

УДК 81'42:004.89:070
DOI 10.31652/2521-1307-2025-41-10

Лінгвістичний аналіз висвітлення штучного інтелекту в новинах

Марія Чадюк <https://orcid.org/0000-0002-6727-1126>

Національний університет «Києво-Могилянська академія», м. Київ, Україна

Надійшла до редакції: 17.09.2025 • Схвалено до друку: 25.11.2025

Анотація

У статті наголошено на важливості дослідження висвітлення штучного інтелекту (ШІ) в новинах, оскільки репрезентація в медіа впливає на його сприйняття населенням і регуляцію на законодавчому рівні; окреслено основні тенденції й наративи репрезентації ШІ в закордонних медіа; з'ясовано, що ця тема недостатньо розроблена в лінгвістиці. Щоб визначити, як часто та як саме в заголовках українських медіа (ТСН, Українська правда, УНІАН) у серпні 2024 – липні 2025 легітимують і делегітимують ШІ, застосовано критичний дискурс-аналіз Т. ван Левена. Аналіз дав змогу встановити, що 92 заголовки містять дискурсивні стратегії легітимації, 122 заголовки – стратегії делегітимації ШІ. З'ясовано, які дискурсивні тактики найчастіше використовують для легітимації, а якими послуговуються для делегітимації ШІ в новинах. Також визначено тенденції медіа репрезентувати ШІ як суб'єкта дії, особливо щодо вчинків, що мають негативні наслідки, та персоніфікувати ШІ, застосовуючи для цього лексику, яку зазвичай використовують щодо істот.

Ключові слова: медіа, новини, заголовки, штучний інтелект, критичний дискурс-аналіз, дискурсивна легітимація

Linguistic Analysis of the News Coverage of Artificial Intelligence

Mariia Chadiuk <https://orcid.org/0000-0002-6727-1126>

National University of Kyiv-Mohyla Academy, Kyiv, Ukraine

Received: 17.09.2025 • Accepted: 25.11.2025

Abstract

The coverage of artificial intelligence (AI) in the news plays a significant role in the public perception of AI and the regulation of this technology. This portrayal is especially critical nowadays, since AI is increasingly used in numerous areas. It comes as no surprise that studies on media coverage of AI have recently become widespread. In particular, scientists use content analysis, frame analysis, text tone analysis, and other methods that allow them to research a large volume of texts and determine whether AI is represented positively or negatively. At the same time, few studies examine the representation of AI from a linguistic perspective, which explains the relevance of this article's topic. The purpose of this study is to analyze the tactics of discursive (de)legitimization of AI used in news headlines. The approach of critical discourse analysis, namely the methodology of T. van Leeuwen, was applied because representation in the media can affect the legitimation or delegitimation of AI usage. The research data were headlines of TSN, Ukrainska Pravda, and UNIAN articles displayed on the «artificial intelligence» search query and published from August 2024 to July 2025. The objectives of the study were to outline the main trends and narratives of the representation of artificial intelligence in the foreign press, determine tactics of discursive (de)legitimization of AI used in Ukrainian news headlines, and examine which tactics and for which aim are used the most. Examination of scientific articles on the portrayal of AI in the media proved the idea of negative representation of AI wrong. According to research, the prevailing attitude is optimistic or neutral. However, as a few studies have demonstrated, the media mentions the risks of AI more and more often. The analysis of Ukrainian news has revealed that 92 headlines have strategies that legitimise AI, and 122 headlines contain delegitimation strategies. Usually, the use of AI is legitimized by indicating the positive consequences of its activities (tactics of result and means orientation), its use by influential countries or institutions (a tactic of positional authority), comparing AI with human intelligence (a tactic of analogy) and using adjectives, nouns and verbs with positive connotations to denote the actions of AI (tactics of evaluation and abstraction). The delegitimation of AI is implemented in the media headlines by reporting the negative consequences of its activities or the illegitimate purposes for which it is used (tactics of result, effect, and means orientation), as well as the use of nouns, verbs, adjectives, and adverbs with negative connotations to describe AI's actions (tactics of abstraction and evaluation). It was also noted that AI is often represented in the media as an agent, which can belittle the role of those who use AI (in the case of reporting achievements) and shift responsibility for negative consequences from people to AI, portraying certain outcomes as inevitable or uncontrollable, making the headlines sensational.

Keywords: media, news, headlines, artificial intelligence, critical discourse analysis, discursive legitimation

Постановка проблеми. Нині штучний інтелект (ШІ) усе частіше застосовують у бізнесі (Chuan, Tsai, Cho, 2019), медіа (Newman, Ross Arguedas, Robertson, Nielsen, Fletcher, 2025), медицині (Bunz, Braghieri, 2022) та інших сферах. З появою ChatGPT та безкоштовних ШІ-помічників кожен користувач мережі «Інтернет» може не лише спостерігати за тим, як працює ШІ, а й використати його для виконання власних завдань. Незважаючи на ці численні можливості, ставлення до ШІ залишається контраверсійним: наукові студії засвідчують упередженість алгоритмів (Fleisig et al., 2024), на розробників ШІ подають у суд (наприклад, позов New York Times проти OpenAI щодо використання їхніх статей для тренування моделей штучного інтелекту), чимало сценаріїв реалізації цієї технології викликають занепокоєння (Cave, Coughlan, Dihal, 2019; Sartori, Bocca, 2023) тощо.

Висвітлення штучного інтелекту в медіа, особливо новинах (Chuan, Tsai, Cho, 2019; Nguyen, 2023), відіграє велику роль у сприйнятті ШІ населенням (González-Arias, López-García, 2024; Nguyen, 2023; Ouchchy, Coin, Dubljević, 2020; Sartori, Bocca, 2023) і його регуляції на політичному рівні (Lammar, Horst, Müller, 2025). Науковці зауважують, що коли медіа будуть зосереджуватися на сенсаційних, але маловірогідних ризиках, то можуть оминати реальні або більш вірогідні виклики – фактичні помилки, упередженість, використання ШІ для дезінформації та ін. (Brennen, Howard, Nielsen, 2018; Cave et al., 2018; Ouchchy, Coin, Dubljević, 2020; Sartori, Bocca, 2023). З іншого боку, лише оптимістичне висвітлення штучного інтелекту може зменшити його регулювання, спричинивши, наприклад, незахищеність певних даних, або сформувані надзвичайно високі очікування, які розробники не зможуть задовольнити, що може знизити рівень довіри до можливостей ШІ (Cave et al., 2018; Cave, Coughlan, Dihal, 2019).

