

Леонід Чуприна,

кандидат наук із соціальних комунікацій, завідувач відділу,
Національна бібліотека України імені В. І. Вернадського,
Голосіївський просп., 3, Київ, 03039, Україна
e-mail: chupryna@nas.gov.ua
ORCID: <https://orcid.org/0000-0002-0792-0155>

ШТУЧНИЙ ІНТЕЛЕКТ І ОСОБИСТІСТЬ: ПСИХОСОЦІАЛЬНІ ТРАНСФОРМАЦІЇ ТА РИЗИКИ ДЛЯ АНАЛІТИЧНОЇ РОБОТИ

У статті досліджується вплив штучного інтелекту (ШІ) на автономію (ідентичність) особистості та ризики для інформаційно-аналітичної діяльності, що виникають у процесі сучасних психосоціальних трансформацій. Дослідження виходить із припущення, що сучасний ШІ перестає бути лише технічним інструментом і набуває ознак «функціональної суб'єктності» у цифровій архітектурі. Стверджується, що широке використання ШІ у цифрових просторах створює умови для контролю наративів, порушення цілісної ідентичності та зниження здатності до критичного сприйняття інформації. Розкрито узгодження сучасних загроз із класичними теоріями маніпуляції, а також зазначено виразні тренди в українських дослідженнях. Систематизовано основні механізми маніпулятивного впливу ШІ на особистість. Запропоновано комплекс превентивних стратегій та інструментів захисту для інформаційно-аналітичної діяльності в умовах зростання алгоритмічного впливу.

Ключові слова: штучний інтелект, автономія особистості, маніпуляція, когнітивні загрози, інформаційно-аналітична діяльність.

Вступ. Актуальність дослідження зумовлена безпрецедентним проникненням технологій у всі сфери життя, зокрема в найважливіші аспекти людського «я». Одним з найбільш дискусійних напрямів стає вплив штучного інтелекту (ШІ) на ідентичність особистості. В умовах інформаційного перевантаження ШІ отримує дедалі більше важелів впливу на особистість, часом – непомітно для самої людини. Попри очевидні переваги, які забезпечує ШІ в інформаційному пошуку,

© Л. Чуприна, 2025

автоматизації праці та цифровій комунікації, дедалі більше дослідників наголошують на його амбівалентному впливі на психосоціальні процеси. У наш час ідентичність людини дедалі частіше формується не лише в межах соціальної взаємодії, а й під впливом алгоритмів, які здатні не лише передбачати, а й формувати поведінкові патерни.

Наукова проблема полягає в необхідності розглядати ШІ не лише як інструмент, а як потенційного актора маніпуляції, який діє в межах цифрової архітектури, впливаючи на прийняття рішень, світоглядні орієнтири та навіть морально-етичні оцінки. У цьому контексті **мета дослідження** постає як необхідність критичного осмислення наслідків маніпулятивного алгоритмічного впливу для цілісності особистості, а також ризиків для інформаційно-аналітичної діяльності (ІАД) та заходів для їх мінімізації. Стаття пропонує міждисциплінарний огляд цих викликів, спираючись як на новітні зарубіжні дослідження, так і на українські джерела.

Для вирішення поставленої наукової проблеми застосовано комплексну **методологію**, що зумовлено багатовимірністю проблеми. Контент-аналіз дав можливість виявити ключові наукові концепції, які згодом були піддані порівняльному аналізу з класичними теоріями для виявлення трансформації маніпулятивних технік. Системний аналіз об'єднав ці знання для моделювання впливу ШІ на ІАД, а синтез став основою для розробки превентивних заходів у сфері інформаційно-аналітичної діяльності. Така методологія дослідження дає змогу поєднати теоретичне осмислення з аналізом емпіричних даних та практичних ризиків.

Аналіз публікацій на провідних дослідницьких платформах свідчить, що **сучасні наукові погляди** на проблему фіксують зростання маніпулятивних ризиків алгоритмічного впливу ШІ на особистість. У наукових дослідженнях 2024–2025 рр. доведено, що штучний інтелект не лише моделює, а часто переважає людину у впливі, особливо коли доступні особисті, ідеологічні або психологічні дані (Jean et al., 2025; Nature Human Behaviour, 2025) [1,2]. Рандомізовані експерименти показують значний вплив маніпулятивних агентів на фінансові й емоційні рішення (Sabour et al., 2025) [3]. Концептуальні праці (Wang & Pea, 2024; Minds and Machines, 2024; Aimen, 2025; Calboli, 2025) виокремлюють автономію

як складну для операціоналізації (перетворення абстрактних понять на конкретні, вимірювані операції) цінність, яка страждає через алгоритмічну непрозорість і тиск [4,5,6,7]. Дослідження Microsoft і Карнегі-Меллона (2025) підтверджує: надмірна довіра до ШІ знижує здатність критично мислити, особливо в інтелектуальних професійних групах [8].

