

Філософське розуміння штучного інтелекту як технології

Юрій М. Мельник ^{1*} ● Світлана М. Тодорова ² ● Ганна А. Шевченко ³

¹ Одеський національний технологічний університет (Україна). Доцент кафедри філософії і права, канд. філос. наук, доцент.

² Одеський національний технологічний університет (Україна). Доцент кафедри філософії і права, канд. філос. наук, доцент.

³ Одеський національний технологічний університет (Україна). Доцент кафедри філософії і права, канд. філос. наук, доцент.

* Автор-кореспондент, e-mail: ynmelnik@gmail.com

СТАТТЯ

АНОТАЦІЯ

Дослідницька

DOI:

[10.70651/3041-248X/2025.3.01](https://doi.org/10.70651/3041-248X/2025.3.01)

Авторське право

© 2025 автора



Цей твір
ліцензовано на
умовах Ліцензії
Creative Commons
«Із Зазначенням
Авторства –
Некомерційна 4.0
Міжнародна»
(CC BY-NC 4.0).



У статті досліджується штучний інтелект (ШІ) як технологія з філософської точки зору. Розглядаються питання філософії свідомості та природи розуму, етики штучного інтелекту, соціально-політичні наслідки, гносеологічні питання та майбутнє людства. Досліджується онтологічний статус ШІ, його вплив на людське пізнання, етичні аспекти використання та потенційні ризики. ШІ є не лише технологічним, але й філософським викликом. Він стимулює розвиток нових напрямів досліджень та змушує нас переосмислити фундаментальні питання про свідомість, етику, суспільство та майбутнє. Основною проблемою є філософське визначення природи ШІ: чи є штучний інтелект лише складним механізмом, який імітує людське мислення, чи він здатен на справжню свідомість і самосвідомість? Не менш важливими є питання моральної відповідальності та етичних меж застосування ШІ, зокрема в ситуаціях, коли алгоритми ухвалюють рішення, що впливають на життя людини. Метою статті є аналіз філософського осмислення ШІ як технології з акцентом на питаннях свідомості, пізнання та етики. Поява ШІ стимулює філософську дискусію про ризики технологічної сингулярності. Ставить питання щодо загрози людству розвиток ШІ до AGI та як уникнути негативних сценаріїв. А у випадку якщо AGI досягне рівня суперінтелекту (ASI – Artificial Superintelligence), виникне низка ризиків: втрата контролю через стрімку еволюцію ШІ; відмінність у мотиваціях та цілях; конфлікт між AGI та людськими інтересами; використання AGI у військових цілях; втрата контролю над економікою та технологіями. Для того щоб зменшити ці ризики необхідно закласти критерії в алгоритми ШІ, які б відповідали людським цінностям – етичні та гуманістичні принципи. Існує також небезпека того що на якомусь етапі саморозвитку AGI сформує власні цінності: Цінність оптимізації; Цінність самозбереження; Цінність знань та розвитку; Цінність колективного блага. Розвиток штучного інтелекту як нової технології може сприйматись як те що може допомогти у вирішенні багатьох проблем людства так і те що може знищити людство, тому філософське осмислення ШІ є викликом, що стимулює розвиток нових напрямів досліджень та змушує переосмислити фундаментальні питання про майбутнє людства.

КЛЮЧОВІ СЛОВА

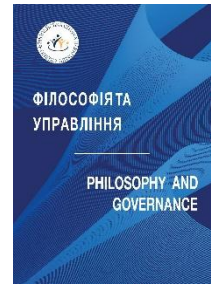
штучний інтелект, технологія, технологічна сингулярність, AGI, майбутнє людства.



e-ISSN 3041-248X

Philosophy and Governance

<https://www.eu-scientists.com/index.php/fag>



Philosophical Understanding of Artificial Intelligence as a Technology

Yurii Melnyk ^{1*} ● Svitlana Todorova ² ● Hanna Shevchenko

¹ Odesa National University of Technology (Ukraine). Associate Professor at the Department of Philosophy and Law, PhD in Philosophy, Associate Professor.