Дослідження висвітлення ШІ в медіа набувають нині все більшого поширення (Lammar, Horst, Müller, 2025). Найчастіше іноземні науковці зосереджуються на вибірковості висвітлення та здатності медіа фреймувати події (Chuan, Tsai, Cho, 2019; Nguyen, 2023; Sartori, Bocca, 2023) і застосовують контент-аналіз, фреймовий аналіз, аналіз тональності тексту та інші методи, що дають змогу проаналізувати великий обсяг текстів (Bunz, Braghieri, 2022; Ouchchy, Coin, Dubljević, 2020). Українські наукові праці фокусуються здебільшого на темі впровадження ШІ в медіасферу та медіаосвіту (Черемних, 2024), з'ясуванні ознак дезінформації, створеної ШІ (Кузнецова, 2024), мовних маркерів, що вказують на текст, написаний ШІ (Fedoriv, Pirozhenko, Shuhai, 2023), застосуванні ШІ для виявлення дезінформації (Легомінова, Тищенко, Недодай, Дьячук, Капелюшна, 2024) та ін. Студій, які вивчають репрезентацію ШІ з погляду лінгвістики, мало (Lammar, Horst, Müller, 2025), що й зумовлює актуальність статті.

Аналіз останніх досліджень і публікацій.

Репрезентація ШІ в текстах журналістів стала предметом вивчення порівняно недавно, що може бути зумовлено збільшенням уваги медіа до ШІ після 2016 року (Chuan, Tsai, Cho, 2019). Спершу доволі поширеним було переконання, що медіа висвітлюють ШІ негативно й сенсаційно, формуючи обережне ставлення до цієї технології (Brennen, Howard, Nielsen, 2018; Weber, 2025). Проте, як засвідчили дослідження, насправді висвітлення ШІ в англомовних медіа є оптимістичним (Fast, Horvitz, 2017): за деякими підрахунками, публікацій, які висвітлюють ШІ позитивно, вдвічі-втричі більше від статей, що наголошують на ризиках (Fast, Horvitz, 2017). Аналіз дописів Reddit за 2010–2015 рр. підтверджує оптимістичну репрезентацію (Fast, Horvitz, 2017). З іншого боку, науковці, які розподіляють публікації на

три категорії (ті, що висвітлюють лише переваги ШІ; ті, що згадують лише ризики; нейтральні, що наводять переваги й ризики), з'ясували, що переважають нейтральні статті (Ouchchy, Coin, Dubljević, 2020). Деякі науковці помічають кореляцію між темою та способом її висвітлення, наприклад, у публікаціях про бізнес та економіку наводять здебільшого переваги ШІ, а в статтях, присвячених етиці, наводять лише ризики (Chuan, Tsai, Cho, 2019).

Водночас дослідники зазначають, що частка статей, які висвітлюють негативні аспекти ШІ, поступово зростає (Nguyen, 2023), а ризики, хоч і обговорюють рідше, проте більш детально (Chuan, Tsai, Cho, 2019). Найбільш згадуваними ризиками в американських і британських медіа є порушення приватності, упередження алгоритмів, кібербезпека й дезінформація (Nguyen, 2023), а для медіа романськими мовами – маніпуляція суспільною думкою, непевність, порушення авторських прав, брак робочих місць (González-Arias, López-García, 2024).

Науковці також зауважували відмінності у висвітленні ШІ залежно від: а) позиції медіа – праві та ліві медіа Великобританії зосереджуються на різних аспектах (Brennen, Howard, Nielsen, 2018); б) політичної та економічної ситуації в країні (González-Arias, López-García, 2024).

Тему висвітлення ШІ в медіа досліджують і з погляду використаних наративів. Однією з найвідоміших є класифікація чотирьох «надій» і відповідних «страхів» (Cave, Dihal, 2019): а) довголіття – нелюдськість: ШІ може допомогти нам жити довше, але це може змінити нашу ідентичність; б) легкість – застарілість: ШІ може полегшити наше життя, проте чи не замінить він людей; в) винагорода – відсторонення: ШІ може задовольнити наші бажання, але це може віддалити нас одне від одного; г) панування – повстання: ШІ може надати військові переваги тим, хто його застосовує, але ці технології можуть повстати проти людей (Cave, Dihal, 2019). Опитування

показали, що «надіями» ці наративи можна назвати лише умовно, бо здебільшого лише легкість і довголіття викликали в аудиторії позитивні емоції; решта наративів викликали занепокоєння (Cave, Coughlan, Dihal, 2019). За даними інших досліджень, позитивно сприймалися лише винагорода й легкість (Sartori, Bocca, 2023).

Також предметом аналізу ставали джерела, на які покликаються журналісти. Зокрема, на матеріалі преси Великобританії з'ясовано, що журналісти набагато частіше покликаються на представників бізнесу, ніж на політиків чи науковців. Це пов'язано з тим, що близько 60 % публікацій стосуються бізнесу – нових пристроїв, оновлень (Brennen, Howard, Nielsen, 2018). Дослідники американської преси дійшли до подібного висновку: бізнес як джерело інформації згадують у медіа вдвічі частіше від науковців (64,7 % та 29,1 % відповідно), а на політиків покликаються у 23,6 % випадків (Chuan, Tsai, Cho, 2019).

Загалом, дослідники зауважують, що публікації нечасно ставлять під сумнів твердження про можливості ШІ та не заглиблюються в аналіз ризиків цієї технології (Brennen, Howard, Nielsen, 2018; Chuan, Tsai, Cho, 2019). Як зазначають науковці: «Медіа повинні детально досліджувати перспективи та ризики штучного інтелекту, проте їм потрібно ставитися до нього не як до революції, що змінить світ, а радше як до набору технологій, що перебувають у процесі розробки, набору рішень у процесі прийняття та набору проблем, що є в процесі колективного вирішення» (Brennen, Howard, Nielsen, 2018, p. 10). Студії, що застосовують дискурс-аналіз, здебільшого зосереджуються на аналізі тем, у зв'язку з якими висвітлюють ШІ (Lammar, Horst, Müller, 2025); значно менше уваги приділено вивченню дискурсивних стратегій, прийомів чи мовних засобів, якими послуговуються для зображення ШІ в медіа.

У цьому дослідженні прагнемо проаналізувати висвітлення ШІ в медіа з лінгвістичного погляду за допомогою підходу критичного дискурс-аналізу (КДА), а саме соціосеміотичного КДА Т. ван Левена, у межах якого розглядають стратегії дискурсивної (де)легітимації (van Leeuwen, 2008), оскільки репрезентація в медіа впливає на сприйняття ШІ суспільством, законодавчі рішення, а отже і на легітимацію чи делегітимацію використання цієї технології. Дискурсивну легітимацію визначають як «процес, за допомогою якого мовець уповноважує чи надає дозвіл певному типу соціальної поведінки. Отже, це обґрунтування поведінки (ментальної чи фізичної)» (Reyes, 2011, p. 782).

