвідок вбачаємо у розробці галузевих специфікацій корпоративного AI-комплаєнсу для окремих секторів економіки, дослідженні впливу стандарту ISO/IEC 42001 на конкурентоспроможність підприємств, а також у кількісному вимірюванні ефекту автоматизаційного упередження серед управлінського персоналу вітчизняних організацій.

Список використаних джерел

1. Kesting, P. (2024). How artificial intelligence will revolutionize management studies: A Savagean perspective. *Scandinavian Journal of*



Management, 40(2), 101330. URL:
<https://doi.org/10.1016/j.scaman.2024.101330> (дата звернення: 15.05.2026).

2. Artificial intelligence in public governance: evaluating opportunities, risks, and policy frameworks. (2024). International Journal of Social Sciences Bulletin, 1(1). URL:
<https://www.pjssrjournal.com/index.php/Journal/article/view/716> (дата звернення: 15.05.2026).

3. European Union. (2024). Regulation (EU) 2024/1689 of the European Parliament and of the Council laying down harmonised rules on artificial intelligence (Artificial Intelligence Act). Official Journal of the European Union, L 1689. URL: <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX:32024R1689> (дата звернення: 15.05.2026).

4. Azenzoul, A., Mahouat, N., & Mokhlis, K. (2026). Artificial intelligence and corporate governance: strengthening oversight mechanisms to mitigate tax avoidance. Corporate Governance: The International Journal of Business in Society, 26(1), 1–27. URL:
<https://www.emerald.com/insight/content/doi/10.1108/CG-11-2025-0931> (дата звернення: 17.05.2026).

5. Kertechian, K. S. (2026). A bibliometric analysis of algorithmic governance and employee experience in the modern workplace. Cogent Business & Management, 13(1), 2655925. URL:
<https://www.researchgate.net/publication/403496936> (дата звернення: 17.05.2026).

6. Берназюк С. (2026). Штучний інтелект у правосудді: від практики Суду ЄС до європейських стандартів [презентація]. Верховний Суд України. URL:
https://court.gov.ua/storage/portal/supreme/prezent2026/177_AI_in_Justice_CJ_EU_Practice_Bernaziuk.pdf (дата звернення: 17.05.2026).



7. Ніценко, В. С., Герасимчук, Н. А., Мазур, Ю. В., Деньгуб, В. В., & Кульганік, О. М. (2026). Цифровізація маркетингової діяльності підприємств в умовах сучасної економіки. *Актуальні проблеми сталого розвитку*, 3(4), 252–261. URL: [https://doi.org/10.60022/3\(4\)-30S](https://doi.org/10.60022/3(4)-30S) (дата звернення: 18.05.2026).

8. Білик, В., Фурсова, О., & Кредісов, В. (2026). Ділова репутація бізнесу як інструмент інтеграції України до Європейського економічного простору. *Development Service Industry Management*, 2, 220–227. URL: [https://doi.org/10.31891/dsim-2026-14\(28\)](https://doi.org/10.31891/dsim-2026-14(28)) (дата звернення: 18.05.2026).

9. Міністерство цифрової трансформації України. (2024). Біла книга з регулювання ІІІ в Україні: бачення Мінцифри. URL: <https://backend.hromada.gov.ua/storage/uploads/files/research/bila-kniga-z-regulyuvannya-si-v-ukrayini-bacennya-mincifri/Регулювання%20ІІІ.pdf> (дата звернення: 18.05.2026).

10. National Institute of Standards and Technology. (2023). Artificial Intelligence Risk Management Framework (AI RMF 1.0). NIST AI 100-1. URL: <https://doi.org/10.6028/NIST.AI.100-1> (дата звернення: 18.05.2026).

11. ISO/IEC 42001:2023. Information technology — Artificial intelligence — Management system. International Organization for Standardization. URL: <https://www.iso.org/standard/81230.html> (дата звернення: 20.05.2026).

12. Crowell & Moring LLP. (2026). Artificial intelligence and human resources in the EU: a 2026 legal overview. URL: <https://www.crowell.com/en/insights/client-alerts/artificial-intelligence-and-human-resources-in-the-eu-a-2026-legal-overview> (дата звернення: 20.05.2026).

13. Міністерство цифрової трансформації України. (2024). Регулювання ІІІ — офіційний ресурс Мінцифри. URL: <https://ai.thedigital.gov.ua/vision> (дата звернення: 20.05.2026).



14. Responsible AI Institute. (2024). Generative AI Best Practices Guide. URL: <https://www.responsible.ai/resources/generative-ai-best-practices> (дата звернення: 20.05.2026).

15. Gibson, R. (2025). Algorithmic bias and generative AI in recruitment frameworks. Proceedings of the IEEE International Conference on Data Mining Workshops (ICDMW), 813200. URL: <https://www.researchgate.net/publication/397621412> (дата звернення: 20.05.2026).

16. Dang, H., & Liu, Y. (2025). The psychological toll of algorithmic management: Cynicism, alienation, and technostress. Frontiers in Psychology, 16, 1745164. URL: <https://doi.org/10.3389/fpsyg.2026.1745164> (дата звернення: 20.05.2026).

17. Intuition Labs. (2026). Developing a corporate AI policy: A comprehensive governance guide. URL: <https://intuitionlabs.ai/articles/developing-corporate-ai-policy-governance-guide> (дата звернення: 20.05.2026).

18. FifthRow. (2026). Turning compliance pressure into opportunity: Navigating the EU AI Act for HR transformation. URL: <https://www.fifthrow.com/blog/turning-compliance-pressure-into-opportunity-navigating-the-eu-ai-act-for-hr-transformation> (дата звернення: 20.05.2026).