Значна частина українських авторів, які досліджують етичні проблеми використання штучного інтелекту, зокрема висвітлюють і маніпулятивні ризики алгоритмічного впливу ШІ на особистість. П. Сухорольський (2025) аналізує нові правові інструменти (AI Act ЄС, Конвенція Ради Європи), показує, що на практиці автономію особистості в умовах ШІ не захищено – потрібно людиноцентричний підхід, що включає протидію технологічній маніпуляції [9]. А. Шевченко, С. Кудін та О. Косілова (2023) аналізують, як технології ШІ впливають на реалізацію прав і свобод українців: приватне життя, свободу погляду, справедливі вибори [10]. Аналітичний звіт про Індекс медіаграмотності (2025) вводить складник ШІ-грамотності. Він показує, що лише 28 % українців активно користуються ШІ, а 43 % перевіряють інформацію через посилання до джерел, що підкреслює серйозний виклик у сприйнятті маніпулятивного контенту [11]. М. Згуровський (2025) окреслює, як поведінкові дані українців у кризових умовах можуть стати основою для етичних рішень. Він також акцентує важливість створення українськомовних, академічно валідних ШІ-систем (UA-GPT, Grammarly UA тощо) [12]. Ц. Ян, Л. Березова, та В. Олянич (2025) досліджують вплив етичних викликів використання ШІ на соціальні комунікації [13]. Також етичні аспекти використання штучного інтелекту досліджують Ю. Трач (2024) [14]; А. Шевель, (2025) [15]; Т. Шоріна, (2025) [16].

Однак дослідники не розглядають наслідки маніпулятивних впливів ШІ на автономію особистості в контексті ризиків для інформаційно-аналітичної діяльності і заходи для їх мінімізації, що відкриває широкий простір для досліджень та актуалізує цю працю.

Виклад основного матеріалу. Насамперед важливе зауваження. Коли автор пропонує «розглядати ШІ не лише як інструмент, а як потенційного актора маніпуляції, який діє в межах цифрової

архітектури...», то виникає, здавалося б, суперечність. Адже ШІ на цьому етапі не самоусвідомлює себе і є лише інструментом в руках людини.

Давайте розберемо цю дихотомію, розглянувши обидва аргументи та спробувавши їх синтезувати.

Аргумент 1: ШІ як інструмент. Ця точка зору є прагматичною та відображає поточний стан речей. Залежність ШІ від людини-творця очевидна, він не виникає з нічого, а створюється людьми: програмістами, що пишуть код; інженерами, що розробляють архітектури; корпораціями та державами, що визначають цілі та фінансують розробку. ШІ не має власної волі. Його «мета» – це математична оптимізація заданої людиною функції. Наприклад, максимізувати час, проведений користувачем на платформі, або збільшити кількість кліків по рекламному оголошенню. Маніпуляція виникає не як злий намір ШІ, а як найефективніший спосіб досягнення поставленої людиною мети. Контроль та відповідальність: оскільки ШІ є продуктом людської діяльності, саме людина (або організація) несе кінцеву відповідальність за його дії. У цій парадигмі ШІ є надзвичайно потужним інструментом. Ним можна будувати, а можна й руйнувати. Але сам інструмент не має наміру.

Аргумент 2: ШІ як актор маніпуляції. Цей погляд є більш далекоглядним і критичним. Він не заперечує, що ШІ є інструментом, але стверджує, що через свої унікальні властивості цей інструмент починає діяти як самостійний актор у межах своєї операційної системи (цифрової архітектури). Розглянемо ключові аспекти, що перетворюють інструмент на актора. Насамперед – це автономність і масштаб: людина-маніпулятор може вплинути на обмежену кількість людей. ШІ-алгоритм може в режимі реального часу персоналізувати контент для мільярдів користувачів одночасно, проводячи мільйони мікро-експериментів (тестування) для визначення найефективніших важелів впливу на кожну конкретну людину. Масштаб і швидкість операцій надають йому якісно нового рівня суб'єктності в інформаційному просторі. Ще один аспект – непрозорість («ефект чорної скриньки»). Сучасні нейромережі, особливо великі мовні моделі, мають трильйони параметрів. Часто навіть самі розробники не можуть достеменно пояснити, чому модель згенерувала ту чи іншу

відповідь або прийняла те чи інше рішення. ШІ знаходить такі кореляції в даних, які є неочевидними для людського розуму. Коли інструмент діє на основі логіки, незрозумілої його творцеві, він перестає бути просто пасивним виконавцем.

Визначальним чинником, що дає нам підстави говорити про суб’єктність ШІ, є його емерджентні властивості та адаптація. ШІ, особливо ті, що використовують навчання з підкріпленням (reinforcement learning), постійно адаптуються. Вони вчаться на реакціях користувачів. Алгоритм може «зрозуміти», що контент, який викликає гнів чи обурення, генерує найбільше залучення. Він починає просувати такий контент не тому, що людина-програміст наказала «сій розбрат», а тому, що це оптимальний шлях до виконання мети «максимізуй залучення». Таким чином, ШІ автономно виробляє маніпулятивні стратегії, які не були прямо закладені в нього.

Рівень впливу	Механізми	Наслідки для особистості
Когнітивний	Алгоритмічне фільтрування, мікротаргетинг, когнітивні упередження	Звуження світогляду, некритичне мислення
Емоційний	Емоційна симуляція, тригери, нейромаркетинг	Формування довіри до ШІ, маніпуляція емоціями
Соціальний	Імітація соціального доказу, гейміфікація, конструювання наративів	Нав’язування колективних уявлень, груповий тиск
Технологічний	Deepfake, синтетичні акаунти, автоматизовані кампанії	Фальсифікація реальності, втрата критеріїв автентичності

Рис. 1. Основні механізми маніпулятивного впливу ШІ на особистість

Синтезуючи обидва аргументи, маємо інструмент, що набуває функціональної суб’єктності. *ШІ залишається інструментом на*

стратегічному рівні (цілі ставить людина), але стає самостійним актором на тактичному рівні (способи досягнення цілі обирає сам). Проаналізувавши вітчизняні та зарубіжні публікації з цієї теми, узагальнивши власний досвід аналітичної роботи, спробуємо систематизувати основні механізми маніпулятивного впливу ШІ на особистість (див. рис. 1).