Серед численних класифікацій стратегій (де)легітимації (Cap, 2008; Hart, 2011; Maskau, 2015; Rojo, van Dijk, 1997; Reyes, 2011 та ін.), класичною й однією з найпоширеніших (Чадюк, 2023; Шевко, 2020; Björkvall, Nyström Höög, 2019; Ross, Rivers, 2017; Vaara, 2014) є класифікація Т. ван Левена (van Leeuwen, 2008). Науковець виокремлює чотири стратегії (де)легітимації – авторизація (тактики звернення до персонального авторитету, авторитету експертів, лідерів думок, більшості, традиції та імперсонального авторитету¹), стратегія моральної оцінки (тактики оцінки, аналогії й абстракції), раціоналізація (тактики засобу, звернення до мети, наведення результату, ефекту, дефініція, пояснення, передбачення) та міфопоетика² – (де)легітимація через наративи

¹ Деякі дослідження пропонують розширити імперсональний авторитет технологічним авторитетом, зараховуючи до нього покликання на інші ресурси (онлайн-видання, традиційні медіа тощо) за допомогою наведення знімків екрану та уривків відео як доказів власної позиції (Inwood, Zappavigna, 2024). Водночас постає питання, чи функцію легітимації виконує покликання (технологічний авторитет), чи те, на що це покликання спрямоване (авторитет медіа, автора відео та ін.).

² Стратегію міфопоетики деякі науковці ставлять під сумнів (Zhu, McKenna, 2012), вважаючи, що вона виокремлена на підґрунті інших принципів, ніж три

(детальніше про їхні особливості та мовні засоби – Чадюк, 2023). Науковці також запропонували спосіб автоматизації з'ясування тактик (де)легітимації на матеріалі державної комунікації (Hansson, Page, 2022), проте застосувати їх до медійних текстів не можна, оскільки вони ґрунтуються на кліше офіційно-ділового стилю. Тож у цьому дослідженні послуговуємося маркерами (де)легітимації, що виокремлені на матеріалі українських медіа (Чадюк, 2023).

Мета дослідження – проаналізувати тактики дискурсивної (де)легітимації ШІ в новинних заголовках. Матеріалом стали 311 заголовків ТСН, видань та проєктів медіагрупи Української правди (www.pravda.com.ua, epravda.com.ua, life.pravda.com.ua, mezha.media, оскільки їхні статті висвітлюються в стрічці новин Української правди) та УНІАН за період серпень 2024 – липень 2025, які відображаються на пошуковий запит «штучний інтелект». Заголовки підсумовують інформацію, наведену в публікації, мають сугестивний потенціал (Руденко, 2022), а також великою мірою формують розуміння того, що відбулося, адже вони є першим, що читають. Їхній вплив посилюється і форматом стрічки новин, що змінив споживання інформації: як зазначає К. Сіріньок-Долгарьова, «в 60–80 відсотках випадків саму новину навіть не читають, а отримують уявлення про неї із заголовка» (Сіріньок-Долгарьова, 2012, с. 27). Вибір медіа зумовлений їхньою популярністю в зазначений період, згідно з дослідженнями USAID-Internews «Українські медіа, ставлення та довіра у 2024 році» та Інституту масової інформації (Машкова, Романюк, 2025; Українські медіа, ставлення та довіра у 2024 р.).

Під час аналізу було з'ясовано, що 97 з 311 заголовків не містять стратегій (де)легітимації

попередні стратегії. А С. Бенетт вважає її визначення занадто широким і пропонує розглядати цю стратегію як звернення до певних суспільно значущих наративів історії (Bennett, 2022).

(повідомляють про оновлення моделей, є розважальними новинами), напр.: *Японці створюють «прапну машину для людей» зі штучним інтелектом* (24.11.2024, www.unian.ua); *У Microsoft Edge з'явився новий режим Copilot із розширеними функціями ШІ* (29.07.25, mezha.media). Решту заголовків було розподілено за стратегіями: 92 заголовки містять стратегії легітимації ШІ, 122 заголовки – стратегії делегітимації.

Виклад основного матеріалу дослідження. Аналіз заголовків засвідчив, що найчастіше в медіа легітимують тактиками наведення результату, оцінки, абстракції, засобу, звернення до позиційного авторитету. Делегітимують здебільшого тактиками абстракції, наведення результату, оцінки, ефекту, засобу.

Стратегія моральної оцінки. (Де)легітимацію тактикою абстракції реалізують за допомогою іменників і дієслів, значення яких має позитивні чи негативні конотації та пов'язане з дотриманням чи порушенням цінностей (van Leeuwen, 2008). Для легітимації застосовують іменники на зразок «помічник», «прорив», «сенсація», дієслово «перевершити», вислови «змінити хід війни», «творити дива» тощо, напр.: *Китайський штучний інтелект творить дива!* (02.02.2025, tsn.ua); *Історична сенсація: ШІ розкрив 500-річну таємницю найвідомішої картини Рафаеля* (05.05.2025, tsn.ua). Для делегітимації застосовують іменники «брехня», «вразливість», «галюцинації», «загроза», «захоплення», «лиходій», «небезпека», «провал», «слабкість», «страх», дієслова «налякати», «напасти», «обманювати», «підвести», «провалити» та ін., напр.: *Штучний інтелект може стати ідеальним лиходієм – шокувальний експеримент* (24.01.2025, tsn.ua); *У Китаї робот-гуманоїд зі штучним інтелектом «напав» на своїх розробників* (04.05.2025, tsn.ua).

Часом журналісти послаблюють категоричність висловлювань, репрезентуючи

певну дію як потенційну, напр.: *Алгоритми ШІ можуть обмежити свободу людей* (01.01.2025, tsn.ua); *Штучний інтелект здатен обманювати, шантажувати і мстити: нове дослідження вчених* (23.06.2025, tsn.ua). З цією метою використовують і дієслова у формі майбутнього часу, пор.: *Кого насамперед замінить ШІ в найближчі 5 років: список професій* (14.03.2025, www.unian.ua) / *Штучний інтелект заміняє людей!* (17.10.2024, tsn.ua). Натомість дієслова у формі теперішнього часу зображають дії ШІ як нагальні, напр.: *ШІ вчиться змінювати вашу думку і робить це краще за людей – дослідження* (29.04.2025, www.unian.ua); *ШІ-пошуковик Google «краде» новини та сам відповідає на запитання: світові медіа обурені* (22.05.2025, www.unian.ua).

Зазвичай маркерами (де)легітимації тактикою оцінки є прикметники та прислівники (Чадюк, 2023). Зокрема, для легітимації використовують прикметники «потужний», «розумний», «сильний», «точний», «успішний» тощо, часто наводячи форму вищого чи найвищого ступеня порівняння, напр.: *У ChatGPT з'явився найпотужніший на сьогодні ШІ-асистент для програмування* (17.05.2025, www.unian.ua); *ШІ-лікар Microsoft ставить діагнози в 4 рази точніше за людей і на 20 % дешевше* (01.07.2025, www.unian.ua). Для делегітимації застосовують прикметники «вразливий», «небезпечний», «самовпевнений», «скандальний», «тривожний», «упереджений», «шокувальний», прислівник «тривожно» тощо, напр.: *Захоплено, але тривожно: творець ChatGPT хоче, щоб ШІ пам'ятав «усе ваше життя»* (16.05.2025, www.unian.ua); *В електромобілі Tesla вбудують скандальний ШІ Grok AI Ілона Маска* (11.07.2025, www.unian.ua).

