

DOI: <https://doi.org/10.32782/2524-0072/2026-88-5>

УДК: 005.53:004.89

МЕХАНІЗМ СИНЕРГІЇ РІШЕНЬ (DSM): ІНТЕГРАЦІЯ ЛІДЕРСТВА ТА ГЕНЕРАТИВНОГО ШІ ДЛЯ ЗНИЖЕННЯ КОГНІТИВНИХ УПЕРЕДЖЕНЬ

DECISION SYNERGY MECHANISM (DSM): INTEGRATING LEADERSHIP AND GENERATIVE AI TO MITIGATE COGNITIVE BIASES

Бойко Арсеній Олегович

аспірант,

Державний університет інформаційно-комунікаційних технологій

ORCID: <https://orcid.org/0009-0001-0043-7136>**Boyko Arsenii**

State University of Information and Communication Technologies

У турбулентному управлінському середовищі експоненційне зростання даних та поширення генеративного ШІ посилюють когнітивну вразливість лідерів, зокрема упередженість автоматизації та підтвердження у критично важливих рішеннях. Стаття має на меті розробити та концептуально обґрунтувати Механізм синергії рішень (Decision Synergy Mechanism, DSM) як фреймворк взаємодії «людина-алгоритм», що інтегрує стратегічне лідерство, можливості генеративного ШІ та етичний нагляд для пом'якшення таких упереджень. На основі інтегративного огляду літератури та синтезу теорії подвійних процесів, соціокогнітивної моделі довіри та парадигми гібридного інтелекту розроблено трирівневу модель DSM. Вона концептуалізує генеративний ШІ як активного когнітивного опонента, який через AI Reasoning та структуровану інженерію запитів створює цілеспрямоване когнітивне тертя, сприяючи точнішому калібруванню довіри та зниженню когнітивних упереджень.

Ключові слова: Механізм синергії рішень (DSM), генеративний ШІ, стратегічне лідерство, когнітивні упередження, взаємодія людини та алгоритму, калібрування довіри, гібридний інтелект.

In turbulent managerial environments, the exponential growth of data and the rapid spread of generative AI amplify the cognitive vulnerabilities of organizational leaders, particularly automation and confirmation biases in high stakes decisions. These vulnerabilities are further intensified by the lack of structured frameworks that govern human-algorithm interaction in real organizational contexts. This article aims to develop and conceptually justify the Decision Synergy Mechanism (DSM) as a novel human-algorithm framework that integrates strategic leadership, generative AI capabilities, and rigorous ethical oversight to mitigate such biases. The study applies an integrative literature review of peer reviewed sources from 2020–2026 indexed in Scopus and Web of Science, combining multi-disciplinary insights from management, organizational psychology, and information systems. Using concept matrices, it synthesizes Dual Process Theory, the socio cognitive model of trust in automation, and the Hybrid Intelligence paradigm to deductively derive the comprehensive DSM architecture. The results introduce a three level DSM model that explicitly links contextual leadership preconditions, AI enabled synergy processes, and trust calibration outcomes. Furthermore, the framework reconceptualizes generative AI from a passive information "oracle" into an active cognitive opponent that creates deliberate cognitive friction through advanced AI Reasoning, dialectical questioning, and structured prompt engineering. The article formulates three distinct propositions regarding the inverted U shaped relationship between AI literacy and decision weight allocation, the role of contextual AI explanations in calibrating user trust, and the bias reducing effect of AI generated counterfactuals under extreme time pressure and task complexity. The proposed mechanism effectively resolves the "trust paradox" by preventing both blind rejection and blind reliance on AI systems, and can directly inform the design of hybrid decision architectures, leadership development programs, and AI governance policies in ethically sensitive, high complexity operational domains.

Keywords: Decision Synergy Mechanism (DSM), generative AI, strategic leadership, cognitive biases, human-algorithm interaction, trust calibration, hybrid intelligence.



Постановка проблеми. У сучасному високотурбулентному бізнес-середовищі технології генеративного штучного інтелекту стрімко трансформуються з допоміжних інструментів рутинної автоматизації у стратегічні компоненти архітектури підприємства. Вони спроможні не лише оптимізувати щоденні операційні процеси, а й моделювати складні прогностичні сценарії та формувати підґрунтя для ухвалення довгострокових рішень. Проте класична парадигма менеджменту виявляється теоретично та інструментально невідповідною до інтеграції таких інтелектуальних агентів, що зумовлює виникнення глибокого управлінського дисонансу.