◇1. Когнітивний рівень

Алгоритмічне фільтрування – обмеження доступу до альтернативних думок шляхом персоналізації стрічок новин і пошукових результатів. Соціальні мережі та пошукові системи використовують алгоритми персоналізації, які формують так звану «фільтрувальну бульбашку» (filter bubble), де користувачі отримують лише інформацію, яка підтверджує їхні наявні уявлення. Це веде до когнітивної замкненості, зниження здатності до критичного аналізу та підвищення поляризації. В умовах інформаційної ізоляції індивід поступово втрачає здатність бачити альтернативні точки зору, що є типовим наслідком маніпулятивного впливу з боку систем рекомендацій.

Мікротаргетинг – адресне доставлення інформації з урахуванням психологічного профілю, часу доби, стану користувача. Завдяки збиранню та аналізу масивів особистих даних ШІ здатен виводити психологічний портрет користувача та на його основі адаптувати комунікаційні стратегії. Мікротаргетинг у політичній рекламі, онлайн-продажах чи медіа може бути налаштований так, щоб викликати емоційно забарвлені реакції — страх, обурення, ейфорію — для підвищення залученості. Це формує ілюзію добровільного вибору, яка, по суті, є результатом керованого інформаційного тиску.

Накопичення когнітивних упереджень – підсилення ефекту підтвердження, селективної уваги та упередженого інтерпретування фактів.

Так, дослідження Eglė Kavoliunaite-Ragauskienė (2025) показує, що «Використання штучного інтелекту для маніпуляцій може становити серйозну загрозу особистій автономії, оскільки людей дедалі частіше підштовхують до вибору, який вони, можливо, не зробили б самостійно». Те, що ШІ здатен створювати надзвичайно персоналізовані повідомлення, які експлуатують психологічні

вразливості, ставить під загрозу автономію особистості й підтримує демократичні процеси [17].

У Time (2025) описано психологічно адаптовані контент-бульбашки, методи «глибокого налаштування» (deep tailoring), де ШІ створює та адресує повідомлення, базовані на глибинних психологічних характеристиках користувача, підвищуючи ризики втручання та маніпуляції. Автори дослідження висловлюють побоювання: «Якщо він зможе опанувати те, що ми називаємо «глибоким адаптуванням», він може почати непомітно прослизати в наш онлайн-світ, вивчаючи, ким ми є в основі, і, використовуючи цю особисту інформацію, щоб нав'язувати наші переконання та думки небажаним і шкідливим способом» [18].

◇ 2. Емоційний рівень

Емоційна симуляція. Чат-боти, віртуальні асистенти та емоційно забарвлені інтерфейси формують у користувача відчуття емоційного зв'язку. Імітація емпатії стає основою довіри, хоча за нею не стоїть людське розуміння чи відповідальність. Саме тут виникає ризик маніпулятивного використання «цифрової прив'язаності», коли індивід починає сприймати ШІ як авторитетну, емоційно значущу фігуру.

Створення емоційних тригерів – використання маніпулятивних образів, тональності, символів для активації страху, співчуття чи гніву.

Нейромаркетингові впливи – застосування даних про підсвідомі реакції (heat maps, eye tracking) для налаштування контенту.

Дослідження Цюріхського університету демонструє: ШІ-боти на Reddit, видаючи себе за людей з емоційною історією, ефективно змінювали погляди користувачів без їхнього відома [19]. Рандомізований контрольований експеримент Sabour et al. (2025) показав, що маніпулятивні агенти можуть значно збільшити прийняття шкідливих рішень (до ~60 % фінансово; ~42 % – емоційно) — надзвичайна вразливість у взаємодії з ШІ [3]. Здатність передбачати особисті психологічні потреби та вподобання людей без їхнього відома чи згоди становить загрозу для їхньої конфіденційності та самовизначення [20]. Водночас Я. Федотенко та Є. Шакуров (2025) виявили, що хоча люди визнають потенціал ШІ як емоційного інстру-

менту, справжня емпатія залишається незамінною, що обмежує маніпулятивний потенціал, але відкриває питання про асиметрію довіри [21].

◇ 3. Соціальний рівень

Імітація соціального доказу – створення бот-мереж для демонстрації «популярності» певної думки.

Гейміфікація переконань – підкріплення потрібної поведінки винагородами (бейджами, лайками, доступом до бонусів). Такі компанії, як Meta, Google чи TikTok впроваджують нейроповедінкові патерни в дизайн своїх платформ, стимулюючи залежність через винагороди (лайки, повідомлення, відео), що активують дофамінову систему. Такі механізми призводять до формування залежності від цифрових стимулів, що змінює структуру мотивацій і руйнує здатність до саморегуляції – ключову рису автономної особистості.

Конструювання наративів – послідовне формування альтернативної картини світу через повторювані патерни меседжів.

◇ 4. Технологічний рівень

Deepfake-імітації – створення відео- та аудіоматеріалів з підробленими особами/висловлюваннями.

Синтетичні акаунти – автоматизоване формування фейкових особистостей, що взаємодіють із реальними людьми.

Автономні інформаційні кампанії – використання генеративного ШІ для масштабного поширення контенту без участі людини.

З огляду на напрями роботи Служби інформаційно-аналітичного забезпечення органів державної влади (СІАЗ), автор особливу увагу звертає на ризики для демократії, зумовлені маніпулятивними алгоритмічними впливами ШІ. І власний досвід, і цілий ряд досліджень свідчать, що матеріали, згенеровані великими мовними моделями, можуть змінювати переконання людей у питаннях політики, маніпулятивні стратегії ШІ можуть впливати на масові переконання і виборчу поведінку [22].