Для (де)легітимації тактикою аналогії певну дію чи об'єкт порівнюють з іншою дією, об'єктом чи суб'єктом, які сприймають як (не)правильні або (не)легітимні. У такий спосіб оцінку об'єкта порівняння прагнуть поширити на суб'єкт порівняння (van Leeuwen, 2008). Для легітимації штучний інтелект порівнюють з

людським, напр.: *OpenAI випустила першу версію ШІ з «мозком» на рівні доктора наук* (13.09.2024, www.unian.ua); *Штучний інтелект виявляє схожість із людським мисленням – дослідження китайських вчених* (16.06.2025, tsn.ua). Для єдиного прикладу делегітимації ШІ в медіа тактикою аналогії штучний інтелект теж порівнюють з людьми, але зосереджуються на негативних рисах характеру: *Штучний інтелект може бути самовпевненим і упередженим, як і люди – вчені* (05.05.2025, www.unian.ua).

У деяких заголовках ШІ репрезентовано як технологію, що виконує завдання краще за людей, напр.: *ШІ-Олімпіада чекає: робот-тенісист від Google зумів обіграти людину* (11.08.2024, www.unian.ua); *ШІ від Google за два дні розв'язав проблему, над якою вчені билися десятиліттями* (17.03.25, www.unian.ua). Такий спосіб репрезентації ШІ не новий: ще в 1984 році медіа зауважували, що «експертні системи» можуть перевершувати фахівців (Bunz, Braghieri, 2022). Водночас, хоч і значно рідше, у новинах згадують, що ШІ не може виконати прості завдання, напр.: *Штучний інтелект провалив завдання, яке під силу кожному школяреві яку слабкість виявили дослідники* (18.05.2025, tsn.ua); *Найрозумніші моделі штучного інтелекту не склали українське ЗНО – дослідження* (17.07.2025, www.unian.ua).

Стратегія раціоналізації. Тактики ефекту й наведення результату (де)легітимують учинки через наведення їхніх позитивних чи негативних наслідків. Відмінність між ними полягає в суб'єкті дії: у тактиці результату він експліцитний, «натомість під час використання тактики наведення ефекту, суб'єкт дії чітко не визначений або імпліцитно представлений у реченнєвій структурі чи тексті» (Чадюк, 2023, с. 121). Легітимацію тактикою наведення результату здійснюють, указуючи на позитивні наслідки дій ШІ, напр.: *Штучний інтелект підвищив рівень влучності українських дронів*

майже вдвічі – ЗМІ (14.10.2024, www.unian.ua); *Штучний інтелект розшифрував стародавні написи, що були недосяжні для вчених* (26.07.2025, tsn.ua). Зазвичай указують на внесок в оборону України³, науку, покращення якості життя (збільшення його тривалості, дводенний робочий тиждень). Натомість для делегітимації вказують на негативні наслідки дій ШІ, напр.: *ШІ зібрав сотні людей на парад на честь Геловіна, якого не існувало: деталі* (01.11.2204, tsn.ua); *ChatGPT і психоз: як штучний інтелект довів чоловіка до лікарні* (23.07.2025, tsn.ua). Найчастіше в новинах згадують про втрату робочих місць, наслідки помилок і обмежень штучного інтелекту. Водночас наведені позитивні й негативні результати часто є потенційними, напр.: *Сучасна реальність: як ШІ може вберегти родину від побутових сварок* (15.05.2025, tsn.ua); *Близький конкурент OpenAI заявив, що ШІ подвоїть тривалість життя людини за 10 років* (27.01.2025, www.unian.ua). Як зазначають дослідники, наведення потенційних результатів у репрезентації ШІ може розмивати межі між тим, що ШІ справді здатен зробити, і тим, що він, можливо, зможе досягти (Brennen, Howard, Nielsen, 2018).

Поширеність (де)легітимації тактикою результату означає, що штучний інтелект часто зображають суб'єктом дії, що зумовлено низкою об'єктивних причин, як-от відсутність даних про тих, хто його використовує, стислість заголовка. Водночас така репрезентація призводить до кількох наслідків. По-перше, у тактиках легітимації це може применшувати роль тих, хто скеровував дії ШІ, наприклад, науковців, урядовців, і формувати хибне враження, що ШІ може виконувати певні завдання без будь-якої участі людей. Для уникнення цього журналісти послуговуються дієсловом «допомогти», яке увиразнює, що ШІ – це інструмент, а головну

³ Кількісні дослідження підтверджують, що ШІ найчастіше згадують у контексті допомоги для оборони України (Marketing Media Review, 2024).

роль відіграють ті, хто його використовує, напр.: *Штучний інтелект допоміг знайти пошкодження базилики Святого Петра у Ватикані* (12.11.2024, www.unian.ua); *Штучний інтелект допоміг розшифрувати обвуглений 2000-річний сувій* (06.02.2025, tsn.ua).

По-друге, висвітлення штучного інтелекту як суб'єкта дії в тактиках делегітимації призводить до репрезентації його як єдиного відповідального за певні вчинки, пор.: *Туристи застрягли в горах, довірившись ChatGPT* (10.01.2025, www.unian.ua) / *ChatGPT відправив туристів у пастку: у Польщі мандрівники застрягли на гірському перевалі* (10.01.2025, tsn.ua); *Alibaba і Tencent тимчасово вимкнули свої чат-боти з ШІ через вступні іспити в Китаї* (09.06.2025, mezha.media) / *У Китаї чат-боти перестали допомагати школярам під час важливих іспитів* (09.06.2025, www.unian.ua). У других варіантах висвітлення події штучний інтелект репрезентовано як ініціатора дії, що посилює сенсаційність заголовка.

Дослідники зауважують, що в медіа важливо вказувати тих, хто ухвалив рішення, наприклад, запрограмувати чи використати ШІ в певний спосіб; до того ж це репрезентувало б цю технологію як більш зрозумілу й контрольовану (Weber, 2025), пор.: *Білл Гейтс спрогнозував, що ШІ забере роботу у людей вже в найближчі 10 років* (29.03.2025, tsn.ua) / *Айтішників масово звільняють по всьому світу, ШІ виявився ефективнішим – Layoffs* (05.05.2025, www.unian.ua). У першому прикладі зменшення робочих місць репрезентовано як неминуче й неконтрольоване явище, а в другому – (імпліцитно) як рішення керівників компаній послуговуватися ШІ, а не наймати людей.

Персоніфікація штучного інтелекту в медіа може посилюватися й вибором лексики: часом щодо технології застосовують слова, які зазвичай використовують щодо істот, напр.: *Чи здатний штучний інтелект відчувати біль: науковці провели низку експериментів* (25.01.2025, tsn.ua); *ШІ довірили керувати*

справжнім магазином – за місяць він збанкрутував і збожеволів (30.06.2025, www.unian.ua). Убачаємо потребу в подальшому дослідженні висвітлення ШІ як автономного, оскільки, як зазначають дослідники (Sartori, Восса, 2023), така репрезентація може посилити занепокоєння щодо ШІ серед населення.