19. Allen, R., & Choudhury, P. (2022). Algorithmic authority vs. professional judgment: Theoretical frameworks for human-in-the-loop systems. Strategic Management Journal, 44(4), 412–430. URL: <https://doi.org/10.1002/smj.3440> (дата звернення: 20.05.2026).

References

1. Kesting, P. (2024). How artificial intelligence will revolutionize management studies: A Savagean perspective. Scandinavian Journal of Management, 40(2), 101330. URL: <https://doi.org/10.1016/j.scaman.2024.101330> (accessed: 15.05.2026).



2. Artificial intelligence in public governance: evaluating opportunities, risks, and policy frameworks. (2024). International Journal of Social Sciences Bulletin, 1(1). URL: <https://www.pjssrjournal.com/index.php/Journal/article/view/716> (accessed: 15.05.2026).

3. European Union. (2024). Regulation (EU) 2024/1689 of the European Parliament and of the Council laying down harmonised rules on artificial intelligence (Artificial Intelligence Act). Official Journal of the European Union, L 1689. URL: <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX:32024R1689> (accessed: 15.05.2026).

4. Azenzoul, A., Mahouat, N., & Mokhlis, K. (2026). Artificial intelligence and corporate governance: strengthening oversight mechanisms to mitigate tax avoidance. Corporate Governance: The International Journal of Business in Society, 26(1), 1–27. URL: <https://www.emerald.com/insight/content/doi/10.1108/CG-11-2025-0931> (accessed: 17.05.2026).

5. Kertechian, K. S. (2026). A bibliometric analysis of algorithmic governance and employee experience in the modern workplace. Cogent Business & Management, 13(1), 2655925. URL: <https://www.researchgate.net/publication/403496936> (accessed: 17.05.2026).

6. Bernaziuk, S. (2026). Artificial intelligence in justice: from the practice of the Court of Justice of the EU to European standards [presentation]. Supreme Court of Ukraine. URL: https://court.gov.ua/storage/portal/supreme/prezent2026/177_AI_in_Justice_CJ_EU_Practice_Bernaziuk.pdf (accessed: 17.05.2026).

7. Nitsenko, V. S., Herasymchuk, N. A., Mazur, Yu. V., Denhub, V. V., & Kulhanyk, O. M. (2026). Digitalization of enterprise marketing activities in the conditions of the modern economy. Current Issues of Sustainable Development, 3(4), 252–261. URL: [https://doi.org/10.60022/3\(4\)-30S](https://doi.org/10.60022/3(4)-30S) (accessed: 18.05.2026).



8. Bilyk, V., Fursova, O., & Kredisov, V. (2026). Business reputation as a tool for Ukraine's integration into the European economic space. *Development Service Industry Management*, 2, 220–227. URL: [https://doi.org/10.31891/dsim-2026-14\(28\)](https://doi.org/10.31891/dsim-2026-14(28)) (accessed: 18.05.2026).
9. Ministry of Digital Transformation of Ukraine. (2024). White Paper on AI Regulation in Ukraine: the vision of Diia. URL: <https://backend.hromada.gov.ua/storage/uploads/files/research/bila-kniga-z-regulyuvannya-si-v-ukrayini-bacennya-mincifri/> (accessed: 18.05.2026).
10. National Institute of Standards and Technology. (2023). Artificial Intelligence Risk Management Framework (AI RMF 1.0). NIST AI 100-1. URL: <https://doi.org/10.6028/NIST.AI.100-1> (accessed: 18.05.2026).
11. ISO/IEC 42001:2023. Information technology — Artificial intelligence — Management system. International Organization for Standardization. URL: <https://www.iso.org/standard/81230.html> (accessed: 20.05.2026).
12. Crowell & Moring LLP. (2026). Artificial intelligence and human resources in the EU: a 2026 legal overview. URL: <https://www.crowell.com/en/insights/client-alerts/artificial-intelligence-and-human-resources-in-the-eu-a-2026-legal-overview> (accessed: 20.05.2026).
13. Ministry of Digital Transformation of Ukraine. (2024). AI Regulation — Official Resource. URL: <https://ai.thedigital.gov.ua/vision> (accessed: 20.05.2026).
14. Responsible AI Institute. (2024). Generative AI Best Practices Guide. URL: <https://www.responsible.ai/resources/generative-ai-best-practices> (accessed: 20.05.2026).
15. Gibson, R. (2025). Algorithmic bias and generative AI in recruitment frameworks. *Proceedings of the IEEE International Conference on Data Mining Workshops (ICDMW)*, 813200. URL: <https://www.researchgate.net/publication/397621412> (accessed: 20.05.2026).



16. Dang, H., & Liu, Y. (2025). The psychological toll of algorithmic management: Cynicism, alienation, and technostress. *Frontiers in Psychology*, 16, 1745164. URL: <https://doi.org/10.3389/fpsyg.2026.1745164> (accessed: 20.05.2026).
17. Intuition Labs. (2026). Developing a corporate AI policy: A comprehensive governance guide. URL: <https://intuitionlabs.ai/articles/developing-corporate-ai-policy-governance-guide> (accessed: 20.05.2026).
18. FifthRow. (2026). Turning compliance pressure into opportunity: Navigating the EU AI Act for HR transformation. URL: <https://www.fifthrow.com/blog/turning-compliance-pressure-into-opportunity-navigating-the-eu-ai-act-for-hr-transformation> (accessed: 20.05.2026).
19. Allen, R., & Choudhury, P. (2022). Algorithmic authority vs. professional judgment: Theoretical frameworks for human-in-the-loop systems. *Strategic Management Journal*, 44(4), 412–430. URL: <https://doi.org/10.1002/smj.3440> (accessed: 20.05.2026).