

Методологічні виклики штучного інтелекту в політичній науці: епістемологічний аналіз

Анотація. У статті досліджено основні проблеми, які виникають під час використання штучного інтелекту (ШІ) в політичній науці, та вплив нових технологій на традиційні методи вивчення політики та труднощі, які виникають у зв'язку з цим у дослідників.

Мета роботи – з'ясувати можливості та обмеження ШІ як інструменту політичної науки, виявити головні напрями покращення дослідницьких методів та окреслити перспективи поєднання традиційних підходів із новими алгоритмічними технологіями.

Дослідження базується на філософському підході до аналізу методологічних проблем та використовує методи концептуального аналізу й порівняння традиційних і цифрових методів політичної науки. Застосовано міждисциплінарний підхід, який поєднує інструментарій політичної теорії, філософії науки та теорії пізнання. Особлива увага приділяється проблемам перевірки достовірності алгоритмічних висновків та складнощам інтерпретації результатів машинного навчання.

Виявлено п'ять ключових методологічних викликів: проблеми перевірки та тлумачення алгоритмічних результатів, труднощі поєднання якісних і кількісних підходів, складності встановлення причинно-наслідкових зв'язків алгоритмічними методами, часові та контекстуальні обмеження машинного навчання, а також етичні проблеми автоматизації політичного аналізу.

Встановлено, що ШІ створює нову реальність для політичної науки, яка потребує розробки інноваційних процедур перевірки та гібридних методологічних підходів. Показано, що ефективне використання ШІ вимагає не заміни традиційних методів, а їх творчого поєднання з алгоритмічними підходами.

Інтеграція ШІ у політичну науку потребує переосмислення основ дисципліни та розробки нових стандартів наукової достовірності. Формування гібридних методологій, що поєднують людське судження з машинним аналізом, є ключовою умовою для збереження наукової строгості та практичної користі політичних досліджень у цифрову епоху.

Ключові слова: штучний інтелект, методологія політичної науки, епістемологія, валідація, каузальність, гібридні методології, машинне навчання, алгоритмічний аналіз.

Автор висловлює щире подяку фахівцям Вінницького національного технічного університету за цінні експертні консультації щодо технічних аспектів штучного інтелекту, які дозволили адаптувати складні технологічні концепції до завдань політологічного дослідження

Валерій Корнієнко,
доктор політичних наук,
професор,
Вінницький національний
технічний університет

Valerii Kornienko,
Doctor of Political Sciences,
Professor,
Vinnitsia National Technical
University

ORCID: 0000-0001-6285-5707
valkorney1958@gmail.com

Methodological Challenges of Artificial Intelligence in Political Science: An Epistemological Analysis

Abstract. The article examines the main challenges of applying artificial intelligence (AI) in political science. The author explains how new technologies influence traditional methods of studying politics and the difficulties this creates for researchers. The study aims to clarify AI's potential and limitations as a tool in political science, identify key areas for improving research methods, and outline prospects for combining traditional approaches with new algorithmic techniques.

It takes a philosophical approach to analyzing methodological issues, using methods of conceptual analysis and comparing traditional and digital political science approaches. An interdisciplinary approach is employed, integrating tools from political theory, philosophy of science, and epistemology. Special attention is given to issues related to verifying the reliability of algorithmic conclusions and interpreting results from machine learning.

Five major methodological challenges are identified: verifying and interpreting algorithmic results, combining qualitative and quantitative methods, establishing causal relationships through algorithms, understanding the temporal and contextual limits of machine learning, and addressing ethical concerns related to the automation of political analysis.

The conclusion is that AI introduces a new reality for political science that requires developing innovative verification techniques and hybrid research methods. Effectively utilizing AI does not mean replacing traditional methods, but instead creatively blending them with algorithmic techniques. Incorporating AI into political science calls for rethinking core principles and establishing new standards of scientific reliability. Developing hybrid methodologies that combine human judgment with machine analysis is essential to maintaining scientific rigor and practical relevance in the digital age.

Keywords: artificial intelligence, political science methodology, epistemology, validation, causality, hybrid methodologies, machine learning, algorithmic analysis.

Постановка проблеми. Ще в шістдесяті роки минулого століття наша вчителька математики здивувала нас: японські школярі не знають таблицю множення напам'ять – у них є калькулятори... Навіть попри те, що японські рахунки (соробан) викладалися в школах понад 500 років і були глибоко вкорінені в культурі Японії, введення технологій змінило освітню систему.

Ця історія ілюструє важливу закономірність: технічний прогрес неможливо зупинити. Питання не в тому, чи варто це робити, а в тому, як використати правильно та ефективно.