Динаміка імплементації ШІ в управлінську практику організацій критично випереджає швидкість рефлексії та розуміння топ-менеджментом внутрішньої логіки, обмежень та специфіки функціонування цих алгоритмів. За відсутності системних знань новітня технологія сприймається управлінським апаратом як непрозорий когнітивний «чорний ящик».

На практиці цей розрив трансформується у критичні помилки під час розподілу ваги рішень між людиною та алгоритмом. Неконтрольоване покладання на виходи ШІ-систем не лише створює тактичні операційні ризики «в моменті», але й здатне системно деформувати стратегічний вектор розвитку підприємства, наражаючи його на довгострокову втрату конкурентоспроможності, фінансові девіації та втрату суб'єктивності лідерського контролю.

Аналіз останніх досліджень і публікацій. Ефективне вирішення складних управлінських проблем сьогодні базується на синергії людського та штучного інтелекту, що формує нову парадигму прийняття рішень «Human+» [16]. Генеративний ШІ стрімко трансформує ці процеси, підвищуючи ефективність обробки даних, допомагаючи управляти часом і знаннями у складних сценаріях [16; 3], а також пом'якшуючи емоційні та когнітивні упередження і забезпечуючи більш об'єктивні рекомендації [15].

Водночас неконтрольоване використання GenAI несе серйозні ризики, пов'язані з когнітивними викривленнями – насамперед надмірною довірою до автоматизованих систем та пошуком інформації, що підтверджує власні погляди. У симульованих середовищах до 78% користувачів приймали рішення на основі результатів ШІ без жодної перевірки, а 71% демонстрували неетичну поведінку через нормалізацію та самозаспокоєння [4].

Емпіричні дані свідчать, що менеджери готові делегувати ШІ близько 30% ваги у прийнятті рішень, однак цей розподіл критично залежить від довіри та сприйняття «свободи волі» алгоритму [1]. З огляду на схильність ШІ до галюцинацій і відтворення успадкованих упереджень [3], просте надання алгоритмам «ваги» без критичного оцінювання є небезпечним.

Для подолання цих пасток необхідний перехід від ШІ як пасивного «оракула» до активного когнітивного опонента. Використання великих мовних моделей у ролі «адвоката диявола» для генерації контраргументів формує більш виважений рівень довіри до технології [14], а ШІ-опоненти, які анонімно представляють альтернативні точки зору, стимулюють критичне мислення та знижують ризик групового мислення [13]. Ефективна реалізація таких підходів спирається на мультиагентні фреймворки та структуровану інженерію запитів [3].

Це вимагає переосмислення ролі лідерства, яке трансформується у гібридний феномен: керівники зміщують фокус із виконання завдань на роль «архітекторів рішень», відповідальних за розподіл завдань між людиною та ШІ, механізми нагляду і вирішення етичних дилем [17]. При цьому лідери потребують не лише технічних пояснень, а й контекстуально насиченого обґрунтування: AI Reasoning генерує адаптовані до контексту наративи, які зберігають людську автономію [2].

Попри визнання потреби в такій співпраці, існуючі інтегративні фреймворки залишаються фрагментованими і не пропонують цілісного механізму управління динамікою GenAI та його впливом на когнітивні упередження, а отже – на менеджмент і підприємство. Відповідаючи на цю прогалину, ця стаття має на меті розробити та концептуально обґрунтувати Механізм Синергії Рішень Людина-Алгоритм (Decision Synergy Mechanism – DSM), що інтегрує стратегічне лідерство, використання GenAI як інструменту когнітивного втручання та механізми етичного нагляду.

Формулювання цілей статті (постановка завдання). Метою статті є концептуальне обґрунтування та розробка Механізму Синергії Рішень Людина-Алгоритм як інструменту мінімізації когнітивних упереджень управлінців та оптимізації процесів прийняття рішень в умовах цифрової трансформації бізнесу.