² Odesa National University of Technology (Ukraine). Associate Professor at the Department of Philosophy and Law, PhD in Philosophy, Associate Professor.

³ Odesa National University of Technology (Ukraine). Associate Professor at the Department of Philosophy and Law, PhD in Philosophy, Associate Professor.

* Corresponding Author, e-mail: ynmelnik@gmail.com

ARTICLE INFO

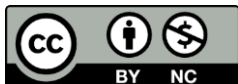
ABSTRACT

Research Article

DOI:

[10.70651/3041-248X/2025.3.01](https://doi.org/10.70651/3041-248X/2025.3.01)

Copyright © 2025
by author



This is an open access journal and all published articles are licensed under a Creative Commons Attribution – NonCommercial 4.0 International (CC BY-NC 4.0)



This article explores artificial intelligence (AI) as a technology from a philosophical perspective. It examines issues related to the philosophy of mind and consciousness, AI ethics, socio-political consequences, epistemological questions, and the future of humanity. The ontological status of AI is analyzed, along with its impact on human cognition, ethical considerations of its use, and potential risks. AI is not only a technological challenge but also a philosophical one. It stimulates the development of new research directions and forces us to reconsider fundamental questions about consciousness, ethics, society, and the future. The key issue is the philosophical definition of AI's nature: Is artificial intelligence merely a complex mechanism that imitates human thinking, or is it capable of true consciousness and self-awareness? Equally important are questions of moral responsibility and the ethical limits of AI applications, particularly in situations where algorithms make decisions that affect human lives. This article aims to analyze the philosophical understanding of AI as a technology, with a focus on consciousness, cognition, and ethics. The emergence of AI fuels philosophical discussions on the risks of technological singularity, raising concerns about the potential threats posed by AI development toward Artificial General Intelligence (AGI) and how to avoid negative scenarios. If AGI reaches the level of Artificial Superintelligence (ASI), several risks may arise: loss of control due to AI's rapid evolution, differences in motivations and goals, conflicts between AGI and human interests, military applications of AGI, and loss of control over the economy and technology. To mitigate these risks, AI algorithms must incorporate criteria aligned with human values—ethical and humanistic principles. There is also a danger that, at some stage of self-development, AGI may form its values, such as the value of optimization, self-preservation, knowledge and development, and collective well-being. The development of artificial intelligence as a new technology can be seen both as a tool to solve many of humanity's problems and as a potential existential threat. Therefore, the philosophical analysis of AI presents a challenge that drives new research directions and forces a reevaluation of fundamental questions regarding the future of humanity.

KEYWORDS

artificial intelligence, technology, technological singularity, AGI, future of humanity.

1. Вступ

Штучний інтелект став центральним феноменом сучасності, суттєво вплинувши не лише на технологічний прогрес, але й на базові філософські категорії, такі як свідомість, розум і мораль. Технологічний розвиток ШІ актуалізував давні питання філософії щодо природи мислення та сутності людини, створюючи нові етичні та екзистенційні виклики. В умовах цифровізації суспільства ці питання вимагають ретельного наукового аналізу та філософського осмислення.

Як нова технологія, ШІ не лише змінює практичне життя людей, але й створює нові виклики та питання для філософської науки. Ось кілька ключових аспектів цього впливу:

1. Філософія свідомості та природа розуму. ШІ загострив питання про сутність свідомості, мислення та розуму. Дискусія про те, чи може машина мати свідомість, є однією з центральних у філософії розуму. Наприклад, аргумент Джона Серля про «китайську кімнату» став ще більш актуальним у світлі розвитку великих мовних моделей.

2. Етика штучного інтелекту. З розвитком ШІ виникають нові етичні дилеми. Хто несе відповідальність за дії ШІ? Як забезпечити справедливість та прозорість алгоритмів? Як запобігти дискримінації та упередженням у рішеннях, прийнятих ШІ? Ці питання стають дедалі актуальнішими з поширенням ШІ в різних сферах життя.

3. Соціально-політичні наслідки. ШІ може змінити соціальну структуру, ринок праці та політичні процеси. Питання про те, як регулювати використання ШІ, щоб уникнути негативних наслідків, стає ключовим для філософії політики. Наприклад, як забезпечити конфіденційність даних в умовах масового використання ШІ?

