

Філософія

УДК 165:004.8

DOI <https://doi.org/10.5281/zenodo.20746518>

**Про пізнавальні межі штучного інтелекту: здивування і суперечність у
діалектичній перспективі**

Кулик Віталій Олександрович

аспірант кафедри філософії, Національний технічний університет
України «Київський політехнічний інститут імені Ігоря Сікорського», пров.
Індустріальний 1, Кіровоградська обл. м Світловодськ, 27500, Україна,
<https://orcid.org/0009-0006-2471-304X>

Прийнято: 22.05.2026 | Опубліковано: 30.05.2026

***Анотація:** Мета дослідження полягає в обґрунтуванні того, що здатність до здивування є необхідною умовою мислення у діалектико-логічному сенсі, і що трансформерна архітектура великих мовних моделей (LLM) структурно виключає цю здатність через відсутність конститутивних умов здивування як пізнавального акту.*

Методологічною основою дослідження є діалектична логіка у традиції Канта і Гегеля. Застосовується аналіз через категорії визначення, заперечення, суперечності та зняття, які дозволяють описати структуру пізнавального акту незалежно від питання про суб'єктивний досвід. Архітектура трансформера розглядається не як технічний об'єкт, а як реалізація певного способу діяльності з поняттями, який зіставляється з тим, що діалектична логіка описує як мислення. Таке зіставлення стає можливим завдяки тому, що сучасна технічна документація дозволяє концептуально



точно описати операції моделі при генерації. Результати дослідження показують, що здивування є не психологічним станом, а структурним моментом пізнання, що вимагає трьох конститутивних умов. По-перше, наявності стійкого поняття, здатного вступити в конфлікт з дійсністю — на відміну від статистичного розподілу імовірностей, яким оперує трансформер. По-друге, здатності утримувати попередні визначення як зняті моменти у сенсі гегелівського Erinnerung — на відміну від марківського вікна контексту, що щоразу починається з нуля. По-третє, здатності перетворити виявлену суперечність на імпульс руху до більш конкретного поняття — на відміну від нейтралізації суперечності через вибір статистично когерентного продовження. Наукова новизна дослідження полягає у тому, що вперше здійснено систематичний аналіз пізнавальних можливостей LLM засобами діалектичної логіки через поняття здивування як структурного моменту пізнання. На відміну від існуючих підходів, які апелюють до відсутності свідомості або — як у Вудса — визнають у машині часткове логічне визначення, запропонований підхід дозволяє провести якісну, а не лише кількісну межу між обчисленням і мисленням. Показано, що ця межа є іманентною — вона випливає з самого поняття трансформерної архітектури — а не зовнішньою технічною характеристикою. Висновки. Трансформерна архітектура LLM структурно виключає всі три умови здивування, і це виключення є іманентним, а не технічним. Визначальна здатність судження реалізується трансформером з ефективністю і масштабом, недосяжними для людини. Рефлектуюча здатність судження — формування нового загального для даного часткового, що починається зі здивування перед суперечністю — залишається людською задачею не з технічних, а з архітектурно-логічних причин. Це розмежування є адекватною основою для розуміння того, чим є і чим не є сучасний штучний інтелект.

Ключові слова: діалектична логіка, здивування, великі мовні моделі, трансформер, здатність судження, суперечність, *Erinnerung*, штучний інтелект.

On the Cognitive Limits of Artificial Intelligence: Wonder and Contradiction in a Dialectical Perspective

Vitalii Kulyk

PhD Student, Department of Philosophy, Igor Sikorsky Kyiv Polytechnic Institute,
1 Industrialnyi Lane, Svitlovodsk, Kirovohrad Region, 27500, Ukraine,
<https://orcid.org/0009-0006-2471-304X>