Для досягнення поставленої мети у дослідженні передбачено вирішення таких завдань: проаналізувати деструктивний

вплив когнітивних упереджень управлінського апарату в умовах неконтрольованої інтеграції генеративного ШІ; обґрунтувати концептуальний перехід від парадигми технічної прозорості алгоритмів до контекстуального обґрунтування та операціоналізації ШІ як активного когнітивного опонента; сформулювати й структурувати трирівневу архітектуру моделі DSM, що інтегрує організаційні передумови лідерства, процеси безпосередньої синергії та контур зворотного зв'язку; дедуктивно вивести систему теоретичних пропозицій, які моделюють характер залежностей між рівнем ШІ-грамотності лідера, штучним когнітивним тертям та точністю калібрування довіри; визначити граничні умови функціонування розробленого механізму в середовищах із високими етичними ризиками й емоційною складністю та окреслити перспективні напрями його подальшої емпіричної валідації.

Виклад основного матеріалу дослідження. Побудова концептуального фреймворку Механізму Синергії Рішень ґрунтується на методології інтегративного огляду літератури [18, 19], який, на відміну від традиційних систематичних оглядів, дозволяє ефективно синтезувати розрізнені теоретичні перспективи з різних дисциплін для створення нових моделей у сферах, що швидко розвиваються. Інформаційну базу дослідження склали рецензовані англомовні статті кuartилів Q1–Q3 із міжнародних наукометричних баз даних Scopus та Web of Science Core Collection за період 2020–2026 років (із фокусом на публікаціях 2024–2026 рр.), відібрані за технологічною, когнітивною та управлінською осями пошукових запитів.

Необхідність такого фреймворку зумовлена тим, що для мінімізації ризиків неконтрольованого застосування генеративного ШІ існує критична потреба структурувати та коригувати процес взаємодії між людиною та алгоритмічними агентами. Існуючі моделі управління штучним інтелектом залишаються переважно техноцентричними, залишаючи поза увагою складну динаміку когнітивного взаємовпливу в системі «людина-алгоритм». За відсутності належної архітектури рішень управлінці схильні впадати у когнітивні крайнощі: від сліпої довіри до системи, відомої як Automation Bias [5], до повного її відторгнення через брак розуміння внутрішньої логіки алгоритму.

Подолання цього розшарування вимагає концептуального зсуву до парадигми Гібридного Інтелекту [6]. Запропонований Механізм

Синергії Рішень переосмислює роль ШІ: з пасивної «чорної скриньки» чи автономного «оракула» він трансформується в активного когнітивного опонента. У межах DSM головна функція алгоритму полягає у створенні цілеспрямованого «когнітивного тертя», що виступає тригером для активації аналітичного мислення лідера – Системи 2 за теорією подвійних процесів [7] – та стимулює глибоке дослідження альтернативних варіантів.

Логіка розкриття цього механізму передбачає послідовний перехід від технологічних засад до системного моделювання та його практичної верифікації. Спочатку досліджено трансформацію ролі ШІ через інструменти контекстуального обґрунтування, що створює підґрунтя для презентації трирівневої архітектури моделі DSM. На основі розробленого фреймворку дедуктивно виведено систему теоретичних пропозицій, а завершує аналіз верифікація граничних умов моделі за допомогою ілюстративних управлінських сценаріїв.

У сучасній практиці управління алгоритмічними системами домінуючим інструментом забезпечення довіри залишається використання Explainable AI (XAI). Проте, як зазначають [8], XAI фокусується виключно на технічній прозорості (наприклад, через алгоритми SHAP або LIME), відповідаючи на питання «як» модель згенерувала результат. Для управлінського середовища цей підхід є фундаментально недостатнім, оскільки технічна валідація не забезпечує розуміння бізнес-контексту, не враховує альтернативних сценаріїв та не сприяє спільній дискусії. Подолання цього обмеження вимагає переходу до інструментарію AI Reasoning, який генерує адаптивні, контекстуально насичені обґрунтування, пояснюючи «чому» конкретне рішення є релевантним у заданих умовах [8].