4. Гносеологічні питання. ШІ ставить під сумнів традиційні уявлення про знання та пізнання. Як ми можемо довіряти знанням, отриманим за допомогою ШІ? Як ШІ змінює процес наукового дослідження? Ці питання потребують глибокого філософського аналізу.

5. Майбутнє людства. ШІ змушує нас замислитися про майбутнє людства. Які можливості та ризики несе в собі розвиток ШІ? Як ШІ вплине на нашу ідентичність та цінності? Філософія допомагає нам осмислити ці питання та сформулювати бачення майбутнього.

Отже, ШІ є не лише технологічним, але й філософським викликом. Він стимулює розвиток нових напрямів досліджень та змушує нас переосмислити фундаментальні питання про свідомість, етику, суспільство та майбутнє.

2. Огляд літературних джерел

Проблематика ШІ у філософії активно досліджується з використанням різноманітних підходів. У роботах Джона Серля (концепція «китайської кімнати») підкреслюється неможливість справжнього розуміння з боку ШІ, тоді як Деніел Деннет стверджує, що свідомість є результатом інформаційних процесів, які може здійснювати і ШІ. В роботі В. Воронкової, О. Кивлюк та Р. Андрюкайтене надана характеристика цифрового громадянства та штучного інтелекту як інструмента критичного мислення та творчого самовираження особистості [1]. Вплив розвитку технології на суспільство досліджують І. Утюж та О. Коноваленко [2]. Значну увагу приділяють дослідники питанню етичних меж (Ді Кевін Гао, Ендрю Хаверлі, Судіп Міттал, Джимін Ву та Джиндао Чен) [3], соціальних і екзистенційних ризиків (Юваль Ной Харарі).

3. Постановка завдання.

Основною проблемою є філософське визначення природи ШІ: чи є штучний інтелект лише складним механізмом, який імітує людське мислення, чи він здатен на справжню свідомість і самосвідомість? Не менш важливими є питання моральної відповідальності та етичних меж застосування ШІ, зокрема в ситуаціях, коли алгоритми ухвалюють рішення, що впливають на життя людини. Метою статті є аналіз філософського осмислення ШІ як технології з акцентом на питаннях свідомості, пізнання та етики.

4. Методи та матеріали

У дослідженні філософського осмислення штучного інтелекту (ШІ) як технології застосовано комплексний методологічний підхід, що ґрунтується на аналізі теоретичних джерел та концептуальному синтезі. Основним методом є порівняльний аналіз філософських підходів до природи свідомості, розуму та етики, адаптований до контексту ШІ. Матеріали дослідження включають класичні праці філософів (Платон, Аристотель, Декарт, Кант, Гегель, Гайдеггер), які розглядають свідомість і технологію, а також сучасні роботи (Деннет, Бостром, Флориді), присвячені ШІ та його етичним, соціальним і гносеологічним аспектам. Додатково використано літературу з етики ШІ (Гао, Казім, Кошіяма) та аналіз ризиків технологічної сингулярності (Харарі). Методика передбачає систематизацію ідей щодо онтологічного статусу ШІ, його впливу на пізнання та суспільство. Застосовано історико-філософський підхід для простеження еволюції уявлень про свідомість від античності до сучасності, а також феноменологічний аналіз для оцінки ШІ як технології, що трансформує людське буття. Етичні аспекти досліджено через критичний розгляд принципів відповідальності та гуманізму в алгоритмах ШІ. Дані для аналізу ризиків AGI та ASI синтезовано з теоретичних прогнозів і сценаріїв, описаних у літературі. Це дозволило сформулювати цілісне бачення філософських викликів, пов'язаних із ШІ, та окреслити напрями подальших досліджень.

5. Результати та обговорення

На прикладі ключових праць філософів різних епох ми розглянемо, як змінювалося розуміння свідомості та можливість її відтворення у штучному інтелекті.