Abstract. Objective. *Objective. The study aims to demonstrate that the capacity for wonder is a necessary condition of thought in the dialectical-logical sense, and that the transformer architecture of large language models (LLMs) structurally excludes this capacity through the absence of the constitutive conditions of wonder as a cognitive act. Methods. The methodological basis of the study is dialectical logic in the tradition of Kant and Hegel. The analysis proceeds through the categories of determination, negation, contradiction, and sublation, which allow the structure of the cognitive act to be described independently of questions concerning subjective experience. The transformer architecture is approached not as a technical object but as the realization of a particular mode of activity with concepts, which is compared with what dialectical logic describes as thinking. This comparison is made possible by the fact that contemporary technical documentation allows a conceptually precise description of the model's operations during generation. Results. The study shows that wonder is not a psychological state but a structural moment of cognition that requires three constitutive conditions. First, a stable*

concept capable of entering into conflict with reality — as opposed to the statistical probability distribution with which the transformer operates. Second, the ability to retain previous determinations as sublated moments in the sense of Hegel's Erinnerung — as opposed to a Markovian context window that restarts with each session. Third, the ability to transform the discovered contradiction into an impulse toward a more concrete concept — as opposed to the neutralization of contradiction through the selection of a statistically coherent continuation. Scientific novelty. For the first time, a systematic analysis of the cognitive possibilities of LLMs has been carried out through the lens of dialectical logic, using the concept of wonder as a structural moment of cognition. Unlike existing approaches — which appeal to the absence of consciousness, or which, as in Woods, acknowledge partial logical determination in the machine — the proposed approach allows a qualitative rather than merely quantitative boundary to be drawn between computation and thinking. This boundary is shown to be immanent, following from the very concept of the transformer architecture, rather than an external technical characteristic subject to future correction. Conclusions. The transformer architecture of LLMs structurally excludes all three conditions of wonder, and this exclusion is immanent rather than technical. Determining judgment is realized by the transformer with an efficiency and scale unattainable for any human. Reflecting judgment — the formation of a new universal for a given particular, beginning with wonder at contradiction — remains a human task not for technical but for architectural-logical reasons. This distinction is the adequate basis for understanding what modern artificial intelligence is — and what it is not.

Keywords: *dialectical logic, wonder, large language models, transformer, power of judgment, contradiction, Erinnerung, artificial intelligence.*

Постановка проблеми. На перший погляд питання може здаватися дивним – де здивування, а де штучний інтелект? Зазвичай ставлять інше питання – чи може машина мислити та чи є штучний інтелект – справді інтелектом.

Стрімкий розвиток великих мовних моделей (LLM) — систем на кшталт GPT-4, Claude, Gemini — знову підняв питання про природу штучного інтелекту та його відношення до людського мислення. Переважна більшість цієї дискусії ведеться у термінах свідомості, інтенціональності, суб'єктивного досвіду або втіленості. Проте такий підхід залишає осторонь питання, яке є, можливо, більш фундаментальним: яку структуру має пізнавальний акт цих систем? Чи може, хоча б в теорії, множина навичок ШІ співпадати чи навіть бути більшою за множину розумових здібностей людини? Від відповіді на нього залежить, де саме пролягає межа між задачами, які можна делегувати штучному інтелекту, і задачами, що вимагають людської присутності. Якщо пізнавальні обмеження LLM є технічними — їх можна подолати вдосконаленням архітектури. Якщо вони іманентні самому алгоритму роботи існуючих моделей — їх не можна подолати в принципі в межах існуючої парадигми.

Аналіз останніх досліджень і публікацій. Дослідження сучасних моделей засобами діалектичної логіки залишається малодослідженим напрямом, хоча саме цей підхід дозволяє поставити питання про структуру пізнавального акту точніше, ніж це робить феноменологія або аналітична філософія свідомості.