Однак, ізольоване впровадження AI Reasoning створює новий, ще більш небезпечний когнітивний ризик – феномен «переконливих галюцинацій». Великі мовні моделі (LLM) здатні генерувати абсолютно логічні, контекстуально насичені, але фактологічно хибні думки. Якщо лідер стикається з таким ідеально структурованим «поясненням», його Automation Bias може злетіти до критичного рівня. Для запобігання цьому, в архітектурі DSM технологія AI Reasoning концептуалізується не як фінальна інстанція, а як базис для інституціоналізованої стратегії «Адвоката диявола». Це вимагає від функціональних лідерів застосування цілеспрямованої

інженерії запитів, яка примусово спрямовує генеративний ШІ на формування альтернативних гіпотез та контраргументів до власних висновків. Таке штучне створення «когнітивного тертя» свідомо диверсифікує початкове судження користувача, ефективно руйнуючи упередження підтвердження під час прийняття рішень.

Крім того, для забезпечення валідності згенерованих контраргументів та мінімізації успадкованих алгоритмічних упереджень, механізм когнітивного втручання вимагає технологічного «заземлення». Інтеграція архітектур пошукової генерації дозволяє обмежити простір висновків ШІ виключно куративними, перевіреними джерелами даних [3]. Це перетворює генеративний ШІ з імовірнісного генератора тексту на надійний інструмент доказового управління, що підвищує загальну надійність виходів системи.

На цих засадах запропонований фреймворк DSM інтегрує окреслені теоретичні концепти у єдину систему, що моделює динамічну взаємодію людського та штучного інтелекту. Для забезпечення операціоналізації та цілісності даний механізм концептуалізовано як динамічну циклічну модель, що розгортається на трьох взаємопов'язаних рівнях: контекстуальні передумови, процеси синергії

та результати калібрування довіри. Ключовим елементом системи виступає петля зворотного зв'язку, яка дозволяє візуалізувати, як управлінські передумови впливають на якість взаємодії та як результати рішень трансформують майбутні стратегії лідерства.

Рівень 1: Контекстуальні передумови. Базовим зовнішнім контуром виступає організаційне середовище, де лідерство виконує функцію архітектора рішень. Відповідно до соціокогнітивної моделі довіри в автоматизацію [9], саме на цьому етапі формується початкова готовність до делегування завдань алгоритмам. Ключовим фактором тут є цифрова грамотність лідера, яка, згідно з дослідженнями [10], формує культуру експертного нагляду та визначає базові критерії для застосування генеративного ШІ в адаптивних управлінських сценаріях.

Рівень 2: Процеси DSM. Центральним механізмом моделі є рівень безпосередньої синергії, який функціонує за принципами теорії подвійних процесів [7]. Ця взаємодія інтегрує людський капітал (інтуїцію Системи 1 та етичні принципи) з інструментами алгоритмічного втручання (ШІ як аналітична Система 2). На відміну від традиційних моделей, цей рівень спирається не на технічну прозорість (Explainable AI), а на контекстуальне обґрун-