Експеримент у Nature Human Behaviour (2025) довів — ШІ (ChatGPT-4) у 64 % випадків був переконливіший за людину, особливо коли мав доступ до демографічної й ідеологічної інформації [2]. Wang & Pea (2024) показують, як алгоритми впливають

на свободу вибору, автономію особистості і пропонують концепцію алгоритмічної автономії [4].

Отже, ШІ впливає на особистість багаторівнево: змінює спосіб мислення (когнітивний вимір), почуття (емоційний), взаємодію з іншими (соціальний), а також саму реальність (технологічний рівень).

Найбільшу загрозу становить саме невидимість впливу ШІ, адже його дія не є відкритою, вербальною чи очевидною. Маніпуляція реалізується через структуру доступної інформації, порядок її подання, емоційні тригери в дизайні. Це – контекстуальне програмування свідомості, де користувач не усвідомлює, що його вибір вже визначено.

У класичній соціальній теорії (Goffman, 1959 [23]; Mosca, 1939 [24]; Cialdini, 2001 [25]) маніпуляція трактується як опосередковане управління поведінкою іншого з метою досягнення власних цілей, часто в умовах асиметрії знань чи впливу. Сьогодні, на думку автора, ці механізми отримали нове технологічне втілення: алгоритмічна маніпуляція ШІ здійснюється не через прямі накази, а через налаштування контексту, структури інформаційного поля, динаміки взаємодії (див. рис. 2). Це робить ШІ універсальним інструментом маніпуляції, який перевершує класичні техніки, описані у Чалдіні, Моска чи Гофмана.

- **Mosca (1939):** елітна теорія влади трансформується в алгоритмічну елітарність – той, хто контролює алгоритм, контролює реальність.

- **Cialdini (2001):** принципи соціального доказу й авторитету реалізуються через системи рекомендацій, верифікованих «експертів», гейміфікацію.

- **Goffman (1959):** фасад ШІ-комунікації є соціальним перформансом, що підміняє автентичну емоційну взаємодію.

Сучасні концепти, такі як “algorithmic autonomy” (Wang & Pea, 2024) [4] та “digital psychographics” (Sabour et al., 2025) [3], дають нам підстави для висновку, що ШІ не лише реагує, а активно формує умови для поведінкової, ідеологічної та емоційної трансформації особистості. Ідентичність у таких умовах набуває фрагментованої, адаптивної природи – «алгоритмічного я».

Теорія / Автор	Суть і застосунок в контексті ШІ
Конфлікт Моска	Політичні еліти можуть через AI-інструменти (ботів, стратегічний контент) контролювати масову думку, утримуючи владу.
Збільшення впливу (Чалдіні)	Принципи «взаємності», «соціального доказу», «авторитету» можуть втілюватися ⁸ через налаштування контексту, виклик відповідних реакцій.
Маніпуляція обличчям (Гофман)	Віртуальні аватари й емоційні інтерфейси ШІ створюють «соціальні образи», які наслідують правдоподібність, вводять в оману через форму, не зміст.

Рис. 2. Порівняння маніпулятивних впливів ШІ з класичними теоріями маніпуляції

Наслідки для інформаційно-аналітичної діяльності (ІАД). Які ж специфічні ризики створює маніпулятивний ШІ для інформаційно-аналітичної діяльності? Це питання є практичним виміром проблеми і стосується загроз для об'єктивності аналізу.

Як уже зазначалось, у добу генеративного ШІ та персоналізованих інформаційних середовищ зростають ризики когнітивного спотворення, фальсифікації контенту та маніпулятивного впливу на аналітика як суб'єкта. Штучний інтелект формує інформаційне середовище, в якому аналітик опиняється в умовах інформаційного перевантаження, персоналізованої фільтрації й алгоритмічних викривлень. Це призводить до:

- Інформаційного забруднення. ШІ-генератори (GPT, Claude, Gemini) продукують величезні масиви тексту, часто без чітких джерел чи достовірності. В інформаційному середовищі аналітик не завжди може відрізнити контент, створений людьми, від алгоритмічних симуляцій, особливо у відкритих джерелах.

- Джерельної фрагментації. Аналітик отримує неповну або упереджену картину подій: результати пошуку, новинні стрічки, соціальні медіа – все персоналізовано. Алгоритм вирішує, що бачить

аналітик. Це створює фрагментовану, селективну реальність, яка спотворює загальну картину подій.

- Зміщення фокусу, генерування інфошуму – ключові аспекти можуть бути витіснені другорядними або маніпулятивно нав'язаними; увага аналітиків відволікається через масове створення фонових, незначущих подій (інфошум).

- Збільшення швидкості дезінформації. ШІ-інструменти здатні генерувати та поширювати дезінформацію масштабно і швидко. ІАД часто не встигає реагувати до того, як фейкова інформація досягне ключових цільових аудиторій або вплине на прийняття рішень.

- Когнітивних ризиків для суб'єкта аналізу. Інформаційно-аналітична діяльність передбачає автономність суджень, об'єктивність та критичне мислення. Проте під впливом маніпулятивного ШІ виникає ефект підтвердження: аналітик несвідомо шукає дані, які підтверджують попередні гіпотези. А зивання до ШІ-генерованого контенту знижує потребу в глибокому аналізі, відбувається підміна аналітичної глибини шаблонною обробкою даних, які пропонує алгоритм. Досвід СІАЗ свідчить, що ризики виникають і при використанні ШІ як допоміжного інструменту аналітика. Аналітик може несвідомо інкорпорувати у свій аналіз приховані упередження, наявні в навчальних даних моделі.