Серед сучасних досліджень, що торкаються суміжної проблематики, слід відзначити кілька напрямів. Найбільш ґрунтовною є стаття Вудса «Prompt, negate, repeat: a Hegelian meditation on AI» [1], яка стверджує, що LLM здійснює часткове діалектичне визначення, спираючись на піппінівське читання Гегеля, в якому логічний рух можливий без суб'єкта. Однак ця інтерпретація зводить

гегелівське заперечення до формальної схеми, відриваючи його від іманентного самозаперечення поняття. Статистичне виключення токенів є зовнішньою операцією над розподілом ймовірностей, а не внутрішнім рухом поняття через власну суперечність — між ними не кількісна, а якісна різниця. Близькою до нашої проблематики є також стаття Чухров «The Philosophical Disability of Reason» [2], яка показує, що «вимога алгоритмічної логіки полягає в тому, щоб або вирішити, або нейтралізувати парадокс, але не в тому, щоб його екстраполювати» [2, р. 75]. Проте питання про умови можливості здивування — сформульоване ще Кантом через розрізнення визначальної і рефлектуючої здатності судження — залишається у цьому аналізі нерозгорнутим.

Загальний стан дискусії про ментальність LLM систематизує Шевлін [3], який розрізняє три конкуруючі підходи: «безмозкові машини» (архітектурний дебанкінг), функціоналізм і соціальний підхід. Жоден з цих підходів не звертається до структури пізнавального акту в діалектико-логічному сенсі — що і визначає прогалину, яку заповнює дане дослідження.

Паралельний напрям представлений дослідженнями, що аналізують LLM крізь призму філософії мови та когнітивної науки. Шанаган [4] показує, що використання філософськи навантажених термінів — «знає», «вірить», «думає» — стосовно LLM є методологічно некоректним: модель здійснює лише передбачення статистично правдоподібних послідовностей, не маючи поняття істини чи хибності. Міллер і Бакнер [5] розглядають це питання детальніше, аналізуючи, чи є LLM «просто Blockhead» у сенсі Блока — системою, що відтворює завчені відповіді без розуміння — і приходять до висновку, що моделі здатні гнучко комбінувати паттерни, однак межа між цим і справжнім пізнанням залишається принципово нез'ясованою.

Окремий напрям пов'язаний із застосуванням кантівської теорії судження до ШІ. Сео [6] аналізує судження ШІ через кантівську таблицю категорій і показує, що функція SoftMax структурно перетворює всі судження



моделі на проблематичні — судження можливості — що унеможливорює асерторичне або аподиктичне судження в кантівському сенсі.

З боку технічних досліджень, що мають філософське значення, Абдалі та ін. [7] пропонують метод «самовідображення» LLM, натхненний гегелівською діалектикою. Характерно, що автори самі визнають: «LLM не "відображає" так, як це робить людина» — і тим самим підтверджують від протилежного нашу тезу про те, що діалектичне самозаперечення поняття не є властивістю трансформерної архітектури. Ху [8] намагається формалізувати «поняття» для ШІ через теорію алгоритмічної інформації, пропонуючи визначення поняття як структури, що може переглядатися і узгоджуватися між агентами — але сама необхідність такої формалізації свідчить про те, що поняття у філософському сенсі в LLM відсутнє.

З позицій діалектичної логіки та епістемології значущою є стаття Ель Маоша і Цзіна [9], яка застосовує культурно-історичну теорію діяльності Виготського і діалектичну логіку до аналізу ШІ, стверджуючи, що ШІ успадкував епістемологічну кризу психології, засновану на дуалістичній онтології. Автори показують, що емпіристська парадигма, панівна в дослідженнях ШІ, структурно нездатна вирішити проблему формування понять. Суміжну проблему досліджує стаття про автоматизацію епістемології [10], яка показує, що те, що LLM видає за «знання», є статистично закодованим відображенням соціальних і культурних структур, а не результатом пізнавального акту.

Таким чином, невирішеною залишається задача систематичного застосування діалектичної логіки до аналізу LLM через поняття здивування як ключового маркера відмінності між визначальним і рефлексуючим судженням.