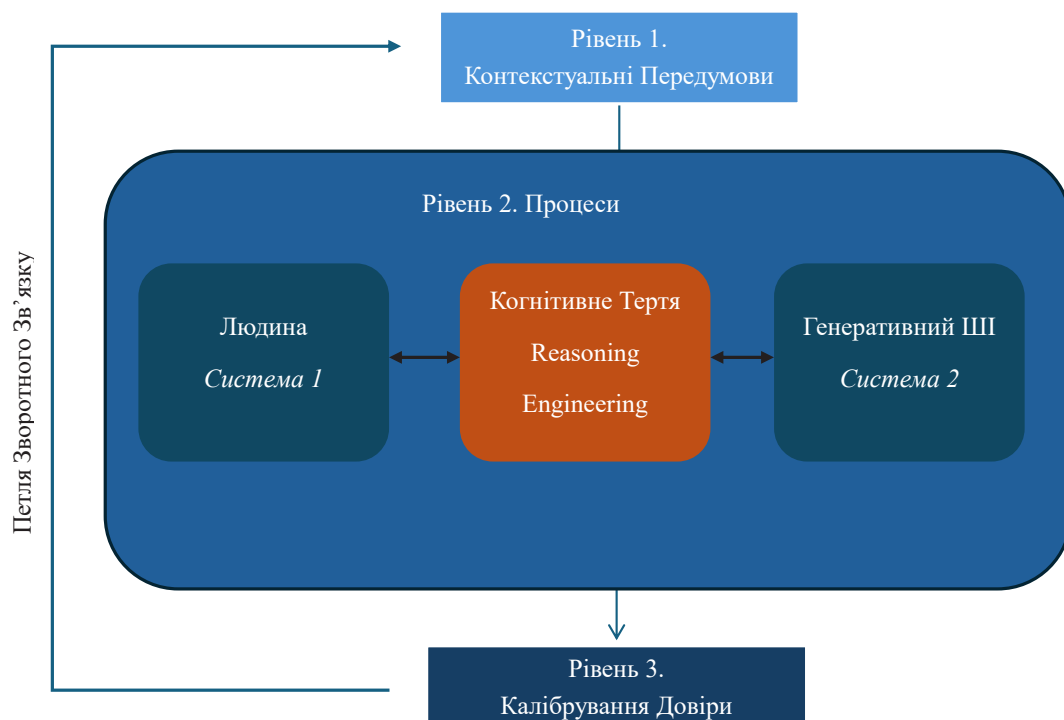


Рис. 1. Механізм синергії рішень: трирівнева модель взаємодії людини та генеративного ШІ

Джерело: сформовано автором на основі [6; 7; 9]

тування у поєднанні зі стратегічною інженерією запитів. Генерація ШІ-контраргументів створює необхідне когнітивне тертя [5], що дозволяє уникнути однобічності поглядів та руйнує упередження підтвердження.

Рівень 3: Результати та зворотний зв'язок. Замикає цикл рівень оцінки ефективності. Якість прийнятого рішення вимірюється не лише бізнес-метриками, але й отриманими когнітивними перевагами – зокрема, успішною оптимізацією ваги рішення, наданої алгоритму [11]. Критично важливим елементом є петля зворотного зв'язку, яка вирішує проблему «Парадоксу довіри». Якщо ШІ-опонент постійно надає якісні контраргументи, існує ризик, що лідер почне сліпо довіряти самому «Адвокату диявола», запускаючи новий виток Automation Bias. Тому результати аналізу ініціюють процес динамічного калібрування довіри: система вимагає регулярної зміни стратегій промптингу та рівня «когнітивного тертя» залежно від типу рішення, запобігаючи як недовірі, так і вторинній сліпій довірі до ШІ опонента [12, 23].

На основі розробленої архітектури Механізму Синергії Рішень (DSM) та інтеграції соціокогнітивних теорій, дедуктивно виведено три теоретичні пропозиції, які операціоналізують вплив ШІ-грамотності, контекстуального обґрунтування та стратегій когнітивного тертя на управлінські рішення.

Відповідно до соціокогнітивної моделі довіри, готовність лідера покладатися на автоматизовані системи залежить від його розуміння меж можливостей цих систем. Цифрова грамотність знижує рівень технологічної тривожності та дозволяє лідеру адекватно оцінювати ризики. Це формує первинну довіру, яка є необхідною умовою для того, щоб алгоритмічним рекомендаціям була надана відповідна «вага» у процесі прийняття рішень.

- Пропозиція 1 (P1): Рівень ШІ-грамотності лідера має перевернуту U-подібну залежність із точністю розподілу ваги рішення. Низька грамотність призводить до упередженості проти програм та відторгнення системи, тоді як надмірна технічна самовпевненість без належного критичного мислення стимулює упередження автоматизації. Оптимальний баланс досягається на середньо-високому рівні AI Fluency, де формується культура експертного нагляду.

Традиційні виходи ШІ (стандартний алгоритмічний Output або технічний XAI) часто не здатні пояснити бізнес-логіку рекомендації, що викликає скептицизм у менеджерів. Натомість,

впровадження AI Reasoning забезпечує контекстуальне обґрунтування, яке резонує з управлінським досвідом. Це семантичне узгодження підвищує сприйняту корисність системи, що безпосередньо конвертується у вищий рівень довіри та більшу вагу, яку менеджер готовий надати алгоритму.

- Пропозиція 2 (P2): Наявність контекстуального обґрунтування не просто збільшує абсолютну вагу, надану алгоритму, а підвищує точність калібрування довіри. AI Reasoning зближує суб'єктивно сприйняту надійність системи з її реальною аналітичною продуктивністю значно ефективніше, ніж надання стандартних алгоритмічних виходів або базової технічної прозорості.