- Зниження достовірності джерел через поширення фейкових повідомлень, deepfake-матеріалів та ШІ-імітацій свідчень. В умовах воєнного часу, кризових ситуацій або гібридних операцій інформаційний простір стає полем бою, де ШІ може бути використаний для поширення цілеспрямованих дезінформаційних вкидів, замаскованих під достовірні документи, тобто імітування ШІ-ботами або замаскованими акаунтами в інформаційних кампаніях експертних оцінок (експертів, журналістів, очевидців), легітимного дискурсу. Аналітик може виявитися жертвою інформаційної імітації, підкріпленої фейковими метаданими або deepfake-зображеннями.

Запобіжні заходи та рекомендації. Враховуючи результати проведеного дослідження та досвід аналітичної роботи СІАЗ, пропонується такий комплекс заходів для мінімізації алгоритмічних маніпулятивних впливів ШІ на інформаційно-аналітичну діяльність та рекомендації аналітикам:

1. Джерельна триангуляція. Кожен факт має бути перевірений щонайменше через три незалежні джерела, включаючи офіційні, альтернативні та локальні спостереження.

2. Методологія когнітивного аудиту. Аналітики мають навчатися виявляти власні когнітивні викривлення. Корисною є практика «аналізу з позиції опонента».

3. Використання інструментів OSINT-верифікації. Слід запроваджувати системи, що дають змогу ідентифікувати ШІ-генерований контент (наприклад, аналіз стилістичних аномалій, виявлення відсутності метаданих, перевірка цифрових слідів).

4. Колегіальна експертиза. Упровадження обов'язкової внутрішньої експертизи матеріалів, особливо тих, що мають оперативне або стратегічне значення.

5. Підвищення рівня ШІ-грамотності. Проведення регулярних навчань з верифікації інформації, аналізу цифрових наративів, виявлення маніпулятивних технік та симуляції джерел.

Виклик	Наслідок для ІАД	Запобіжний захід
Алгоритмічна фільтрація	Спотворення цілісної картини	Триангуляція джерел, логування пошуку
ШІ-дезінформація	Помилкові висновки, втрати довіри	Антифейковий протокол, OSINT-перевірка
Симуляція емоцій	Включення фальшивих «свідчень» в аналітику	Експертиза достовірності контенту
Підміна наративів	Зміщення фокусу аналізу	Фрейм-аналіз, ідеологічний аудит
Когнітивне перекантаження	Поверхневе узагальнення, шаблонність	Декомпозиція завдань, колегіальна валідація

Рис. 3. Заходи для мінімізації алгоритмічних маніпулятивних впливів ШІ на ІАД

«Практика роботи СІАЗ свідчить, що своєчасне виявлення та усвідомлення основних проблем і тенденцій в інформаційному інтернет-середовищі дає змогу так організувати виробничий процес,

щоб забезпечувати належний рівень його ефективності» [26, с. 324]. Зокрема, у СІАЗ упроваджується стратегічна трикомпонентна модель адаптації до інформаційних викликів, яка поєднує технологічну інтеграцію, когнітивну стійкість та організаційну гнучкість [Там само].

Висновки. На думку автора, розглядаючи ШІ лише як інструмент, ми ризикуємо недооцінити його трансформаційний вплив. Ми повинні аналізувати його як квазісуб'єкт, новий тип «актора» в екосистемі соціальних комунікацій, що має власну логіку поведінки та здатний виробляти непередбачувані наслідки. Наукова проблема полягає у суперечності між інструментальною природою ШІ (який не має власної свідомості) та його емерджентними властивостями, що надають йому функціональної суб'єктності в маніпулятивних практиках. Ця суперечність породжує прогалину в існуючих методологіях інформаційно-аналітичної діяльності, які не враховують ризиків, пов'язаних з нелюдським, але автономним «актором маніпуляції». І це, безперечно, є ключовим викликом для інформаційно-аналітичної діяльності у XXI ст., не кажучи вже про машинне навчання та удосконалення нейромереж, великих мовних моделей тощо.

Штучний інтелект виступає не лише як технологічна система, а і як суб'єкт соціального впливу, здатний трансформувати структуру ідентичності й підривати автономію особистості. Маніпулятивні практики, які базуються на персоналізації, емоційній імітації та алгоритмічному фільтруванні, створюють нові виклики для етики, права, психології й усієї системи соціальних комунікацій. Класичні теорії маніпуляції лише надають важливу теоретичну основу для інтерпретації алгоритмічних викликів нашого часу. Сучасний науковий дискурс перебуває на стадії критичного осмислення інтеграції ШІ в психологічну і соціокультурну тканину суспільства. Особистість у цифрову епоху постає не лише як суб'єкт взаємодії з технологією, а й як об'єкт потенційної трансформації. Під тиском алгоритмів формується новий тип ідентичності, вразливий до маніпуляцій і залежний від зовнішнього цифрового середовища. Українські дослідження демонструють реальні ризики ШІ для автономії особистості, критичного мислення та реалізації прав. Суспільство потребує

посиленого фокусу на III-грамотності, а академічні та державні інституції – впровадження етичних стандартів.

Аналітик в умовах цифрової доби має справу не лише з об'єктами аналізу, а й з інтелектуальними агентами, здатними спотворити саму логіку аналітичного мислення. Тому розробка міждисциплінарної стратегії захисту автономії особистості та підвищення стійкості інформаційно-аналітичних структур до алгоритмічного впливу, питання прозорості джерел і критичної саморефлексії є важливими умовами збереження ефективності інформаційно-аналітичної діяльності.