Мета цього дослідження — обґрунтувати, що здатність до здивування є необхідною умовою мислення у діалектико-логічному сенсі, і що

трансформерна архітектура LLM структурно виключає цю здатність через відсутність трьох конститутивних умов: стійкого поняття, здатності утримувати попередні визначення як зняті моменти і здатності до зняття суперечності. Для досягнення цієї мети необхідно: реконструювати поняття здивування як пізнавальної структури в його розвитку від Арістотеля до Гегеля — від першої постановки питання до його діалектичного розгортання; описати архітектуру трансформера концептуально в термінах, що дозволяють зіставити її з виявленими логічними умовами; і показати, що відсутність кожної з трьох умов є структурною, а не технічною.

Методологічною основою дослідження є діалектична логіка у традиції Канта і Гегеля. Конкретно — аналіз через категорії визначення, заперечення, суперечності та зняття, які дозволяють описати структуру пізнавального акту незалежно від питання про суб'єктивний досвід. Архітектура трансформера розглядається не як технічний об'єкт, а як реалізація певного способу діяльності з поняттями — і саме цей спосіб діяльності зіставляється з тим, що діалектична логіка описує як мислення. Таке зіставлення стає можливим завдяки тому, що сучасна технічна документація дозволяє концептуально точно описати, що саме робить модель при генерації.

Наукова новизна. Питання про здивування як пізнавальну структуру має тривалу філософську традицію — від арістотелівської постановки через кантівський аналіз умов його можливості до гегелівського розуміння суперечності як рушія мислення. Однак це питання досі не ставилося у зв'язку з аналізом конкретної архітектури сучасних LLM. Так само залишалося нерозгорнутим питання про те, які саме умови необхідні для того, щоб здивування як пізнавальний акт було взагалі можливим — і чи присутні ці умови в трансформерній архітектурі. Заповнення цієї прогалини і є предметом даного дослідження.

Виділення невирішених раніше частин загальної проблеми. Попри зростаючий інтерес до філософського осмислення штучного інтелекту, питання про структуру пізнавального акту LLM залишається нерозв'язаним. Наявні дослідження або зупиняються на рівні феноменологічної критики — апелюючи до відсутності свідомості, втіленості чи інтенціональності — або, як у випадку Вудса, визнають у машині часткове логічне визначення, не розгортаючи питання про умови його можливості. В обох випадках поза увагою залишається те, що саме робить пізнавальний акт мисленням у діалектичному сенсі — а не лише обчисленням або симуляцією.

Зокрема, невирішеним залишається питання про здивування як структурний момент пізнання. Жодна з розглянутих робіт не ставить його в центр аналізу і не виводить з нього систематичного діагнозу архітектури трансформера. Між тим саме здивування — як момент виявлення недостатності наявного поняття — є точкою, в якій визначальна здатність судження переходить у рефлектуючу, формальне обчислення — у мислення, а протиріччя стає наявним, усвідомленим.

Виклад основного матеріалу дослідження.

Здивування як пізнавальна структура

Здивування — не психологічна реакція і не риторичний прийом. У філософській традиції від Арістотеля до Гегеля воно є структурним моментом пізнання: точкою, в якій пізнавальна система виявляє межі власного розуміння і отримує імпульс для руху вперед. Арістотель формулює це у «Метафізиці»: «Саме через здивування люди і тепер починають, і спочатку почали філософствувати — дивуючись спочатку очевидним труднощам, а потім поступово ставлячи питання про більш значні речі. Той, хто дивується і збентежений, відчуває своє незнання; тому любитель міфів є в певному сенсі філософом, бо міф складається з дивовижного» [11, 982b]. Важливо, що

Арістотель характеризує здивування через відчуття незнання. Але це специфічне незнання — не просто відсутність інформації, а виявлення того, що наявне знання є недостатнім. Щоб дивуватися, потрібно мати поняття, яке може виявитися недостатнім.