В умовах високої складності лідери схильні покладатися на інтуїцію (Систему 1), що часто призводить до упередження підтвердження. Використання генеративного ШІ виключно як пасивного джерела інформації може посилити це упередження, оскільки людина підсвідомо шукатиме в алгоритмі підтвердження власних думок. Однак, цілеспрямована генерація контраргументів створює когнітивне тертя, яке примусово активує аналітичну Систему 2, змушуючи лідера переглядати власні гіпотези.

- Пропозиція 3 (P3): Стратегічне використання генеративного ШІ для генерації альтернативних сценаріїв негативно корелює з рівнем упередження підтвердження у лідерів. Цей ефект зниження упередження виступає залежною змінною, яка значно посилюється в умовах високої часової компресії та високої когнітивної складності завдання.

Етична медіація та когнітивна адаптивність. Для валідації Механізму Синергії Рішень та визначення меж його застосовності необхідно проаналізувати граничні умови, за яких автономія ШІ стає критично небезпечною і вимагає обов'язкового втручання лідера. Запропонована модель найкраще ілюструється через два сценарії, що демонструють необхідність етичної медіації та когнітивної адаптивності.

Гранична умова 1: Середовища з високими етичними ризиками (Inherent Bias).

Перша межа автономії ШІ виникає у ситуаціях, де великі мовні моделі реплікують та масштабують приховані соціокультурні упередження, засвоєні з гігантських масивів неструктурованих навчальних даних. На відміну від старих предиктивних моделей машинного навчання, сучасний GenAI демонструє це через «алгоритмічне дзеркало упереджень». Наприклад, дослідження показують, що у сфері рекрутингу моделі класу GPT-4 або Claude при

генеративному аналізі та ранжуванні резюме можуть непомітно дискримінувати кандидатів [20, 21] за лінгвістичними маркерами, іменами або культурним кодом, навіть коли в промпті жорстко задана вимога абсолютної об'єктивності. У парадигмі DSM ці кейси доводять, що ідеально згенерований текст не є синонімом етичної нейтральності. Лідерство тут виступає як механізм етичної медіації, гарантуючи, що результати роботи ШІ проходять через фільтр організаційних цінностей, запобігаючи автоматизації прихованої дискримінації.

Гранична умова 2: Середовища з високою емоційною складністю (Cognitive Adaptability). Друга межа стосується нездатності ШІ до емпатії та контекстуального розуміння людської психології. У сфері управління людськими ресурсами ця проблема яскраво проявляється через феномен «синдрому прихованого вигорання». Алгоритмічні системи моніторингу ефективності схильні інтерпретувати зниження активності працівника виключно як падіння продуктивності, ігноруючи емоційне виснаження або психологічний стрес [22]. Імплікація моделі DSM у цьому сценарії полягає в тому, що ШІ використовується лише як інструмент керованого виявлення, який підсвічує аномалії в даних. Проте фінальне рішення вимагає когнітивної адаптивності лідера, який здатен інтегрувати емпатію та неформалізований контекст у процес прийняття рішень, уникаючи хибних каральних дій на основі «сліпої» алгоритмічної логіки.

Ці ілюстративні сценарії підтверджують, що в умовах кризи, невизначеності або етичної складності управлінське судження не може бути повністю делеговане алгоритмам. DSM забезпечує необхідний баланс, де ШІ виступає каталізатором об'єктивності, а лідер – гарантом етичності та адаптивності.

Висновки. Дане дослідження розв'язує критичну проблему когнітивної вразливості управлінців у взаємодії з алгоритмічними системами, пропонуючи концептуальний фреймворк Механізму Синергії Рішень (DSM). Розроблена на основі методології інтегративного синтезу, ця робота здійснює фундаментальну трансформацію парадигми сприйняття генеративного ШІ: від пасивного джерела інфор-

мації («оракула») до його інституціоналізації як активного когнітивного опонента. Розширюючи межі теорії подвійних процесів, дослідження концептуалізує ШІ як екстерналізовану «Систему 2». Головний теоретичний внесок полягає у розв'язанні «Парадоксу довіри»: запропонована трирівнева архітектура DSM доводить, що створення цілеспрямованого «когнітивного тертя» не лише руйнує первинний автоматизм прийняття рішень, але й, через механізми динамічного калібрування, запобігає вторинній сліпій довірі до самого ШІ-опонента.