Сьогодні ми стоїмо перед аналогічною ситуацією зі штучним інтелектом (ШІ) – він активно впроваджується в усі сфери життя, включно з політикою та політичною наукою. Це "зіткнення" створює нові методологічні виклики: як валідувати результати алгоритмічного аналізу політичних явищ, коли традиційні методи перевірки не завжди працюють? Чи здатні машини адекватно "розуміти"

природу політичних процесів, враховуючи іронію, контекст та приховані мотиви? Як поєднувати традиційні якісні методи з алгоритмічними підходами?

Проблема використання ШІ в політології виходить далеко за межі технічних питань, торкаючись фундаментальних засад наукової об'єктивності, етичної відповідальності дослідників та демократичних цінностей суспільства. Нам важливо зберегти гуманітарну природу дисципліни, водночас використовуючи переваги сучасних технологій для кращого розуміння політичних явищ.

Аналіз останніх досліджень і публікацій. Проблематика методологічних викликів ШІ у політичній науці порівняно нещодавно стала предметом систематичного наукового аналізу. На особливу увагу заслуговують думки зарубіжних дослідників, які намагаються концептуалізувати епістемологічні наслідки "алгоритмічного повороту" у соціально-гуманітарних науках.

Важливим є дослідницький доробок китайських науковців. У статті "Від ChatGPT, DALL-E 3 до Sora: як генеративний штучний інтелект змінив дослідження і сервіси в цифрових гуманітарних науках?" [Liu et al., 2024] автори розглядають, як великі генеративні моделі штучного інтелекту (наприклад, *ChatGPT*, *DALL-E 3*) формують новий парадигматичний "п'ятий" підхід у науковому дослідженні, особливо в гуманітарних науках. Вони описують використання ШІ для обробки й збереження стародавніх текстів, їх класифікації, генерування контенту та культурної інтеграції, а також обговорюють потенціал цих технологій у культурній спадщині й художніх інноваціях.

Чжен Хуанг [Huang, 2022] розглядає можливості впровадження нейро-символічного ШІ в гуманітарні та соціальні науки. Стаття розкриває зв'язки між нейро-символічним ШІ та гуманітарними і соціальними науками, узагальнює останні тенденції розвитку та типові застосування, а також пропонує можливий шлях розширення плюралістичних методологій у гуманітарних і соціальних науках для адаптації до епохи великих даних.

Хаоюе Чжао [Zhao, 2023] аналізує потенційні переваги і ризики застосування ШІ у гуманітарних науках. Він звертається до прикладів з використанням ChatGPT і проводить діалектичну рефлексію: ШІ – не заміна людині, але інструмент, що змінює методологію досліджень та текстову продукцію.

Лін Чжіці [Lin, 2024] розглядає ризики, пов'язані з використанням ШІ, вказуючи на необхідність збереження балансу між перевагами технологій та потенційними загрозами: підривом довіри, поширенням помилкових даних, нерівністю доступу. У політичному контексті ця робота демонструє, що регулювання ШІ стосується не лише технічних чи етичних питань, але й управління ключовими інститутами науки та суспільного знання.

Дослідження Мередіт Дедема та Ронгцянь Ма [Dedema & Ma, 2024] ґрунтується на міжнародному опитуванні дослідників у сфері цифрових гуманітаристик щодо їхнього ставлення та досвіду використання генеративного ШІ (наприклад, ChatGPT) в науковій роботі. В статті аналізуються емпіричні дані про репутацію і позицію ШІ в гуманітарних науках.

Привертає увагу "Огляд застосування штучного інтелекту в дослідженнях гуманітарних наук" [Chapinal-Heras & Díaz-Sánchez, 2024], який узагальнює, як саме застосовують технології штучного інтелекту в гуманітарних дисциплінах. Проаналізовано основні алгоритми, програмні інструменти та їхні досягнення – і водночас підкреслено нестачу емпіричних кейсів.

Варто звернути увагу на дослідження італійських науковців Маріфлора Карузо та Алессандро Спададо [Caruso & Spadaro, 2024]. Автори підкреслюють, що сучасні інструменти ШІ докорінно змінюють методологію досліджень у гуманітаристиці, відкриваючи нові можливості для інтерпретації культурної спадщини, проте водночас породжують етичні виклики. Особливу увагу автори приділяють питанням відповідальності дослідників у використанні ШІ, закликаючи до міждисциплінарного та етичного обговорення для забезпечення балансу між інноваціями та збереженням наукової достовірності.