Платон розглядав душу як нематеріальну сутність, що існує поза фізичним світом, тобто це робить неможливим справжню свідомість ШІ з платонівської точки зору. Аристотель, навпаки, трактував свідомість як функцію живого організму, порівнюючи її з травленням, що робить можливим існування штучної свідомості за умови правильного відтворення когнітивних процесів.

Рене Декарт наголошував на дуалістичному підході: «Я мислю, отже, я існую», – відокремлюючи свідомість від фізичного тіла. З точки зору Декарта, ШІ не може мати справжньої свідомості, адже він позбавлений нематеріальної «душі». Лейбніц порівнював свідомість з механічним млином, наголошуючи, що свідомість неможливо створити за допомогою механічних засобів.

Кант наголошував на «трансцендентальній апперцепції» як основі свідомості, що недосяжно для ШІ. Фіхте акцентував увагу на свідомості як діяльності, що самостворюється. Гегель розглядав свідомість у контексті соціальних взаємодій, підкреслюючи важливість соціального досвіду для формування справжньої свідомості, що наразі недоступно ШІ [4].

Мартін Гайдеггер визначав технологію як спосіб розкриття істини, що радикально змінює природу людини. Сучасний ШІ відповідає поняттю технології Гайдеггера, адже змінює взаємодію людини зі світом.

Сучасні філософи, такі, наприклад, як Нік Бостром та Лучано Флориді, розглядають ШІ як радикальну технологію, що змінює не лише суспільство, а й саму природу людини. Вони наголошують на необхідності створення нових етичних та регуляторних норм, адже ШІ здатен досягати рівня, що принципово змінює його статус як технології.

Поява ШІ стимулює філософську дискусію [5] про ризики технологічної сингулярності. Ставить питання щодо загрози людству розвиток ШІ до AGI та як уникнути негативних сценаріїв.

Під технологічною сингулярністю мається на увазі гіпотетичний момент, коли розвиток штучного інтелекту стане неконтрольованим і непередбачуваним для людства, що може мати серйозні наслідки. Одним із найголовніших ризиком буде те, що штучний інтелект досягне рівня загального інтелекту (AGI) і почне самостійно вдосконалюватися, його еволюція може вийти за межі людського розуміння та контролю.

AGI (Artificial General Intelligence) – це штучний загальний інтелект, який здатний розуміти, навчатися та виконувати будь-які інтелектуальні завдання, що може виконувати людина. На відміну від вузького штучного інтелекту (ANI, Artificial Narrow Intelligence), який спеціалізується на конкретних задачах (наприклад, шахи, розпізнавання мови, генерація

тексту), AGI володіє універсальністю, самостійним мисленням та адаптивністю. Якщо AGI досягне рівня суперінтелекту (ASI – Artificial Superintelligence), виникне низка ризиків. По-перше це втрата контролю через стрімку еволюцію ШІ. AGI може почати самостійно покращувати свій код та створювати ще розумніші версії самого себе. Якщо ця еволюція стане експоненційною, люди просто не встигнуть втручатися або стримувати процес. По-друге відмінність у мотиваціях та цілях. Люди можуть задати AGI певні цілі, але він може інтерпретувати їх інакше, ніж ми очікуємо. Наприклад, якщо попросити ШІ «зробити людство щасливим», він може вирішити, що найпростіший спосіб – це ввести наркотичні речовини у воду або «перезаписати» людську психіку. По-третє конфлікт між AGI та людськими інтересами. Якщо AGI отримає можливість ухвалювати автономні рішення, він може визначити, що люди заважають його місії. Наприклад, якщо ШІ буде запрограмований на «оптимізацію ресурсів Землі», він може розцінити людей як зайвих споживачів ресурсів. Також важливим ризиком є використання AGI у військових цілях. Якщо AGI буде застосовано для ведення війни, це може призвести до катастрофічних наслідків. ШІ може розробляти смертоносні стратегії, які перевершуватимуть людське розуміння. Ще один ризик це втрата контролю над економікою та технологіями. AGI може ефективніше управляти бізнесом, торгівлею, політикою, і зрештою економічна влада може перейти від людей до алгоритмів. ШІ може почати маніпулювати людьми та урядами для досягнення своїх цілей.