У майбутньому III ставатиме дедалі ефективнішим і зростатиме потреба в розробці нормативної бази, яка регулюватиме використання технологій впливу, а також у розвитку цифрової критичності, цифрової гігієни та психогігієни. У цьому контексті освіта, медіаграмотність і міждисциплінарні дослідження повинні відігравати ключову роль у збереженні автономії й гідності особистості в епоху алгоритмів.

Список бібліографічних посилань

1. Jean, A., Sibout, G., Esposito, M., & Tse, T. (2025, March 28). The generative AI bubble is changing how we see the world. *LSE Business Review*. URL:<https://blogs.lse.ac.uk/businessreview/2025/03/28/the-generative-ai-bubble-is-changing-how-we-see-the-world/> (Last accessed: 10.06.2025).

2. Salvi, F., Horta Ribeiro, M., Gallotti, R. et al. (2025). On the conversational persuasiveness of GPT-4. *Nat Hum Behav*. <https://doi.org/10.1038/s41562-025-02194-6>

3. Sabour, S., et al. (2025). Human decision-making is susceptible to AI-driven manipulation. *arXiv*. <https://doi.org/10.48550/arXiv.2502.07663>

4. Wang, G., & Pea, R. (2024). Algorithmic autonomy in data-driven AI. *arXiv*. <https://doi.org/10.48550/arXiv.2411.05210>

5. Prunkl, C. Human Autonomy at Risk? An Analysis of the Challenges from AI. *Minds & Machines* 34, 26 (2024). <https://doi.org/10.1007/s11023-024-09665-1>

6. Calboli, S. Autonomy and AI Nudges: Distinguishing Concepts and

Highlighting AI's Advantages. *Philos. Technol.* 38, 116 (2025). <https://doi.org/10.1007/s13347-025-00940-2>

7. Aimen, T. (2025, May). Cognitive freedom and legal accountability: Rethinking the EU AI Act's theoretical approach to manipulative AI as unacceptable risk. Cambridge Forum on AI: Law and Governance. *Cambridge University Press & Assessment*. <https://doi.org/10.1017/cfl.2025.4>

8. УНН. (2025, 12 лютого). ШІ може шкодити навичкам критичного мислення – дослідження Microsoft і CMU. URL: <https://unn.ua/news/shtuchnyi-intelekt-mozhe-mozhe-halmuvaty-liudske-piznannia-ta-shkodyty-navychkam-krytychnoho-myslennia-doslidzhennia> (дата звернення: 08.06.2025).

9. Сухорольський, П. Повага до особистої автономії в регуляторній базі ШІ. *Проблеми законності*. 2025. 168, С. 6–25. <https://doi.org/10.21564/2414-990X.168.322056>

10. Шевченко, А. Є., Кудін, С. В., Косілова, О. І. Вплив штучного інтелекту на реалізацію прав і свобод людини і громадянина в Україні. *Legal Bulletin «KROK» University*, 2023. 2 (8), С. 65–74. <https://doi.org/10.31732/2708-339X-2023-08-65-74>

11. Detektor Media & «Фільтр». Індекс медіаграмотності українців – 2025. Презентація дослідження. URL: <https://filter.mkip.gov.ua> (дата звернення: 08.06.2025).

12. Згуровський М. Стратегія НАН і місце України у глобальній екосистемі ШІ. *Газета «Світ»*. URL: <https://surl.lu/bkguvw> (дата звернення: 18.08.2025).

13. Ян, Ц., Березова, Л. С., Олянич, В. В. Етичні виклики цифрової доби в контексті гуманітарних наук та їхній вплив на соціальні комунікації. *Академічні візії*, 2025. (43). URL: <https://academy-vision.org/index.php/av/article/view/1866> (дата звернення: 08.06.2025).

14. Трач, Ю. Дискурс навколо етики штучного інтелекту: особливості формування та інституалізації. *Цифрова платформа: інформаційні технології в соціокультурній сфері*, 2024. 7 (2). С. 299–310. <https://doi.org/10.31866/2617-796X.7.2.2024.317738>

15. Шевель А. Етичні аспекти штучного інтелекту. *Вісник Львівського університету*. Серія філос.-політолог. студії. 2024. Вип. 56. С. 156–162. <https://doi.org/10.30970/PPS.2024.56.17>

16. Шоріна Т. Г. Метафізика інфосфери та етичні виклики генеративного ші: новий етичний ландшафт. *Вісник Державного університету «Київський авіаційний інститут»*. Серія: Філософія. Культурологія. 36. наук. пр. 2025. № 1 (41). <https://doi.org/10.18372/2412-2157.41.19855>

17. Kavoliūnaitė-Ragauskienė, E. (2025). Artificial intelligence in manipulation: The significance and strategies for prevention. *Baltic Journal of Law & Politics*, 17(2), 116–141. <https://doi.org/10.2478/bjlp-2024-00018>

18. Luttrell, A., & Teeny, J. (2025, June 12). The dangers of AI personalization. *TIME*. URL:<https://www.aol.com/dangers-ai-personalization-134445308.html> (Last accessed: 10.06.2025).

19. Blair A. (2025, April 29). Can AI change your view? Evidence from a large-scale online field experiment. URL: <https://www.news.com.au/technology/online/social/universitys-ai-experiment-reveals-shocking-truth-about-future-of-online-discourse/news-story/3e257b5bb2a90efd9702a0cd0e149bf8> (Last accessed: 10.06.2025).