Кант аналізує умови можливості цього структурного моменту з максимальною точністю. У «Критиці здатності судження» він вводить розрізнення, принципове для нашого аналізу: «Здатність судження взагалі є здатністю мислити особливе як таке, що міститься під загальним. Якщо загальне (правило, принцип, закон) дано, то здатність судження, яка підводить під нього особливе, є визначальною. Якщо ж дано лише особливе, для якого потрібно знайти загальне, то здатність судження є лише рефлектуючою» [12, р. 66]. Визначальна здатність судження — це застосування відомого правила до конкретного випадку. Загальне вже є; потрібно лише підвести часткове під нього. Рефлектуюча — принципово інша операція: загального ще немає, є лише конкретний випадок, що не вкладається в наявні категорії, і потрібно знайти нове загальне. Саме в цьому — виявленні того, що жодне з наявних загальних не підходить — і полягає здивування як пізнавальний акт.

Гегель поглиблює цей аналіз, показуючи, що здивування є внутрішнім моментом самого мислення. У «Малій логіці» він описує, як думка «мислення заплутується в суперечностях, тобто, губиться в постійній нетотожності думки й, таким чином, не доходить до самого себе [...] Розуміння того, що діалектика являє природу самого мислення, що як розсудок воно має впасти в заперечення самого себе, в суперечність, розуміння цього становить одну з головних сторін логіки» [13, с. 50]. Суперечність тут — не помилка і не дефект мислення, а нормальний і необхідний момент його руху. У «Великій науці логіки» Гегель формулює: «Щось є живим лише тією мірою, якою воно містить у собі суперечність, і більш того — є цією силою утримувати і витримувати суперечність у собі» і «Спекулятивне мислення полягає виключно в тому, що

думка міцно тримається суперечності і в ній — самої себе, але не дозволяє суперечності панувати над собою, як у звичайному мисленні, де її визначення вирішуються суперечністю лише в інші визначення або в ніщо» [14, р. 439].

Тепер звернемося до архітектури трансформера. Сучасні великі мовні моделі засновані на архітектурі, запропонованій Васвані та ін. у 2017 році: «Ми пропонуємо нову просту мережеву архітектуру — трансформер, засновану виключно на механізмах уваги, повністю відмовляючись від рекурентності та згорток» [15, р. 5998]. Модель вирішує одну задачу — передбачити наступний токен на основі всіх попередніх. Центральний механізм — само-увага: «Само-увага — це механізм уваги, що пов'язує різні позиції однієї послідовності для обчислення її представлення» [15, р. 6]. Кожен токен статистично зважує релевантність усіх інших у контексті і на основі цього зваженого представлення модель виробляє розподіл імовірностей по всьому словнику. Goodfellow, Bengio та Courville описують загальний принцип: «Глибоке навчання досягає великої потужності та гнучкості шляхом представлення світу у вигляді вкладеної ієрархії понять, де кожне поняття визначається відносно простіших понять» [16, р. 8]. Але ці «поняття» є статистичними паттернами, індукованими з трильйонів токенів, а не поняттями у філософському сенсі.

Якщо перекласти це мовою Канта: трансформер завжди підводить новий контекст під вже вивчені паттерни. Загальне завжди вже дане — у вигляді ваг і розподілів. Механізму виявити, що жоден паттерн не підходить і потрібно сформувати нове загальне, в архітектурі немає. Це визначальна здатність судження в чистому архітектурному вигляді — без рефлектуючої.

Перевіримо тепер, чи присутні у трансформері ті моменти, без яких здивування як пізнавальний акт неможливе.

Арістотель характеризує здивування через відчуття незнання — але специфічного: не відсутності інформації, а виявлення того, що наявне уявлення про предмет діяльності недостатнє. Це передбачає це саме уявлення,

яке може виявитися недостатнім. У трансформері замість уявлення — розподіл імовірностей. Коли модель зустрічає статистично малоімовірний токен, це є відхиленням від вивченого розподілу, але не конфліктом уявлення з дійсністю. Вчений дивується аномальному результату тому, що його теорія передбачала інший — теорія і є те поняття, яке входить у конфлікт. У моделі немає теорій. Аномальний токен — просто точка з низькою вагою, а не сигнал про недостатність знань. Звідси — структурний, а не випадковий характер галюцинацій: за відсутності вірного паттерну модель генерує найправдоподібніше продовження, не маючи механізму, який сигналізував би про власну некомпетентність.