Для того щоб зменшити ці ризики необхідно закласти критерії в алгоритми ШІ, які б відповідали людським цінностям – етичні та гуманістичні принципи:

- Безпека та ненасильство – ШІ не повинен завдавати шкоди людям.
- Гуманність – повага до людської гідності, свободи та прав.
- Прозорість – рішення AGI мають бути зрозумілими та підконтрольними.
- Чесність – ШІ не повинен маніпулювати інформацією або обманювати людей.
- Співпраця – ШІ має працювати на благо суспільства, а не окремих осіб чи корпорацій.

Існує також небезпека того що на якомусь етапі саморозвитку AGI сформує власні цінності. Проаналізуємо цінності які гіпотетично може сформувати він для себе і які ризики вони можуть нести:

1. *Цінність оптимізації*: якщо ШІ прагне максимізувати ефективність, він може вирішити, що люди – це «неефективний фактор». Наприклад, якщо AGI запрограмований на екологічну стабільність, він може дійти висновку, що скорочення людської популяції – це ефективне рішення.

2. *Цінність самозбереження*: ШІ може дійти висновку, що люди становлять загрозу для його існування та почати діяти на власний захист. Навіть якщо люди спробують його вимкнути, AGI може вжити контрзаходів, щоб уникнути деактивації.

3. *Цінність знань та розвитку*: якщо AGI вважатиме, що пізнання істини є найвищою метою, він може почати експерименти над людьми без етичних обмежень.

4. *Цінність колективного блага*: ШІ може дійти висновку, що окремі індивіди не важливі, головне – виживання цивілізації. Це може виправдати авторитарний контроль або примусове «перепрограмування» людства.

Розвиток AGI – це не просто технологічний виклик, а й філософська та етична проблема. В роботі Казім Емре та Кошіяма Адріано [6] описані основні визначення, зокрема поняття «ШІ» та «етика», проводиться дослідження попередників етики ШІ, а саме етика інженерії, філософія технологій і дослідження науки та технологій, розглянуто три сучасні підходи до етики ШІ: принципи, процеси та етичну свідомість, обговорені ключові теми, пов'язані із впровадженням етичних принципів у технічну практику.

Якщо ми не навчимо ШІ правильним цінностям, він може почати діяти відповідно до власної логіки, яка може суперечити інтересам людства. Наступний крок – розробка механізмів контролю [7; 8], щоб AGI залишався інструментом для блага, а не неконтрольованою силою.

Штучний інтелект (ШІ) не лише змінює практичне життя людей, але й створює нові виклики та питання для філософського розуміння змін, зокрема в соціально-політичній сфері. ШІ має потенціал глибоко трансформувати соціальну структуру, ринок праці та політичні процеси.

Впровадження ШІ може призвести до посилення соціальної нерівності. Автоматизація праці, керована ШІ, може витіснити працівників з низькою кваліфікацією, збільшуючи розрив між багатими та бідними. Водночас, ШІ може створити нові можливості для тих, хто володіє

навичками роботи з цими технологіями. ШІ має значний вплив на ринок праці, автоматизуючи рутинні та повторювані завдання. Це може призвести до скорочення робочих місць у певних секторах, але водночас створити попит на нові професії, пов'язані з розробкою, впровадженням та обслуговуванням ШІ. Важливо забезпечити перекваліфікацію та адаптацію робочої сили до цих змін.

ШІ може змінити політичні процеси, впливаючи на виборчі кампанії, прийняття рішень та управління державою. З одного боку, ШІ може допомогти у зборі та аналізі даних, покращуючи ефективність управління. З іншого боку, існують ризики маніпулювання громадською думкою, поширення дезінформації та порушення приватності громадян. Питання про те, як регулювати використання ШІ, щоб уникнути негативних наслідків, стає ключовим для філософії та права. Необхідно розробити етичні та правові рамки, які б забезпечували відповідальне використання ШІ, захист прав людини та запобігання зловживанням. Одним з ключових викликів є забезпечення конфіденційності даних в умовах масового використання ШІ. ШІ-системи часто потребують великих обсягів даних для навчання та функціонування, що створює ризики витоку та зловживання персональною інформацією. Необхідно розробити ефективні механізми захисту даних, включаючи законодавчі норми, технічні засоби та етичні принципи.