Тактику ефекту використовують рідше, ніж тактику наведення результату, і лише для делегітимації. Ефект від дій ШІ часто зображають як потенційний, напр.: *«Нічого особистого»: Duolingo припинить наймати співробітників, чия роботу може робити ШІ* (29.04.2025, www.unian.ua); *Деякі регіони Європи можуть залишитися без водних запасів через розвиток ШІ – Politico* (28.05.2025, www.unian.ua). Зазвичай ШІ делегітимують через вказівку на втрату робочих місць і появу нових способів шахрайства.

(Де)легітимацію *тактикою засобу* реалізують за схемою «здійснення чогось слугує для здійснення чогось іншого» (van Leeuwen, 2008, р. 114), тобто ШІ як інструмент (де)легітимують тим, задля яких цілей його використовують. Легітимація пов'язана з вказівкою на важливість ШІ для медицини, безпеки в інтернеті, захисту України, напр.: *Міноборони: Українські військові виявляють 12 тисяч цілей щотижня завдяки ШІ* (23.09.2024, www.pravda.com.ua); *Учені створили новий метод протидії змійним отрутам за допомогою ШІ* (16.01.2025, tsn.ua). Делегітимують ШІ зазвичай через зображення його як інструмента для неправомірних дій, напр.: *Мін'юст США звинувачує Google у використанні ШІ для зміцнення монополії* (22.04.2025, mezha.media); *У США шахрай за допомогою ШІ видавав себе за Рубіо і спілкувався з чиновниками – ЗМІ* (08.07.2025, www.unian.ua).

Стратегія авторизації. (Де)легітимувати також можуть словами чи діями експертів, лідерів думок, більшості або осіб, що мають високий статус, а також тим, що певний вчинок (не) відповідає законам чи традиціям (Чадюк,

2023; van Leeuwen, 2008). Стратегію авторизації в заголовках застосовують нечасто, але найпоширенішою її тактикою в проаналізованому матеріалі є звернення до позиційного авторитету – осіб, країн, інституцій, що мають високий статус (Чадюк, 2023). Для легітимації послугуються назвами впливових країн (зазвичай США) та компаній чи установ, як-от Netflix, КМДШ. Застосування ними штучного інтелекту може посилити сприйняття цієї технології як надійної, напр.: **Армія США підвищує ставку на ШІ: плани впровадження розширюють** (16.08.2024, www.unian.ua); **ОАЕ запроваджують штучний інтелект у шкільну програму з 2025 року** (05.05.2025, www.pravda.com.ua). Для делегітимації наводять обмеження або заборони, які країни чи відомі компанії накладають на застосування ШІ, напр.: **В Італії обмежили роботу китайського чат-бота DeepSeek** (31.01.2025, www.eurointegration.com.ua); **Microsoft забороняє своїм співробітникам користуватись чат-ботом китайської DeepSeek** (09.05.2025, mezha.media).

Тактикою звернення до авторитету експертів здебільшого делегітмують, напр.: **Експерти закликають терміново видалити DeepSeek з телефону – чат-бот небезпечний** (09.02.2025, www.unian.ua); **Тенісисти Вімблдону скаржаться на ШІ-суддів** (08.07.2025, mezha.media). Винятком є один приклад легітимації: **Isomorphic Labs готує перші випробування на людях ліків, які розробив ШІ** (07.07.2025, mezha.media) – допуск розроблених ШІ ліків до цього етапу тестувань засвідчує їхню якість. Незважаючи на те що тактика звернення до авторитету експертів не є поширеною, журналісти часто згадують у заголовках дослідження для підтвердження достовірності висловленого, напр.: **Шахраї використовують ШІ для зламу Telegram: дослідження** (18.03.2025, tsn.ua); **Штучний інтелект поставив власне «життя» вище за**

людське: тривожне дослідження (27.06.2025, tsn.ua).

Тактикою звернення до авторитету більшості лише легітмують, указуючи на популярність ШІ серед населення, напр.: **ChatGPT використовують 5 % населення світу щотижня – OpenAI** (21.02.2025, mezha.media); **Кількість платних користувачів ChatGPT зросла до 20 млн – The Information** (02.04.2025, mezha.media). Натомість тактикою звернення до імперсонального авторитету лише делегітмують, зазначаючи про порушення етики, судові процеси щодо розробників ШІ, напр.: **Кінокорпорації Disney і Universal судитимуться зі штучним інтелектом Midjourney** (12.06.2025, life.pravda.com.ua); **ШІ порушив одразу три етичні закони – експерти попереджають про небезпеку** (19.07.2025, tsn.ua).

Висновки дослідження та перспективи подальших наукових розвідок. Аналіз публікацій ТСН, Української правди та УНІАН засвідчив, що 29,6% заголовків новин про штучний інтелект містять тактики легітимації, а 39,2% – тактики делегітимації. Зазвичай використання ШІ легітмують через зазначення позитивних наслідків його діяльності (тактики засобу, наведення результату), указівку на його застосування впливовими країнами чи інституціями (тактика звернення до позиційного авторитету), порівняння ШІ з людським інтелектом (тактика аналогії) та використання прикметників, іменників і дієслів з позитивною конотацією на позначення дій ШІ (тактики оцінки та абстракції). Делегітимацію ШІ реалізують у заголовках медіа за допомогою повідомлення про негативні наслідки його діяльності чи неправомірні цілі, задля яких його використовують (тактики засобу, ефекту, наведення результату), а також застосування іменників, дієслів, прикметників і прислівників, що мають негативні конотації (тактики абстракції, оцінки), для опису дій ШІ.

Репрезентація ШІ як суб'єкта дії є ефективним способом забезпечити стислість

повідомлення, що важливо для заголовків, проте таке представлення може применшувати роль тих, хто використовує ШІ (у разі повідомлення здобутків), і репрезентувати ШІ як ініціатора дії, зображаючи певні результати як неминучі або неконтрольовані, посилюючи сенсаційність повідомлення.

Перспективи подальших досліджень убачаємо в більш детальному аналізові

найбільш популярних за кількістю переглядів статей, що дасть змогу зрозуміти, які саме представлення є найбільш впливовими та які прийоми в них застосовано, щоб репрезентувати штучний інтелект.