Робота індійських дослідників Маткунті, Анантанагу і Менон [Mathkunti et al., 2023], які працюють у сфері штучного інтелекту та його застосування в гуманітарних і соціальних науках, також окреслює ШІ як трансформаційний чинник гуманітарних і соціальних наук, демонструючи як інноваційні можливості, так і методологічні ризики, що потребують міждисциплінарного й етичного осмислення. Автори застерігають, що некритичне застосування ШІ може призвести до редукціонізму в гуманітарних дослідженнях, тому пропонують розвивати баланс між алгоритмічними підходами та людською інтерпретацією.

Одним із перших, хто вдався до серйозних роздумів про методологічні виклики ШІ для політичної науки, був Джеймс Еванс – відомий американський соціолог науки, який розглядав проблеми валідації комп'ютерних моделей політичних процесів [Evans, 2008]. Дж. Еванс зауважив, що механічне перенесення стандартів природничих наук на галузь комп'ютерного моделювання політичних явищ може призвести до спотворення їх сутності.

Не можна оминати увагою думки Данкана Воттса, який у своєму фундаментальному дослідженні "Все очевидно: чому здоровий глузд нас підводить" розкрив фундаментальні проблеми каузального аналізу в соціальних науках [Watts, 2011]. Його робота демонструє, як ШІ може як посилювати, так і ускладнювати проблеми встановлення причинно-наслідкових зв'язків у політичних процесах. Примітно, що Д. Воттс одним із перших звернув увагу на "ілюзію розуміння", коли складні алгоритмічні моделі створюють хибне відчуття пояснення політичних явищ.

Дослідження Кеті О'Ніл та Вірджинії Дігнум – "Зброя математичного знищення" [O'Neil, 2016] детально аналізує, як алгоритмічні упередження можуть спотворювати політичний аналіз та впливати на демократичні процеси. В. Дігнум, своєю чергою, у дослідженні "Відповідальний штучний інтелект" запропонувала концептуальну рамку для етично відповідального використання ШІ у суспільних науках.

Не менш важливими у контексті дослідження є робота Пітера Норвіга та Стюарта Рассела, які у своєму підручнику "Штучний інтелект: сучасний підхід" [Russell & Norvig, 2020] присвятили окремий розділ методологічним проблемам застосування ШІ у гуманітарних науках. Хоча їхні міркування стосуються ширшого кола питань, вони надають цінні інсайти для розуміння специфіки політичного застосування ШІ-технологій.

Аві Голдфарб – відомий канадський економіст, який спеціалізується на цифровій економіці, штучному інтелекті та інноваціях, у своїй статті "Призупинити дослідження в галузі штучного інтелекту? Розуміння викликів політики в галузі ШІ" [Goldfarb, 2024] розбирає складну дилему сучасності: чи варто зупинити розвиток штучного інтелекту? А. Голдфарб показує, що суспільство серйозно стурбоване: люди бояться втратити роботу через автоматизацію, хвилюються за безпеку нових технологій, побоюються, що ШІ може зашкодити демократії. Науковець доходить висновку, що повна зупинка – не вихід. Натомість потрібен розумний баланс у регулюванні, який враховуватиме і ризики, і величезні можливості штучного інтелекту.

Інга Ульніцане та Тєро Ерккіля [Ulinicane & Erkkilä, 2023] досліджують ШІ як політичне явище, акцентуючи на потребі вироблення спеціальних політик та стратегій. Автори аналізують ключові напрями формування політики у сфері ШІ, зокрема регулювання етики, безпеки та відповідальності, а також вплив ШІ на управління, економіку і демократію.

Вальтер Йошія [Walter, 2024] пропонує метафору "перегонів до Місяця" для опису глобальної конкуренції у сфері ШІ. Автор аналізує, як різні країни намагаються встановити власні правила гри, які соціально-економічні наслідки це має та які механізми глобального врядування потрібні для уникнення нерівностей і конфліктів. У політичному плані робота показує, що регулювання ШІ – це не лише питання технологій, а й боротьба за геополітичний вплив і соціальну стабільність.

Практичний вимір цих процесів демонструють експертки з міграційної політики Бедуші та МакАліфф [Beduschi & McAuliffe, 2022], які детально аналізують вплив ШІ на весь міграційний цикл – від підготовки до виїзду до інтеграції в країні, що приймає. Дослідниці розглядають застосування для автоматизації візових процедур, управління кордонами та надання послуг мігрантам, водночас підкреслюючи серйозні ризики алгоритмічної упередженості та посилення цифрового розриву між країнами.