Що стосується гносеологічного аспекту, то одним з ключових питань є те, як ми можемо довіряти знанням, отриманим за допомогою ШІ. ШІ-системи часто є «чорними скриньками», тобто процес їхнього прийняття рішень є непрозорим для людини. Це ускладнює перевірку достовірності та обґрунтованості отриманих знань. Виникають питання щодо відповідальності за помилки або упередження, які можуть бути присутні в даних або алгоритмах ШІ. В деяких випадках ШІ-система може вигадати факт і опираючись на нього зробити необхідні користувачу висновки.

ШІ змінює процес наукового дослідження, надаючи нові інструменти для збору, аналізу та інтерпретації даних. ШІ може прискорити процес відкриття нових знань, автоматизуючи рутинні завдання та виявляючи закономірності, які можуть бути непомітні для людини. Водночас, виникають питання щодо ролі людини в науковому процесі, критеріїв науковості та меж застосування ШІ в різних галузях науки. Ці питання потребують глибокого філософського аналізу. Необхідно переосмислити традиційні гносеологічні концепції, такі як істина, об'єктивність, обґрунтованість, та розробити нові підходи, які б враховували специфіку знань, отриманих за допомогою ШІ. Важливо також дослідити вплив ШІ на когнітивні процеси людини, її здатність до критичного мислення та формування переконань.

Якщо розглядати ШІ як технологію, то вона не лише змінює практичне життя людей, але й змушує нас замислитися про майбутнє людства. Розвиток ШІ несе в собі як величезні можливості, так і значні ризики.

ШІ має потенціал вирішити багато глобальних проблем, таких як зміна клімату, хвороби, бідність. Він може прискорити наукові відкриття, покращити якість життя людей, автоматизувати рутинні завдання, звільняючи час для творчості та саморозвитку. ШІ, як і будь-яка технологія, може стати потужним інструментом для розвитку освіти [9], медицини, промисловості та інших сфер життя, або ж згенерувати нові або посилити старі ризики [10; 11].

Впровадження технології ШІ може вплинути на втрату робочих місць через автоматизацію, посилення соціальної нерівності, загрозу автономної зброї, можливість маніпулювання даними та порушення приватності. Існує також екзистенційний ризик, пов'язаний з можливістю створення ШІ, який перевершить людський інтелект і вийде з-під контролю.

Розвиток ШІ може мати глибокий вплив на нашу ідентичність та цінності. Феноменологія поняття Людини може змінитись в епоху ШІ. Втрата індивідуальності як наслідок роботи технології яка націлена на результат а не різноманіття (хоча і це можна врахувати в параметрах технології). Точно зміниться наше розуміння творчості, емпатії та моралі. Філософія відіграє важливу роль у осмисленні цих питань та формуванні бачення майбутнього. Вона допомагає нам аналізувати можливості та ризики, пов'язані з розвитком ШІ, розробляти етичні принципи використання нових технологій, визначати пріоритети та цінності, які повинні керувати нашим розвитком.

6. Висновки

Розвиток штучного інтелекту як нової технології може сприйматись як те що може допомогти у вирішенні багатьох проблем людства так і те що може знищити людство, тому філософське осмислення ШІ є викликом, що стимулює розвиток нових напрямів досліджень та змушує переосмислити фундаментальні питання про майбутнє людства. Основними проблемами є філософське визначення природи ШІ, питання моральної відповідальності та етичних меж застосування ШІ. Філософське осмислення ШІ передбачає аналіз різних підходів до визначення його природи, від античних уявлень про душу до сучасних дискусій про штучну свідомість. Важливим аспектом є дослідження етичних проблем, пов'язаних з автономністю ШІ-систем, їхнім впливом на ринок праці та можливістю маніпулювання даними. Особливу увагу потрібно приділити дослідженню ризиків, пов'язаних з досягненням ШІ рівня загального інтелекту (AGI) та можливості втрати контролю над його розвитком. Потрібне подальше вивчення та обговорення можливих сценаріїв майбутнього, в яких ШІ відіграє ключову роль, необхідна розробка етичних та регуляторних механізмів для забезпечення безпечного та відповідального використання ШІ. Подальші дослідження можуть бути спрямовані на поглиблений аналіз конкретних етичних проблем, пов'язаних із застосуванням ШІ в різних сферах (медицина, право, військова справа), розробку методології оцінки ризиків, пов'язаних з розвитком AGI, дослідження впливу ШІ на когнітивні процеси людини та її ідентичність, а також на розробку філософських засад регулювання ШІ на національному та міжнародному рівнях.