Література

- Кузнецова, О. (2024). Ознаки російської дезінформації, створеної ШІ в інтернет-ЗМІ, соціальних мережах. *Вісник Національного університету «Львівська політехніка»: журналістика*, № 1(7), с. 79–89. <https://doi.org/10.23939/sjs2024.01.079>
- Легомінова, С., Тищенко, В., Недодай, М., Дьячук, О., Капелюшна, Т. (2024). Штучний інтелект та соціальні мережі: підходи до виявлення фейкової інформації. *Телекомунікаційні та інформаційні технології*, № 4(85), с. 83–89. <https://doi.org/10.31673/2412-4338.2024.044748>
- Машкова, Я., Романюк, О. (2025). Аналітики ІМІ фіксують незначне просідання трафіку онлайн-медіа. Аналіз даних Similarweb. URL: <https://imi.org.ua/monitorings/analitky-imi-fiksuyut-neznachne-prosidannya-trafiku-onlajn-media-analiz-danyh-similarweb-i69051> (26.08.2025)
- Руденко, Н. В. (2022). *Сугестія як засіб формування громадської думки в сучасних англомовних інтернет-виданнях: інформаційно-комунікаційні стратегії та способи їх реалізації*: дис. д-ра філос. в галузі журналістики. Суми, 235 с.
- Сірінюк-Долгарьова, К. Г. (2012). *Глобальний новинний дискурс: тенденції функціонування англомовних інтернет-медіа*. Київ, Запоріжжя : Центр вільної преси, Запорізький національний університет.
- Українські медіа, ставлення та довіра у 2024 р. Опитування USAID-Internews щодо споживання медіа. (2024). URL: https://drive.google.com/file/d/1kwsclr3Qm2QaqaIVv0_l4saWRFY_NXWb/view (26.08.2025)
- Чадюк, М. О. (2023). *Дискурсивні стратегії легітимації та делегітимації в новинних текстах*: дис. д-ра філос. зі спеціальності «Філологія». Київ, 286 с.
- Черемних, І. (2024). Штучний інтелект у медіагалузі й медіаосвіті. Основні виклики та конкурентні переваги. *Communications and communicative technologies*, № 24, с. 123–132. <https://doi.org/10.15421/292413>
- Шевко, Д. Г. (2020). Легітимація агресії проти України в російському офіційному дискурсі. *Стратегічна панорама*, № 1–2, с. 48–57.
- Bennett, S. (2022). Mythopoetic Legitimation and the Recontextualisation of Europe's Foundational Myth. *Journal of Language and Politics*, vol. 21, issue 2, pp. 370–389. <https://doi.org/10.1075/jlp.21070.ben>
- Björkqvall, A., Nyström Höög, C. (2019). Legitimation of value practices, value texts, and core values at public authorities. *Discourse & Communication*, vol. 13, issue 4, pp. 398–414. <https://doi.org/10.1177/1750481319842457>
- Brennen, J. S., Howard, P. N., Nielsen, R. (2018). *An industry-led debate: how UK media cover artificial intelligence*. Oxford : Reuters Institute for the Study of Journalism. <https://doi.org/10.60625/risj-v219-d676>

- Bunz, M., Braghieri, M. (2022). The AI doctor will see you now: assessing the framing of AI in news coverage. *AI & Society*, vol. 37, pp. 9–22. <https://doi.org/10.1007/s00146-021-01145-9>
- Cap, P. (2008). Towards the proximization model of the analysis of legitimization in political discourse. *Journal of Pragmatics*, vol. 40, issue 1, pp. 17–41. <https://doi.org/10.1016/j.pragma.2007.10.002>
- Cave, S., Craig, C., Dihal, K., Dillon, S., Montgomery, J., Singler, B., Taylor, L. (2018). *Portrayals and perceptions of AI and why they matter*. London : The Royal Society. <https://doi.org/10.17863/CAM.34502>
- Cave, S., Coughlan, K., Dihal, K. (2019). «Scary Robots»: Examining Public Responses to AI. *Proceedings of the 2019 AAAI/ACM Conference on AI, Ethics, and Society*, pp. 331–337. <https://doi.org/10.1145/3306618.3314232>
- Cave, S., Dihal, K. (2019). Hopes and Fears for Intelligent Machines in Fiction and Reality. *Nature Machine Intelligence*, issue 1, pp. 74–78. <https://doi.org/10.1038/s42256-019-0020-9>
- Chuan, C., Tsai, W. S., Cho, S. (2019). Framing Artificial Intelligence in American Newspapers. *Proceedings of the 2019 AAAI/ACM Conference on AI, Ethics, and Society*, pp. 339–344. <https://doi.org/10.1145/3306618.3314285>
- Fast, E., Horvitz, E. (2017). Long-Term Trends in the Public Perception of Artificial Intelligence. *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 31, no. 1, pp. 963–969. <https://doi.org/10.1609/aaai.v31i1.10635>
- Fedoriv, Y., Pirozhenko, I., Shuhai, A. (2023). Linguistic Analysis of Human- and AI-Created Content in Academic Discourse. *Journal of Vasyl Stefanyk Precarpathian National University. Philology*, vol. 10, pp. 47–67. <https://doi.org/10.15330/jpnuphil.10.47-67>
- Fleisig, E., Smith, G., Bossi, M., Rustagi, I., Yin, X., Klein, D. (2024). Linguistic Bias in ChatGPT: Language Models Reinforce Dialect Discrimination. *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing*, pp. 13541–13564. <https://doi.org/10.18653/v1/2024.emnlp-main.750>
- González-Arias, C., López-García, X. (2024). Rethinking the Relation between Media and Their Audience: The Discursive Construction of the Risk of Artificial Intelligence in the Press of Belgium, France, Portugal, and Spain. *Journalism and Media*, vol. 5, issue 3, pp. 1023–1037. <https://doi.org/10.3390/journalmedia5030065>
- Hansson, S., Page, R. (2022). Corpus-assisted analysis of legitimation strategies in government social media communication. *Discourse & Communication*, vol. 16, issue 5, pp. 551–571. <https://doi.org/10.1177/17504813221099202>
- Hart, C. (2011). Legitimizing assertions and the logico-rhetorical module: Evidence and epistemic vigilance in media discourse on immigration. *Discourse Studies*, vol. 13, no. 6, pp. 751–769. <https://doi.org/10.1177/1461445611421360>
- Inwood, O., Zappavigna, M. (2024). The legitimation of screenshots as visual evidence in social media: YouTube videos spreading misinformation and disinformation. *Visual Communication*. URL: <https://journals.sagepub.com/doi/full/10.1177/14703572241255664>. <https://doi.org/10.1177/14703572241255664>
- Lammar, D., Horst, M., Müller, R. (2025). AI in the German Media: Narratives of AI-in-Particular and AI-in-General in German Media Reporting About Artificial Intelligence. *Digital Journalism*, pp. 1–19. <https://doi.org/10.1080/21670811.2025.2493759>
- Mackay, R. R. (2015). Multimodal legitimation: Selling Scottish independence. *Discourse & Society*, vol. 26, no. 3, pp. 323–348. <https://doi.org/10.1177/0957926514564737>
- Marketing Media Review. Що писали про штучний інтелект українські медіа в 2023 році: дослідження. (2024). URL: <https://mmr.ua/shho-pysaly-pro-shtuchnyj-intelekt-ukrayinski-media-v-2023-rocz-i-doslidzhennya>