Конкретний національний кейс становить робота Лемке, Трейна та Вароне [Lemke et al., 2024], які досліджують, як ШІ став політичною проблемою в Німеччині. Науковці з'ясували, що кожна група по-своєму бачить головні проблеми: одні вважають найважливішою безпеку, інші – економічний розвиток, треті наголошують на етичних питаннях або необхідності конкурувати з іншими країнами. Головний висновок: політичні рішення щодо ШІ залежать не тільки від самих технологій, а набагато важливіше те, як різні політичні сили обговорюють ці питання та взаємодіють між собою.

Проблеми ШІ є важливими і для українських науковців, які активно досліджують його роль у політичних процесах та державному управлінні,

опрацьовують різні аспекти використання ШІ. При цьому Іван Сабов [Сабов, 2018] підкреслює – поки що не існує комплексного дослідження застосування засобів ШІ у політиці.

Сергій Бут визначає головні напрями використання ШІ в суспільно-політичних процесах. Особливу увагу приділено ризикам та викликам, які виникають через недосконале правове регулювання, можливі загрози неконтрольованого застосування ШІ в політичних цілях та етичні проблеми, пов'язані з цими технологіями [Бут, 2025].

Найбільш амбітну концептуалізацію здійснив Михайло Янишівський [Янишівський, 2024], який дослідив роль ШІ та його вплив на демократичні процеси, етичні питання, безпеку та інші соціально важливі сфери. Ним розроблена концептуальна модель "Піраміда особливостей політики у галузі ШІ", що систематизує його фундаментальні характеристики, ключові виклики, стратегії та інструменти, а також унікальні риси.

Ірина Радіонова [Радіонова, 2023] розглядає потенційні загрози, пов'язані з військовим використанням, кібербезпекою, контролем даних та критичною інфраструктурою, і показує, як держави та міжнародні організації реагують на ці виклики. Особливо наголошується на взаємозв'язку між технологічними ризиками та політичними рішеннями, а також на необхідності врегулювання ШІ для гарантування національної та глобальної безпеки.

Значення ШІ в сучасній дипломатії, зокрема як інструменту цифрової підтримки для аналізу даних, прийняття стратегічних рішень і моніторингу міжнародних процесів аналізує Ірина Климчук [Климчук, 2025]. Авторка детально розглядає можливості використання ШІ для прогнозування конфліктів, боротьби з дезінформацією, гарантування кібербезпеки, обробки великих обсягів інформації, аналізу громадської думки та автоматизації рутинних завдань дипломатичних служб.

Системний погляд на національні стратегії пропонує Олексій Костенко [Костенко, 2022], чиє дослідження демонструє важливість стратегічного підходу у розвитку ШІ на національному рівні через порівняльний аналіз провідних світових практик. Автор аналізує понад 60 національних стратегій, особливо відзначаючи США, Китай, Канаду, Велику Британію, Японію, Францію, Німеччину, Південну Корею та країни ЄС. Автор підкреслює, що українська стратегія розвитку ШІ не є ефективною через недостатність інвестицій, відсутність стимулюючих механізмів та конкретних дій, наголошуючи на важливості прийняття Кабінетом Міністрів відповідної стратегії.

Проглядається важлива і одночасно головна мета наукового співтовариства – знайти баланс між потужними можливостями ШІ та збереженням демократичних цінностей і наукової об'єктивності.

Виділення невирішених раніше частин загальної проблеми. Більшість досліджень зосереджується на конкретних застосуваннях ШІ-технологій. Проте системний аналіз того, як штучний інтелект змінює основи політичної науки, залишається недостатньо розробленим. Постає питання: чи формується

нова методологічна парадигма, чи йдеться лише про технічне вдосконалення існуючих підходів?

Особливо проблемною є оцінка достовірності алгоритмічних результатів у політичному контексті. Традиційні критерії наукової валідності виявляються недостатніми для машинного навчання. Алгоритми можуть виявляти складні закономірності у великих масивах політичних даних, але їх інтерпретація з погляду політичної теорії залишається відкритою проблемою.

Складною є проблема поєднання традиційних політологічних підходів з алгоритмічними методами. Політична наука не виробила чітких принципів створення "гібридних методологій", які б органічно інтегрували людське судження та машинний аналіз. Постають фундаментальні питання про межі машинного пізнання політичних явищ.

Проблемною є каузальність в алгоритмічному аналізі. Більшість ШІ-методів виявляє кореляції та закономірності, а не причинно-наслідкові зв'язки. Потрібно пригадати важливий внесок Карла Поппера і Томаса Куна, які застерігали проти змішування кореляції та каузації. Питання про здатність алгоритмів встановлювати справжні каузальні зв'язки у політиці залишається дискусійним.