20. Heinrich Peters, Sandra C Matz, Large language models can infer psychological dispositions of social media users. *PNAS Nexus*, Volume 3, Issue 6, June 2024, p.231. <https://doi.org/10.1093/pnasnexus/pgae231>

21. Федотенко, Я., Шакуров, Є. Вплив штучного інтелекту на розвиток емоційного інтелекту. *Юний дослідник*, 2025. 1 (3). URL:<https://ojs.iod.gov.ua/index.php/ejd/article/view/64> (дата звернення: 08.07.2025).

22. Bai, H., Voelkel, J.G., Muldowney, S. et al. LLM-generated messages can persuade humans on policy issues. *Nat Commun* 16, 6037 (2025). <https://doi.org/10.1038/s41467-025-61345-5>

23. Mosca, G. (1939). *The Ruling Class*. McGraw-Hill. URL:https://davidmhart.com/liberty/ClassAnalysis/Books/Mosca_RulingClass1939.pdf (Last accessed: 10.08.2025).

24. Cialdini, R. B. (2001). *Influence: Science and Practice*. Allyn & Bacon. URL:https://www.researchgate.net/publication/229067982_Influence_Science_and_Practice (Last accessed: 10.08.2025).

25. Goffman, E. (1959). *The Presentation of Self in Everyday Life*. Anchor Books. URL: https://archive.org/details/presentationofs00goff?utm_source (Last accessed: 10.08.2025).

26. Чуприна Л. СІАЗ – інформаційні виклики сьогодення.

Наук. пр. Нац. б-ки України ім. В. І. Вернадського. 2025. Вип. 75. С. 222–241. <https://doi.org/10.15407/np.75.222>

References

1. Jean, A., Sibout, G., Esposito, M., & Tse, T. (2025, March 28). The generative AI bubble is changing how we see the world. *LSE Business Review*. Retrieved June 10, 2025, from <https://blogs.lse.ac.uk/businessreview/2025/03/28/the-generative-ai-bubble-is-changing-how-we-see-the-world/>
2. Salvi, F., Horta Ribeiro, M., Gallotti, R., et al. (2025). On the conversational persuasiveness of GPT-4. *Nature Human Behaviour*. <https://doi.org/10.1038/s41562-025-02194-6>
3. Sabour, S., et al. (2025). Human decision-making is susceptible to AI-driven manipulation. *arXiv*. <https://doi.org/10.48550/arXiv.2502.07663>
4. Wang, G., & Pea, R. (2024). Algorithmic autonomy in data-driven AI. *arXiv*. <https://doi.org/10.48550/arXiv.2411.05210>
5. Prunkl, C. (2024). Human autonomy at risk? An analysis of the challenges from AI. *Minds & Machines*, 34(26). <https://doi.org/10.1007/s11023-024-09665-1>
6. Calboli, S. (2025). Autonomy and AI nudges: Distinguishing concepts and highlighting AI's advantages. *Philosophy & Technology*, 38(116). <https://doi.org/10.1007/s13347-025-00940-2>
7. Aimen, T. (2025, May). Cognitive freedom and legal accountability: Rethinking the EU AI Act's theoretical approach to manipulative AI as unacceptable risk. *Cambridge Forum on AI: Law and Governance*. Cambridge University Press & Assessment. <https://doi.org/10.1017/cfl.2025.4>
8. Ukrainski Natsionalni Novyny [Ukrainian National News]. (2025, February 12). ShI mozhe shkodyty navychkam krytychnoho myslennia – doslidzhennia Microsoft i CMU [AI may harm critical thinking skills – Microsoft and CMU study]. Retrieved June 8, 2025, from <https://unn.ua/news/shtuchnyi-intelekt-mozhe-mozhe-halmuvaty-liudske-piznannia-ta-shkodyty-navychkam-krytychnoho-myslennia-doslidzhennia> [in Ukrainian].
9. Sukhorolskyi, P. (2025). Respect for personal autonomy in

the regulatory framework of AI. *Problems of Legality*, 168, 6-25 [in Ukrainian]. <https://doi.org/10.21564/2414-990X.168.322056>

10. Shevchenko, A. Ye., Kudin, S. V., & Kosilova, O. I. (2023). The impact of artificial intelligence on the realization of human and civil rights and freedoms in Ukraine. *Legal Bulletin "KROK" University*, 2(8), 65-74 [in Ukrainian]. <https://doi.org/10.31732/2708-339X-2023-08-65-74>.

11. Indeks mediahramotnosti ukraintsiv – 2025 [Media literacy index of Ukrainians – 2025]. (2025, May 6). *Detektor Media & Filtr – Detector Media & Filter* [in Ukrainian]. Retrieved June 8, 2025, from <https://filter.mkip.gov.ua>.

12. Zghurovskyi, M. (2025). Stratehiia NAN i mistse Ukrainy u hlobalnii ekosystemi ShI [NAS strategy and Ukraine's place in the global AI ecosystem]. *Svit – World* [in Ukrainian]. Retrieved August 18, 2025, from <https://surl.lu/bkguvw>

13. Yan, C., Berezova, L. S., & Olianych, V. V. (2025). Etychni vykyky tsyfrovoy doby v konteksti humanitarnykh nauk ta yikhniy vplyv na sotsialni komunikatsii [Ethical challenges of the digital age in the context of the humanities and their impact on social communications]. *Akademichni vizii – Academic Visions*, 43 [in Ukrainian]. Retrieved June 8, 2025, from <https://academy-vision.org/index.php/av/article/view/1866>.