- Newman, N., Ross Arguedas, A., Robertson, C. T., Nielsen, R. K., Fletcher, R. (2025). Digital news report 2025. Oxford : Reuters Institute for the Study of Journalism. <https://doi.org/10.60625/risj-8qqf-jt36>
- Nguyen, D. (2023). How news media frame data risks in their coverage of big data and AI. *Internet Policy Review*, vol. 12, issue 2, pp. 1–30. <https://doi.org/10.14763/2023.2.1708>
- Ouchchy, L., Coin, A., Dubljević, V. (2020). AI in the headlines: the portrayal of the ethical issues of artificial intelligence in the media. *AI & Society*, vol. 35, pp. 927–936. <https://doi.org/10.1007/s00146-020-00965-5>
- Reyes, A. (2011). Strategies of legitimization in political discourse: From words to actions. *Discourse & Society*, vol. 22, no. 6, pp. 781–807. <https://doi.org/10.1177/0957926511419927>
- Rojo, L. M., van Dijk, T. A. (1997). «There was a Problem, and it was Solved!»: Legitimizing the Expulsion of «Illegal» Migrants in Spanish Parliamentary Discourse. *Discourse & Society*, vol. 8, issue 4, pp. 523–566. <https://doi.org/10.1177/0957926597008004005>
- Ross, A. S., Rivers, D. J. (2017). Digital cultures of political participation: internet memes and the discursive delegitimization of the 2016 U.S presidential candidates. *Discourse, Context and Media*, vol. 16, pp. 1–11. <https://doi.org/10.1016/j.dcm.2017.01.001>
- Sartori, L., Bocca, G. (2023). Minding the gap(s): public perceptions of AI and socio-technical imaginaries. *AI & Society*, vol. 38, pp. 443–458. <https://doi.org/10.1007/s00146-022-01422-1>
- Vaara, E. (2014). Struggles over legitimacy in the Eurozone crisis: Discursive legitimization strategies and their ideological underpinnings. *Discourse & Society*, vol. 25, no. 4, pp. 500–518. <https://doi.org/10.1177/0957926514536962>
- van Leeuwen, T. (2008). *Discourse and Practice: New Tools for Critical Discourse Analysis*. Oxford : Oxford University Press. <https://doi.org/10.1093/acprof:oso/9780195323306.001.0001>
- Weber, K. (2025). Public images of artificial intelligence: an overview. *Media – Culture – Social Communication*, no. 21, pp. 9–24. <https://doi.org/10.31648/mcsc.10343>
- Zhu, Y., McKenna, B. (2012). Legitimizing a Chinese takeover of an Australian iconic firm: Revisiting models of media discourse of legitimacy. *Discourse & Society*, vol. 23, no. 5, pp. 525–552. <https://doi.org/10.1177/0957926512452971>

Джерела

TCH <https://tsn.ua/>
 Українська правда <https://www.pravda.com.ua/>
 УНІАН <https://www.unian.ua/>

References

- Bennett, S. (2022). Mythopoetic Legitimation and the Recontextualisation of Europe's Foundational Myth. *Journal of Language and Politics*, vol. 21, issue 2, pp. 370–389. <https://doi.org/10.1075/jlp.21070.ben> (in English).
- Björkqvall, A., Nyström Höög, C. (2019). Legitimation of value practices, value texts, and core values at public authorities. *Discourse & Communication*, vol. 13, issue 4, pp. 398–414. <https://doi.org/10.1177/1750481319842457> (in English).
- Brennen, J. S., Howard, P. N., Nielsen, R. (2018). An industry-led debate: how UK media cover artificial intelligence. Oxford: Reuters Institute for the Study of Journalism. DOI: 10.60625/risj-v219-d676 (in English).
- Bunz, M., Braghieri, M. (2022). The AI doctor will see you now: assessing the framing of AI in news coverage. *AI & Society*, vol. 37, pp. 9–22. <https://doi.org/10.1007/s00146-021-01145-9> (in English).

- Cap, P. (2008). Towards the proximization model of the analysis of legitimization in political discourse. *Journal of Pragmatics*, vol. 40, issue 1, pp. 17–41. <https://doi.org/10.1016/j.pragma.2007.10.002> (in English).
- Cave, S., Craig, C., Dihal, K., Dillon, S., Montgomery, J., Singler, B., Taylor, L. (2018). Portrayals and perceptions of AI and why they matter. London: The Royal Society. <https://doi.org/10.17863/CAM.34502> (in English).
- Cave, S., Coughlan, K., Dihal, K. (2019). «Scary Robots»: Examining Public Responses to AI. *Proceedings of the 2019 AAAI/ACM Conference on AI, Ethics, and Society*, pp. 331–337. <https://doi.org/10.1145/3306618.3314232> (in English).
- Cave, S., Dihal, K. (2019). Hopes and Fears for Intelligent Machines in Fiction and Reality. *Nature Machine Intelligence*, issue 1, pp. 74–78. <https://doi.org/10.1038/s42256-019-0020-9> (in English).
- Chadiuk, M. O. (2023). *Discursive strategies of legitimation and delegitimation in news texts*. (PhD thesis). Kyiv, 286 p. (in Ukrainian).
- Cheremnykh, I. (2024). Artificial intelligence in the media industry and media education. Main challenges and competitive advantages. *Communications and communicative technologies*, no. 24, pp. 123–132. DOI: 10.15421/292413 (in Ukrainian).
- Chuan, C., Tsai, W. S., Cho, S. (2019). Framing Artificial Intelligence in American Newspapers. *Proceedings of the 2019 AAAI/ACM Conference on AI, Ethics, and Society*, pp. 339–344. <https://doi.org/10.1145/3306618.3314285> (in English).
- Fast, E., Horvitz, E. (2017). Long-Term Trends in the Public Perception of Artificial Intelligence. *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 31, no. 1., pp. 963–969. <https://doi.org/10.1609/aaai.v31i1.10635> (in English).
- Fedoriv, Y., Pirozhenko, I., Shuhai, A. (2023). Linguistic Analysis of Human- and AI-Created Content in Academic Discourse. *Journal of Vasyl Stefanyk Precarpathian National University. Philology*, vol. 10, pp. 47–67. <https://doi.org/10.15330/jpnuphil.10.47-67> (in English).
- Fleisig, E., Smith, G., Bossi, M., Rustagi, I., Yin, X., Klein, D. (2024). Linguistic Bias in ChatGPT: Language Models Reinforce Dialect Discrimination. *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing*, pp. 13541–13564. DOI:10.18653/v1/2024.emnlp-main.750 (in English).
- González-Arias, C., López-García, X. (2024). Rethinking the Relation between Media and Their Audience: The Discursive Construction of the Risk of Artificial Intelligence in the Press of Belgium, France, Portugal, and Spain. *Journalism and Media*, vol. 5, issue 3, pp. 1023–1037. <https://doi.org/10.3390/journalmedia5030065> (in English).
- Hansson, S., Page, R. (2022). Corpus-assisted analysis of legitimation strategies in government social media communication. *Discourse & Communication*, vol. 16, issue 5, pp. 551–571. <https://doi.org/10.1177/17504813221099202> (in English).
- Hart, C. (2011). Legitimizing assertions and the logico-rhetorical module: Evidence and epistemic vigilance in media discourse on immigration. *Discourse Studies*, vol. 13, no. 6, pp. 751–769. <https://doi.org/10.1177/1461445611421360> (in English).
- Inwood, O., Zappavigna, M. (2024). The legitimation of screenshots as visual evidence in social media: YouTube videos spreading misinformation and disinformation. *Visual Communication*. URL: <https://journals.sagepub.com/doi/full/10.1177/14703572241255664>. <https://doi.org/10.1177/14703572241255664> (in English).
- Kuznetsova, O. (2024). Features of Russian disinformation created by AI on the Internet media, social networks. *Bulletin of Lviv Polytechnic National University: journalism*, no. 1(7), pp. 79–89. <https://doi.org/10.23939/sjs2024.01.079> (in Ukrainian).