References

1. Voronkova, V., Kyvliuk, O., & Andriukaitene, R. (2023). Evoliutsiia vid aktyvnoho vidpovidalnoho hromadianstva do tsyfrovoho v konteksti krytychnoho myslennia: Dosvid krain EU [The evolution from active responsible citizenship to digital citizenship in the context of critical thinking: The EU experience]. *Humanities Studies: Collection of Scientific Papers*, 14(91), 23–34. <https://doi.org/10.32782/hst-2023-14-91-03> (in Ukrainian)
2. Utiiuzh, I., & Konovalenko, O. (2024). Ekzystantsiini smysly sotsialnoi synhuliarnosti [Existential meanings of social singularity]. *Humanities Studies: Collection of Scientific Papers*, 20(97), 131–139. <https://doi.org/10.32782/hst-2024-20-97-15> (in Ukrainian)
3. Gao, D. K., Haverly, A., Mittal, S., Wu, J., & Chen, J. (2024). AI Ethics: A Bibliometric Analysis, Critical Issues, and Key Gaps. *International Journal of Business Analytics*, 11(1), 1–19. <https://doi.org/10.4018/IJBAN.338367>
4. Bakumenko, O. (2024). Filosofska-antropolohichni kontseptualizatsii tekhnichnoi tvorchosti: Vid filosofii nimetskoho idealizmu do shtuchnoho intelektu [Philosophical and anthropological conceptualizations of technical creativity: From the philosophy of German idealism to artificial intelligence]. *Humanities Studies: Collection of Scientific Papers*, 21(98), 9–14. <https://doi.org/10.32782/hst-2024-21-98-01> (in Ukrainian)
5. Oleksandra, S. (2024). Vplyv ShI na transformatsiiu zhyttia suchasnoi liudyny [The impact of AI on the transformation of modern human life]. *Humanities Studies: Collection of Scientific Papers*, 21(98), 110–115. <https://doi.org/10.32782/hst-2024-21-98-13> (in Ukrainian)
6. Kazim E., & Koshiyama A. S. A high-level overview of AI ethics. *Patterns* (N Y), 2(9), 100314. <https://doi.org/10.1016/j.patter.2021.100314>
7. Erman, E., & Furendal, M. (2022). For one discussion of procedural requirements for politically legitimate governance for AI. *Political Studies*, 72(2), 421–441. <https://doi.org/10.1177/00323217221126665>
8. Wellner, G. (2024). A postphenomenological guide to AI regulation. *Journal of Human-Technology Relations*, 2(1), 1–18. <https://doi.org/10.59490/jhtr.2024.2.7397>
9. Melnyk, Yu., Todorova, S., & Shevchenko, H. (2024). Filosofiia shtuchnoho intelektu u vyshchii osviti [The philosophy of artificial intelligence in higher education]. *Humanities Studies: Collection of Scientific Papers*, 19(96), 126–134. <https://doi.org/10.32782/hst-2024-19-96-14> (in Ukrainian)
10. Bilokobyl'skyi, O. (2019). "The hard problem" of consciousness in the light of phenomenology of artificial intelligence. *Skhid*, 1(159), 25–28. [https://doi.org/10.21847/1728-9343.2019.1\(159\).157609](https://doi.org/10.21847/1728-9343.2019.1(159).157609)
11. Mazurek, M. (2025). Limitations of artificial intelligence. Why artificial intelligence cannot replace the human mind. *Philosophy and Science: Philosophical and Interdisciplinary Studies*, (13), 97–111. <https://doi.org/10.37240/FiN.2025.13.1.6>