Недостатньо досліджена проблема часової специфіки політичних процесів. Алгоритми, навчені на історичних даних, часто застосовуються для аналізу сучасних явищ. Це створює методологічну проблему: політичні процеси мають історичну специфіку та здатність до якісних змін, що може робити історичні закономірності нерелевантними.

Етичні аспекти використання ШІ в політичній науці потребують особливої уваги. Автоматизація політичного аналізу може мати далекосяжні наслідки для демократичних процесів. Хто відповідає за алгоритмічні рекомендації у політиці? Як забезпечити прозорість та підзвітність ШІ-систем?

Попри зростаючий інтерес до теми, бракує комплексного методологічного підходу, який би дозволив осмислити роль ШІ як системного фактору трансформації політичної науки. Важливість цього питання зумовлює те, що подальші дослідження повинні спрямовуватися на розробку цілісної методологічної рамки для інтеграції ШІ-технологій у політичну науку з урахуванням як їх потенціалу, так і фундаментальних обмежень.

Формулювання цілей статті (постановка завдання). Враховуючи нові можливості та виклики, які штучний інтелект створює для методології політичних досліджень, **мета статті** полягає в аналізі основних методологічних проблем інтеграції ШІ та формуванні концептуальних основ для розробки гібридних дослідницьких підходів.

Основними завданнями дослідження є:

– систематизувати основні епістемологічні проблеми використання алгоритмів у дослідженні політики, зокрема питання валідації висновків ШІ, "розуміння" політики машинами та інтерпретації результатів у політичному контексті;

- проаналізувати виклики методологічної гібридності – як поєднувати традиційні якісні методи з алгоритмічними підходами та інтегрувати людське судження з машинною логікою;

- з'ясувати проблеми причинності в алгоритмічних дослідженнях політики: чи можуть ШІ-методи встановлювати справжні причинно-наслідкові зв'язки та як відрізнити реальні закономірності від статистичних випадковостей;

- дослідити часові й контекстуальні обмеження алгоритмічного аналізу, враховуючи історичну специфіку та мінливість політичних процесів;

- розглянути етичні аспекти автоматизації політичного аналізу, включаючи нові форми відповідальності та підзвітності в умовах зростаючої ролі експертного знання;

- визначити принципи створення гібридних методологій, що поєднують традиційні політологічні підходи з алгоритмічними методами, включаючи критерії вибору стратегій та етичні стандарти застосування;

- окреслити перспективи того, як розвиток досконаліших ШІ-систем може змінити природу політичного знання.

Виконання означених завдань може сприяти систематизації наявних знань щодо методологічних викликів ШІ у політичній науці та формуванню концептуальних засад для їх вирішення. З огляду на це, дослідження потенційно може мати подвійну цінність: теоретичну та практичну.

Виклад основного матеріалу дослідження

1. Епістемологічні проблеми. Інтеграція штучного інтелекту у політичну науку породжує серйозні епістемологічні виклики, які торкаються самих основ вивчення політичних явищ. Традиційно політична наука спиралася на інтерпретативні методи – дослідники прагнули зрозуміти політичні процеси через контекстуальний аналіз, детальне вивчення конкретних ситуацій та якісне осмислення явищ.

Алгоритмічні системи працюють принципово інакше: аналізують статистичні закономірності у великих масивах даних, виявляють математичні патерни та формують формалізовані правила. Ця різниця створює напругу між двома способами пізнання політичної реальності. Це ставить перед політичною наукою важливе питання: як поєднати глибину розуміння з можливостями масштабного аналізу даних, зберігши наукову строгість дослідження?

Як валідувати алгоритмічні висновки про політичні процеси?

Проблема валідації результатів алгоритмічного аналізу становить одну з найскладніших методологічних дилем сучасної політичної науки. Як перевірити, чи справді те, що "побачив" алгоритм у даних, відповідає реальності? Традиційні критерії достовірності – можливість повторити експеримент, узгодженість висновків різних дослідників, здатність передбачати майбутні події – виявляються недостатніми для оцінки якості висновків машинного навчання у політичному контексті.

У дослідженнях політики з використанням ШІ важливо розрізняти два рівні перевірки результатів. Перший – *технічна валідація*: чи правильно працює алгоритм, чи немає помилок у коді, чи коректно застосовані статистичні методи? Цей рівень перевірки необхідний, але для політичної науки явно недостатній.