- Lammar, D., Horst, M., Müller, R. (2025). AI in the German Media: Narratives of AI-in-Particular and AI-in-General in German Media Reporting About Artificial Intelligence. *Digital Journalism*, pp.1-19. <https://doi.org/10.1080/21670811.2025.2493759> (in English).
- Lehominova, S., Tyshchenko, V., Nedodai, M., Diachuk, O., Kapeliushna, T. (2024). Artificial intelligence and social networks: approaches to detecting fake information. *Telecommunication and information technologies*, no. 4 (85), pp. 83-89. DOI: 10.31673/2412-4338.2024.044748 (in Ukrainian).
- Mackay, R. R. (2015). Multimodal legitimation: Selling Scottish independence. *Discourse & Society*, vol. 26, no. 3, pp. 323-348. <https://doi.org/10.1177/0957926514564737> (in English).
- Marketing Media Review. (2024). What Ukrainian media wrote about artificial intelligence in 2023: research. Available at: <https://mmr.ua/shho-pysaly-pro-shtuchnyj-intelekt-ukrayinski-media-v-2023-roczy-doslidzhennya> (in Ukrainian).
- Mashkova, Ya., Romaniuk, O. (2025). The Institute of Mass Information analysts record a slight decrease in online media traffic. Similarweb data analysis. Available at: <https://imi.org.ua/monitorings/analitky-imi-fiksuyut-neznachne-prosidannya-trafiku-onlajn-media-analiz-danyh-similarweb-i69051> (26.08.2025) (in Ukrainian).
- Newman, N., Ross Arguedas, A., Robertson, C. T., Nielsen, R. K., Fletcher, R. (2025). Digital news report 2025. Oxford: Reuters Institute for the Study of Journalism. DOI: 10.60625/risj-8qqf-jt36 (in English).
- Nguyen, D. (2023). How news media frame data risks in their coverage of big data and AI. *Internet Policy Review*, vol. 12, issue 2, pp. 1-30. <https://doi.org/10.14763/2023.2.1708> (in English).
- Ouchchy, L., Coin, A., Dubljević, V. (2020). AI in the headlines: the portrayal of the ethical issues of artificial intelligence in the media. *AI & Society*, vol. 35, pp. 927-936. <https://doi.org/10.1007/s00146-020-00965-5> (in English).
- Reyes, A. (2011). Strategies of legitimization in political discourse: From words to actions. *Discourse & Society*, vol. 22, no. 6, pp. 781-807. <https://doi.org/10.1177/0957926511419927> (in English).
- Rojo, L. M., van Dijk, T. A. (1997). «There was a Problem, and it was Solved!»: Legitimizing the Expulsion of «Illegal» Migrants in Spanish Parliamentary Discourse. *Discourse & Society*, vol. 8, issue 4, pp. 523-566. <https://doi.org/10.1177/0957926597008004005> (in English).
- Ross, A. S., Rivers, D. J. (2017). Digital cultures of political participation: internet memes and the discursive delegitimization of the 2016 U.S presidential candidates. *Discourse, Context and Media*, vol. 16, pp. 1-11. <https://doi.org/10.1016/j.dcm.2017.01.001> (in English).
- Rudenko, N. V. (2022). *Suggestion as a means of public opinion shaping in modern English-language internet editions: information and communication strategies and ways to implement them*. (PhD thesis). Sumy, 235 p. (in Ukrainian).
- Sartori, L., Bocca, G. (2023). Minding the gap(s): public perceptions of AI and socio-technical imaginaries. *AI & Society*, vol. 38, pp. 443-458. <https://doi.org/10.1007/s00146-022-01422-1> (in English).
- Shevko, D. H. (2020). Legitimation of aggression against Ukraine in Russian official discourse. *Strategic panorama*, no. 1-2, pp. 48-57. (in Ukrainian).
- Sirinok-Dolharova, K. H. (2012). *Global News Discourse: Trends in the Functioning of English-Language Internet Media*. Kyiv and Zaporizhzhia, Center for Free Press; Zaporizhzhia National University. (in Ukrainian).
- Ukrainian media, attitudes and trust in 2024. USAID-Internews survey on media consumption (2024). Available at: https://drive.google.com/file/d/1kwsclr3Qm2QaqaIVv0_l4saWRFY_NXWb/view (in Ukrainian).

- Vaara, E. (2014). Struggles over legitimacy in the Eurozone crisis: Discursive legitimation strategies and their ideological underpinnings. *Discourse & Society*, vol. 25, no. 4, pp. 500–518. <https://doi.org/10.1177/0957926514536962> (in English).
- van Leeuwen, T. (2008). *Discourse and Practice: New Tools for Critical Discourse Analysis*. Oxford : Oxford University Press. <https://doi.org/10.1093/acprof:oso/9780195323306.001.0001> (in English).
- Weber, K. (2025). Public images of artificial intelligence: an overview. *Media – Culture – Social Communication*, no. 21, pp. 9–24. <https://doi.org/10.31648/mcsc.10343> (in English).
- Zhu, Y., McKenna, B. (2012). Legitimizing a Chinese takeover of an Australian iconic firm: Revisiting models of media discourse of legitimacy. *Discourse & Society*, vol. 23, no. 5, pp. 525–552. <https://doi.org/10.1177/0957926512452971> (in English).

Автор

Марія Чадюк, доктор філософії зі спеціальності «Філологія», старший викладач кафедри української мови
e-mail: m.chadiuk@ukma.edu.ua
<https://orcid.org/0000-0002-6727-1126>

Author

Mariia Chadiuk, PhD in Philology, Senior Lecturer at the Department of the Ukrainian Language
e-mail: m.chadiuk@ukma.edu.ua
<https://orcid.org/0000-0002-6727-1126>

Конфлікт інтересів

Автор засвідчує про відсутність конфлікту інтересів

Conflict of Interests

The author declares that she has no conflict of interest