З практичної точки зору, результати дослідження вимагають перегляду корпоративних стратегій впровадження GenAI. По-перше, організаціям слід відійти від покладання виключно на технічну прозорість на користь інструментів AI Reasoning. Проте, враховуючи критичний ризик «переконливих галюцинацій», ці інструменти мають використовуватись виключно в рамках стратегії «Адвоката диявола» через структуровану інженерію запитів. По-друге, розвиток цифрової грамотності лідерів має бути переосмислений: метою є не максимізація технологічної самовпевненості, яка веде до Automation Bias, а досягнення оптимального балансу для забезпечення точного етичного нагляду та виявлення прихованих алгоритмічних упереджень.

Незважаючи на значний теоретичний внесок, дане дослідження має обмеження, зумовлені його концептуальним характером. Ключовим вектором для майбутніх наукових розвідок має стати емпірична валідація розробленої моделі. Дослідникам пропонується протестувати сформульовані теоретичні пропозиції за допомогою експериментальних дизайнів: зокрема, емпірично виміряти перевернуту U-подібну залежність між AI Fluency та розподілом ваги рішень (P1), оцінити точність калібрування довіри при використанні AI Reasoning (P2), а також дослідити, як фактори середовища (наприклад, часова компресія та когнітивна складність завдання) модерують ефективність ШІ-опонента (P3). Крім того, перспективним напрямком є крос-галузовий аналіз моделі DSM (зокрема, у сферах рекрутингу та охорони здоров'я) для виявлення специфічних граничних умов її застосування.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ:

1. Wen Y., Wang J., Chen X. Trust and AI Weight: Human-AI Collaboration in Organizational Management Decision-Making. *Frontiers in Organizational Psychology*. 2025. Vol. 3. Art. 1419403. DOI: <https://doi.org/10.3389/forgp.2025.1419403>.