Набагато складніше завдання – *семантична валідація*. Тут потрібно з'ясувати, чи дійсно закономірності, які виявив алгоритм, відповідають реальним політичним явищам. Адже алгоритм може знайти статистично значущий паттерн, який насправді не має політичного сенсу.

Розглянемо гіпотетичний сценарій, який демонструє ключову проблему валідації алгоритмів машинного навчання в політичному контексті. Уявімо, що сучасний алгоритм аналізує архівні матеріали українських регіональних ЗМІ часів президентських виборів 2004 року і виявляє сильну кореляцію між програмою Віктора Януковича "Нова взаємодія", обговореннями громадських проєктів та його електоральними перевагами. Статистично цей зв'язок може бути значущим – у регіонах, де активно обговорюються проєкти програми, кандидат отримує більше підтримки.

Але цей гіпотетичний сценарій показує, як технічно коректний алгоритм може прийти до хибних висновків без належного розуміння політичного контексту (тим більше, що результат цих виборів нам відомий). Технічна валідація покаже, що алгоритм працює коректно, статистика правильна, кореляція реальна. Але семантична валідація ставить критичне питання: чи має цей паттерн політичний сенс?

При поверхневому аналізі алгоритм може інтерпретувати цю закономірність як "участь у демократичному плануванні → електоральна підтримка". Однак семантична валідація, яка враховує політичний контекст, розкриває справжню природу зв'язку: це був завуальований підкуп виборців через програму "Нова взаємодія". Громадам обіцяли кошти на проєкти, але з неявним посилом, що ці обіцянки будуть виконані лише у разі обрання В. Януковича президентом.

Алгоритм технічно правильно виявив закономірність між програмою та підтримкою, але без розуміння українського політичного контексту інтерпретував її як легітимний зв'язок "соціальні програми → електоральні преференції", тоді як справжній механізм був "*непрямий підкуп через соціальні обіцянки → голосування*".

Семантична валідація вимагає від дослідника розуміння специфіки української політичної культури. Саме таке контекстуальне розуміння допомагає відрізнити справжні демократичні процеси від замаскованих корупційних схем.

Продовжуючи приклад: алгоритм також міг виявити кореляцію між активністю програми "Нова взаємодія" та офіційними результатами виборів, де В. Янукович нібито переміг у першому турі. Технічна валідація підтвердила б статистичну значущість зв'язку між програмою та "електоральним успіхом". Однак семантична валідація, яка враховує повний політичний контекст, розкриває глибшу проблему: офіційні результати виборів були сфабриковані.

Програма "Нова взаємодія" корелювала не з реальною підтримкою виборців, а з регіонами, де відбувалися найбільші фальсифікації результатів голосування. Згодом виявилось: ті громади, які найактивніше обговорювали проекти програми, насправді голосували не за В. Януковича. Це показує, що справжня громадська активність та критичне мислення виявилися сильнішими за спроби маніпуляцій.

Помаранчева революція та повторне голосування показали справжні електоральні преференції — В. Ющенко переміг з відривом більше 7%. Це означає, що алгоритм, навчений на "офіційних" даних першого туру, не просто неправильно інтерпретував політичні процеси, а й базувався на сфальсифікованих вихідних даних. Тобто без розуміння політичного контексту та критичного аналізу джерел даних ШІ ризикує не просто неправильно інтерпретувати реальність, а й легітимізувати фальсифіковані результати. Семантична валідація тут вимагає не лише розуміння політичних процесів, а й критичного ставлення до достовірності самих даних, особливо в умовах політичних маніпуляцій.

Важливим внеском у розв'язання цієї проблеми стало дослідження Джастіна Гріммер та Брендона Стюарта "Текст як дані: перспективи та ризики автоматичних методів аналізу політичних текстів" [Grimmer & Stewart, 2013]. Їхня робота заклала методологічні основи для того, як поєднувати комп'ютерні методи з традиційними підходами до аналізу політичних текстів, і досі служить базовим посібником для дослідників у цій сфері.

Найскладнішою є *теоретична валідація* — перевірка того, чи узгоджуються висновки алгоритму з існуючими політичними теоріями. Що робити, коли ШІ виявляє паттерни, які суперечать усталеним науковим уявленням? Відкидати як помилки чи переосмислювати теоретичні основи дисципліни?

Проблема ускладнюється специфікою політичних даних — вони часто неповні, упереджені або залежать від контексту. Алгоритми можуть знаходити статистично значущі закономірності без політичного сенсу. Наприклад, зв'язок між активністю у соцмережах та голосуванням може відображати не політичні переконання, а демографічні чи технологічні особливості користувачів. Особливо складно валідувати результати під час роботи з "великими даними", які перевищують можливості людської перевірки. Складно перевірити висновки про політичні настрої мільйонів громадян на основі їхніх цифрових слідів.