Але навіть якби модель мала щось подібне до уявлення — дивуватися їй було б нікому. Здивування передбачає суб'єкта, у якого є історія власних визначень. Гегелівське *Erinnerung* — ретроспективне збирання минулих визначень у більш конкретне поняття — є не психологічною пам'яттю, а логічною структурою розвитку мислення. Саме тому «Становлення є першою конкретною думкою і, отже, першим поняттям, буття ж і ніщо суть, навпаки, порожні абстракції» [13, с. 167]: становлення багатше не кількісно, а якісно — воно містить буття і ніщо як зняті моменти. Трансформер позбавлений цієї структури. Вудс точно описує це через протиставлення двох темпоральностей: «темпоральність зняття, в якій думка ретроспективно збирає себе у вищу когерентність, поступається місцем марківській темпоральності безпосередніх переходів без тривалої пам'яті» [1, р. 3147]. Ранні токени впливають на розподіл імовірностей через механізм само-уваги — але це статистичне обумовлення, а не зняття. Розширення контекстного вікна не змінює цього принципово — більше вікно залишається вікном, а не *Erinnerung*.

Нарешті, навіть якби поняття і суб'єкт були — залишається питання про те, що відбувається із суперечністю. Гегель стверджує: «Щось є живим лише тією мірою, якою воно містить у собі суперечність» і «Суперечність є більш

глибоким із двох визначень» [14, р. 439]. Утримати суперечність — означає не вирішити її на користь однієї зі сторін, а рухатися через неї до більш конкретного поняття. Трансформер, натомість, генерує статистично когерентне продовження при зустрічі з суперечністю — обирає одну зі сторін або згладжує її. Це є прямою реалізацією того, що Гегель називає «звичайним мисленням» — де визначення «вирішуються суперечністю лише в інші визначення або в ніщо» [14, р. 440]. Якщо ШІ може взяти на себе задачі формальної логіки — діалектична залишається йому недоступною.

Висновки. Здивування є не психологічним станом, а структурним моментом пізнання, що містить три необхідних умови: наявність стійкого поняття, здатного бути порушеним явищем; здатність утримувати попередні визначення як зняті моменти; здатність перетворити виявлену суперечність на імпульс руху до більш конкретного поняття. Ця реконструкція ґрунтується на синтезі кантівського розрізнення визначальної і рефлектуючої здатності судження [12] та гегелівської концепції суперечності як рушія мислення [13; 14].

Трансформерна архітектура LLM структурно виключає всі три умови здивування. Замість стійкого поняття — статистичний розподіл імовірностей. Замість здатності утримувати минулі визначення як зняті моменти — марківське вікно контексту, що щоразу починається з нуля. Замість зняття суперечності — її нейтралізація через вибір статистично когерентного продовження. Це виключення є структурним, а не технічним: воно не може бути усунуте розширенням контекстного вікна або збільшенням параметрів у межах існуючої парадигми.

Описане обмеження має конкретний практичний зміст. Визначальна здатність судження — підведення конкретного під вже відоме загальне — реалізується трансформером з ефективністю і масштабом, недосяжними для людини [5; 6]. Рефлектуюча здатність судження — формування нового

загального для даного часткового, що починається зі здивування перед суперечністю — залишається людською задачею не з технічних, а з архітектурно-логічних причин. Це розмежування, а не ілюзія взаємозамінності, є адекватною основою для розуміння того, чим є і чим не є сучасний штучний інтелект.