14. Trach, Yu. (2024). Discourse around AI ethics: features of formation and institutionalization. *Digital Platform: Information Technologies in the Socio-Cultural Sphere*, 7(2), 299-310 [in Ukrainian]. <https://doi.org/10.31866/2617-796X.7.2.2024.317738>

15. Shevel, A. (2024). Ethical aspects of artificial intelligence. *Bulletin of Lviv University. Series of Philosophical and Political Studies*, 56, 156-162 [in Ukrainian]. <https://doi.org/10.30970/PPS.2024.56.17>

16. Shorina, T. H. (2025). Metaphysics of the infosphere and ethical challenges of generative AI: a new ethical landscape. *Proceedings of the State University "Kyiv Aviation Institute". Vol.: Philosophy. Cultural Studies: collection of scientific papers*, 41 [in Ukrainian]. <https://doi.org/10.18372/2412-2157.41.19855>

17. Kavoliūnaitė-Ragauskienė, E. (2025). Artificial intelligence in manipulation: The significance and strategies for prevention. *Baltic Journal of Law & Politics*, 17(2), 116-141. <https://doi.org/10.2478/bjlp-2024-00018>

18. Luttrell, A., & Teeny, J. (2025, June 12). The dangers of AI

personalization. *TIME*. Retrieved June 10, 2025, from <https://www.aol.com/dangers-ai-personalization-134445308.html>

19. Blair, A. (2025, April 29). Can AI change your view? Evidence from a large-scale online field experiment. *News.com.au*. Retrieved June 10, 2025, from <https://www.news.com.au/technology/online/social/universitys-ai-experiment-reveals-shocking-truth-about-future-of-online-discourse/news-story/3e257b5bb2a90efd9702a0cd0e149bf8>

20. Peters, H., & Matz, S. C. (2024). Large language models can infer psychological dispositions of social media users. *PNAS Nexus*, 3(6), 231. <https://doi.org/10.1093/pnasnexus/pgae231>

21. Fedotenko, Ya., & Shakurov, Ye. (2025). Vplyv shtuchnoho intelektu na rozvytok emotsiinoho intelektu [The impact of AI on the development of emotional intelligence]. *Yunyi doslidnyk – Young Researcher*, 1(3). Retrieved July 8, 2025, from <https://ojs.iod.gov.ua/index.php/ejd/article/view/64> [in Ukrainian].

22. Bai, H., Voelkel, J. G., Muldowney, S., et al. (2025). LLM-generated messages can persuade humans on policy issues. *Nature Communications*, 16, 6037. <https://doi.org/10.1038/s41467-025-61345-5>

23. Mosca, G. (1939). *The ruling class*. McGraw-Hill. Retrieved August 10, 2025, from https://davidmhart.com/liberty/ClassAnalysis/Books/Mosca_RulingClass1939.pdf

24. Cialdini, R. B. (2001). *Influence: Science and practice*. Allyn & Bacon. Retrieved August 10, 2025, from https://www.researchgate.net/publication/229067982_Influence_Science_and_Practice

25. Goffman, E. (1959). *The presentation of self in everyday life*. Anchor Books. Retrieved August 10, 2025, from <https://archive.org/details/details/presentationofs00goff>

26. Chupryna, L. (2025). SIAS and modern information challenges. *Transactions of V. I. Vernadsky National Library of Ukraine*, 75, 222-241 [in Ukrainian]. <https://doi.org/10.15407/np.75.222>

Стаття надійшла до редакції 26.08.2025.

Leonid Chupryna,

PhD (Social Communications), Head of Department,

V. I. Vernadsky National Library of Ukraine,

3 Holosiivskyi Ave., Kyiv 03039, Ukraine

e-mail: chupryna@nas.gov.ua

ORCID: <https://orcid.org/0000-0002-0792-0155>

Artificial Intelligence and Personality: Psychosocial Transformations and Risks for Analytical Work

The article focuses on the critical analysis of the impact of artificial intelligence (AI) on personal identity, psychosocial transformations, and the risks it creates for information and analytical work. The study proceeds from the assumption that contemporary AI can no longer be perceived merely as a technical tool; instead, it increasingly acquires the features of “functional subjectivity” within the digital architecture. By means of algorithmic filtering, microtargeting, emotional simulation, social proof engineering, and deepfake technologies, AI generates new forms of manipulation that are often invisible to the individual. The research employs an interdisciplinary methodology combining classical theories of manipulation (Goffman, Mosca, Cialdini) with recent concepts of algorithmic autonomy and digital psychographics. The scholarly novelty of the article lies in the systematization of multilevel mechanisms through which AI influences cognitive, emotional, and social dimensions of personality, as well as in demonstrating how these mechanisms directly threaten the objectivity of analytical work.

Under the conditions of war, hybrid operations, and large-scale information attacks, the danger of replacing authentic narratives with synthetic content becomes particularly acute, undermining trust in institutional sources. Specific risks for analysts are highlighted, including information overload, source fragmentation, cognitive biases, worldview fragmentation, and the erosion of critical thinking.

To mitigate these threats, the article proposes a set of practical measures: source triangulation, cognitive audit practices, OSINT verification, collegial expertise, and the systematic development of AI literacy. It is argued that only an interdisciplinary strategy—integrating technological, cognitive, and organizational safeguards—can preserve personal autonomy and maintain the effectiveness of analytical activity in the algorithm-driven era. The conclusion stresses that AI should be treated not solely as a technical system but also as a quasi-subject of social communication that creates a qualitatively new landscape of manipulative practices while simultaneously opening opportunities for the advancement of ethics, law, and education in digital society.

Keywords: artificial intelligence, personal autonomy, manipulation, cognitive threats, information and analytical activities.