Чи можуть машини "розуміти" політику або лише обробляють паттерни? Це питання торкається фундаментальної проблеми природи розуміння та його відмінності від обробки інформації. Сучасні ШІ-системи демонструють вражаючі здібності у розпізнаванні складних паттернів у політичних даних. Проте чи означає це, що вони "розуміють" політику так, як її розуміють люди?

Наше розуміння політики передбачає не лише здатність виявляти закономірності, але й усвідомлення контексту, мотивацій, цінностей та смислів, які лежать в основі політичних дій. Наприклад, політичні поняття —

"демократія", "справедливість", "свобода" – мають не лише дескриптивний, але й нормативний вимір. Вони не просто описують політичну реальність, але й виражають оцінки, цінності та ідеали. Чи може машина адекватно працювати з такими концептами, не розуміючи їх нормативного навантаження?

Методи "контекстуального вбудовування" дозволяють врахувати різні значення політичних термінів залежно від контексту їх вживання. Скажімо, термін "платформа" в політиці може означати політичну програму партії або технічний майданчик для комунікації. Контекстуальне вбудовування допомагає алгоритму моделювати значення цього слова залежно від тексту, в якому воно вживається. "Легітимність" у різних контекстах може мати нюанси – від юридичної законності до соціального визнання влади.

Термін "політичні дослідження" демонструє особливу складність контекстуального розпізнавання, позначаючи академічну дисципліну, конкретні роботи, методологічний підхід або назву видання! Парадоксальність полягає в тому, що аналіз цього терміна в статті для журналу "Політичні дослідження" створює семантичну неоднозначність – алгоритми мають розрізняти галузь знань і саме видання. Це показує, як середовище публікації стає частиною контексту для точної інтерпретації, хоча питання про справжнє машинне "розуміння" політики залишається відкритим.

Тепер щодо проблеми *інтерпретації результатів ШІ у політичному контексті*. Навіть коли алгоритм правильно знаходить закономірності у політичних даних, залишається складне питання: як інтерпретувати ці результати у політично значущому контексті? Проблема не тільки в технічній складності алгоритмів – хоча "чорна скринька" багатьох ШІ-систем справді ускладнює розуміння їхньої роботи. Головна відмінність полягає між алгоритмічною логікою та політичним мисленням.

Алгоритми виявляють паттерни на рівні, недоступному людському сприйняттю – у багатовимірних просторах, нелінійних залежностях, складних взаємодіях між змінними. Але ці паттерни потрібно "перекласти" на мову політичних понять та теорій. Як здійснити такий переклад, зберігши точність і не спотворивши політичний сенс?

Особливо гострою є ця проблема у контексті пояснюваності ШІ-систем. Навіть найскладніші методи пояснення алгоритмічних рішень – LIME, SHAP, градієнтні методи – дають лише наближене уявлення про "логіку" алгоритму. У політиці, де кожне рішення може мати далекосяжні наслідки, така приблизність неприйнятна. Наприклад, у дослідженні виборів 2016 року в США дослідники використали *Twitter* для прогнозування результатів на рівні округів з точністю 81%. Інше дослідження досягло лише 54,8% точності при використанні Naive Bayes класифікатора [Alvi et al., 2023]. Проблема семантичної валідації тут очевидна: алгоритми показують статистично значущі паттерни, але дослідники не можуть пояснити, чому ці паттерни працюють або не працюють. Коли алгоритм виявляє, що певні слова або емоції корелюють з електоральними перевагами, залишається незрозумілим політичний механізм цього зв'язку.

Як зазначається в критичному огляді, "результати залежать від довільних рішень авторів" і навіть додавання однієї партії або одного дня збору даних може кардинально змінити результати [Brito et al., 2024]. Це показовий приклад того, як технічна точність не гарантує політичного розуміння – алгоритм може працювати, але дослідники не розуміють, що саме він "бачить" у даних і чи має це політичний сенс.

2. Методологічна гібридність. Сучасна політична наука стоїть перед необхідністю розробки принципово нових методологічних підходів, які б органічно поєднували традиційні дослідницькі практики з алгоритмічними методами. Цей процес методологічної гібридизації є не просто технічним завданням, а фундаментальною трансформацією дисциплінарних основ політичної науки.