Звідси випливає практична рекомендація, яка стосується не стільки технічного вдосконалення моделей, скільки правильного розуміння їхньої природи. LLM працює виключно з мовним матеріалом — зі слідами чужого мислення, зафіксованими в текстах. Вона відтворює результати пізнавальних актів, що вже відбулися, але не відтворює сам процес, який ці результати породив — процес здивування, виявлення суперечності і руху до більш конкретного поняття. Тому межа між задачами, які доцільно делегувати LLM, і задачами, що вимагають людської присутності, проходить саме тут: там, де потрібне відтворення і систематизація вже здобутого знання — модель є потужним інструментом; там, де потрібне нове здивування перед невідомим — вона структурно відсутня.

Перспективи подальших досліджень у цьому напрямі включають: аналіз проблеми відношення мислення і мови стосовно того, що LLM працює виключно з мовним матеріалом, не маючи доступу до предметної діяльності, яка ці сліди породила; застосування гюнтерівської багатозначної логіки як операціоналізації гегелівської діалектики для опису самореферентних систем; порівняльний аналіз радянської і американської традицій у теорії III крізь призму діалектичної і формальної логіки.

Список використаних джерел

1. Woods J. Prompt, negate, repeat: a Hegelian meditation on AI. AI & Society. 2026. Vol. 41, iss. 4. P. 3143–3157. DOI: <https://doi.org/10.1007/s00146-025-02802-z>.



2. Chukhrov K. The Philosophical Disability of Reason: Evald Ilyenkov's Critique of Machinic Intelligence. *Radical Philosophy*. 2020. Vol. 207. P. 67–78.
3. Shevlin H. Three frameworks for AI mentality. *Frontiers in Psychology*. 2026. Vol. 17. Art. 1715835. DOI: <https://doi.org/10.3389/fpsyg.2026.1715835>.
4. Shanahan M. Talking About Large Language Models. *Communications of the ACM*. 2023. Vol. 67. P. 68–79. DOI: <https://doi.org/10.48550/arXiv.2212.03551>.
5. Millière R., Buckner C. A Philosophical Introduction to Language Models. Part I: Continuity With Classic Debates. *arXiv*. 2024. DOI: <https://doi.org/10.48550/arXiv.2401.03910>.
6. Seo J. Unexplainability of Artificial Intelligence Judgments in Kant's Perspective. *arXiv*. 2025. DOI: <https://doi.org/10.48550/arXiv.2407.18950>.
7. Abdali S., Goksen C., Solodko M. та ін. Self-reflecting Large Language Models: A Hegelian Dialectical Approach. *arXiv*. 2025. DOI: <https://doi.org/10.48550/arXiv.2501.14917>.
8. Hu Z. Dialectics for Artificial Intelligence. *arXiv*. 2025. DOI: <https://doi.org/10.48550/arXiv.2512.17373>.
9. El Maouch M., Jin Z. Artificial Intelligence Inheriting the Historical Crisis in Psychology: An Epistemological and Methodological Investigation of Challenges and Alternatives. *Frontiers in Psychology*. 2022. Vol. 13. Art. 781730. DOI: <https://doi.org/10.3389/fpsyg.2022.781730>.
10. Automating epistemology: how AI reconfigures truth, authority, and verification. *AI & Society*. 2025. DOI: <https://doi.org/10.1007/s00146-025-02560-y>.
11. Aristotle. *Metaphysics* / transl. W. D. Ross. Oxford : Clarendon Press, 1924. 982b.
12. Kant I. *Critique of the Power of Judgment* / transl. P. Guyer, E. Matthews. Cambridge : Cambridge University Press, 2000. 384 p.

13. Hegel G. W. F. Наука логіки. Енциклопедія філософських наук. Мала логіка / уклад. Б. Лис. Київ : Ліра-К, 2023. 373 с.
14. Hegel G. W. F. Science of Logic / transl. A. V. Miller. London : Allen & Unwin, 1969. 844 p.
15. Vaswani A., Shazeer N., Parmar N. та ін. Attention is All You Need. Advances in Neural Information Processing Systems. 2017. Vol. 30. P. 5998–6008.
16. Goodfellow I., Bengio Y., Courville A. Deep Learning. Cambridge : MIT Press, 2016. 800 p.