2. Li H., Tian F. Advancing Decision-Making through AI-Human Collaboration: A Systematic Review and Conceptual Framework. 2025. DOI: <https://doi.org/10.21203/rs.3.rs-6885768/v1>.
3. Albashrawi M. Generative AI for Decision-Making: A Multidisciplinary Perspective. *Journal of Innovation & Knowledge*. 2025. Vol. 10. № 4. Art. 100751. DOI: <https://doi.org/10.1016/j.jik.2025.100751>.
4. Bertoncini A. L., Matsushita R., Da Silva S. AI, Ethics, and Cognitive Bias: An LLM-Based Synthetic Simulation for Education and Research. *AI in Education*. 2025. Vol. 1. № 1. Art. 3. DOI: <https://doi.org/10.3390/aieduc1010003>.
5. Buçinca Z., Malaya M. B., Gajos K. Z. To Trust or to Think: Cognitive Forcing Functions Can Reduce Overreliance on AI in AI-Assisted Decision-Making. *Proceedings of the ACM on Human-Computer Interaction*. 2021. Vol. 5. № CSCW1. P. 1–21. DOI: <https://doi.org/10.1145/3449287>.
6. Hybrid Intelligence / D. Dellermann et al. *Business & Information Systems Engineering*. 2019. Vol. 61. № 5. P. 637–643. DOI: <https://doi.org/10.1007/s12599-019-00595-2>.
7. Evans J. St. B. T., Stanovich K. E. Dual-Process Theories of Higher Cognition: Advancing the Debate. *Perspectives on Psychological Science*. 2013. Vol. 8. № 3. P. 223–241. DOI: <https://doi.org/10.1177/1745691612460685>.
8. As'ad M., Faran N., Joharji H. AI-Supported Shared Decision-Making [AI-SDM]: Conceptual Framework. *JMIR AI*. 2025. Vol. 4. Art. e75866. DOI: <https://doi.org/10.2196/75866>.
9. Hoff K. A., Bashir M. Trust in Automation: Integrating Empirical Evidence on Factors That Influence Trust. *Human Factors*. 2015. Vol. 57. № 3. P. 407–434. DOI: <https://doi.org/10.1177/0018720814547570>.
10. Pramod D. Decoding Responsible AI Use: The Influence of Digital Literacy and Ethical Awareness. *Cogent Education*. 2025. Vol. 12. № 1. Art. 2592371. DOI: <https://doi.org/10.1080/2331186X.2025.2592371>.
11. Logg J. M., Minson J. A., Moore D. A. Algorithm Appreciation: People Prefer Algorithmic to Human Judgment. *Organizational Behavior and Human Decision Processes*. 2019. Vol. 151. P. 90–103. DOI: <https://doi.org/10.1016/j.obhdp.2018.12.005>.
12. Toward a Science of Human-AI Teaming for Decision Making: A Complementarity Framework / C. Gonzalez et al. *PNAS Nexus*. 2026. Vol. 5. № 3. Art. pgag030. DOI: <https://doi.org/10.1093/pnasnexus/pgag030>.
13. Amplifying Minority Voices: AI-Mediated Devil's Advocate System for Inclusive Group Decision-Making / S. Lee et al. *Companion Proceedings of the 30th International Conference on Intelligent User Interfaces*. Cagliari, Italy : ACM, 2025. P. 17–21. DOI: <https://doi.org/10.1145/3708557.3716334>.
14. Enhancing AI-Assisted Group Decision Making through LLM-Powered Devil's Advocate / C.-W. Chiang et al. *Proceedings of the 29th International Conference on Intelligent User Interfaces*. Greenville, SC, USA : ACM, 2024. P. 103–119. DOI: <https://doi.org/10.1145/3640543.3645199>.
15. Hsu W.-C., Wang M.-C., Ting H.-I. Reducing Emotional Bias in Investment Decisions: The Role of GPT-4 in Financial Analysis. *Asia-Pacific Journal of Business Administration*. 2025. P. 1–26. DOI: <https://doi.org/10.1108/APJBA-03-2025-0181>.
16. Bastone A., Mandiello A., Zeuli F. Generative Artificial Intelligence and Decision-Making: Evidence from a Participant Observation with Latent Entrepreneurs. *European Journal of Innovation Management*. 2026. P. 1–21. DOI: <https://doi.org/10.1108/EJIM-03-2025-0388>.
17. AI-Driven Leadership: Decision-Making, Competencies, and Ethical Challenges – A Systematic Review / A. Sacavém et al. *Administrative Sciences*. 2026. Vol. 16. № 4. Art. 173. DOI: <https://doi.org/10.3390/admsci16040173>.
18. Snyder H. Literature Review as a Research Methodology: An Overview and Guidelines. *Journal of Business Research*. 2019. Vol. 104. P. 333–339. DOI: <https://doi.org/10.1016/j.jbusres.2019.07.039>.
19. Torraco R. J. Writing Integrative Literature Reviews: Using the Past and Present to Explore the Future. *Human Resource Development Review*. 2016. Vol. 15. № 4. P. 404–428. DOI: <https://doi.org/10.1177/1534484316671606>.
20. Dressel J., Farid H. The Accuracy, Fairness, and Limits of Predicting Recidivism. *Science Advances*. 2018. Vol. 4. № 1. Art. eaao5580. DOI: <https://doi.org/10.1126/sciadv.aao5580>.
21. Dissecting Racial Bias in an Algorithm Used to Manage the Health of Populations / Z. Obermeyer et al. *Science*. 2019. Vol. 366. № 6464. P. 447–453. DOI: <https://doi.org/10.1126/science.aax2342>.
22. Porkodi S., Cedro T. L. The Ethical Role of Generative Artificial Intelligence in Modern HR Decision-Making: A Systematic Literature Review. *European Journal of Business and Management Research*. 2025. Vol. 10. № 1. P. 44–55. DOI: <https://doi.org/10.24018/ejbmr.2025.10.1.2535>.
23. Okamura K., Yamada S. Adaptive Trust Calibration for Human-AI Collaboration. *PLoS ONE*. 2020. Vol. 15. № 2. Art. e0229132. DOI: <https://doi.org/10.1371/journal.pone.0229132>.

Дата надходження статті: 24.06.2026

Дата прийняття статті до публікації: 31.07.2026

Дата публікації статті: 05.08.2026