Методологічні виклики ШІ відкривають нові горизонти для дослідження, але водночас вимагають переосмислення дисциплінарних основ. У таблиці 1 ми спробували продемонструвати можливі варіанти застосування штучного інтелекту в політологічних дослідженнях.

Таблиця 1.

Застосування штучного інтелекту в політологічних дослідженнях*

Тематична цірина політології	Можливості ШІ	Форми застосування	Основний вид штучного інтелекту	Обмеження та ризики	Очікувані результати
Електоральні дослідження	Аналіз великих масивів даних про виборців, прогнозування результатів	Аналіз настроїв у соцмережах, обробка опитувань, моделювання поведінки виборців, аналіз демографічних трендів	Machine Learning (машинне навчання)	Складність прогнозування "чорних лебедів", упередженість у даних	Точніші прог- нози виборів, виявлення нових електоральних трендів, пер- соналізовані кампанії, раннє виявлення зсувів у громадській думці
Аналіз політичного дискурсу	Обробка текстів, виявлення паттернів риторики, аналіз мови політиків	Обробка при- родної мови для аналізу промов, виявлення популістської риторики, класифікація політичних позицій, аналіз медіа-контенту	Natural Language Processing (обробка природної мови)	Контекстуальні нюанси, культурні особливості мови	Автоматична класифікація політичних позицій, виявлення маніпулятивних технік, аналіз еволюції політичної риторики, моніторинг дискурсивних змін
Міжнародні відносини	Аналіз дип- ломатичної корес- понденції, прогнозування конфліктів	Системи раннього попередження, аналіз договорів і угод, моделювання геополітичних сценаріїв, обробка новинних потоків	Machine Learning + NLP (машинне навчання + обробка природної мови)	Непередба- чуваність людського фактору, секретність інформації	Раннє попередже- ння про конфлікти, покращений аналіз дипломатичних угод, прогнозу- вання геополі- тичних трендів, оптимізація переговорних стратегій

Державне управління	Оптимізація політичних рішень, аналіз ефективності політик	Оцінка впливу політик, аналіз бюджетних даних, моделювання реформ, оцінка публічних програм	Machine Learning (машинне навчання)	Етичні питання автоматизації рішень, відповідальність	Підвищення ефективності державних програм, оптимізація бюджетного планування, evidence-based policy making, швидша оцінка реформ
Політична участь та активізм	Вивчення онлайн-мобілізації, аналіз соціальних рухів	Аналіз соцмереж, виявлення мереж активістів, дослідження віральності контенту, картування протестних рухів	Network Analysis + Machine Learning (мережевий аналіз + машинне навчання)	Приватність, етика спостереження, фільтрові бульбашки	Розуміння механізмів політичної мобілізації, прогнозування соціальних рухів, аналіз ефективності активістських кампаній, виявлення лідерів думок
Політична економія	Аналіз взаємозв'язків економіки та політики, моделювання впливу	Економетричне моделювання, аналіз лобіювання, дослідження корупції, оцінка економічної політики	Machine Learning (машинне навчання)	Складність каузальних зв'язків, мультиколінеарність	Кращі моделі взаємодії економіки та політики, виявлення корупційних схем, оцінка ефективності економічних реформ, аналіз лобістської діяльності
Порівняльна політологія	Зіставлення політичних систем, виявлення закономірностей	Міжнародний аналіз даних, кластерний аналіз країн, типологізація режимів, картування інститутів	Machine Learning (машинне навчання)	Проблема порівняльності, культурні відмінності	Нові типології політичних режимів, виявлення універсальних закономірностей, покращені порівняльні дослідження, прогнозування політичних трансформацій
Політична теорія	Аналіз текстів класиків, виявлення еволюції ідей	Комп'ютерний аналіз текстів, тематичне моделювання, аналіз цитувань, семантичні мережі понять	Natural Language Processing (обробка природної мови)	Інтерпретативна природа теорії, складність смислів	Картування еволюції політичних ідей, виявлення прихованих зв'язків між теоріями, аналіз впливовості концепцій, цифрові гуманітарні дослідження
Громадська думка	Моніторинг настроїв, виявлення трендів, сегментація аудиторій	Відстеження настроїв у реальному часі, аналіз опитувань, моделювання громадської думки, виявлення маніпуляцій	Machine Learning + NLP (машинне навчання + обробка природної мови)	Репрезентативність онлайн-даних, ехо-камери	Безперервний моніторинг суспільних настроїв, виявлення дезінформаційних кампаній, прогнозування соціальних трендів, персоналізований аналіз аудиторій

Політичні інститути	Аналіз функціонування, ефективності та еволюції інститутів	Інституційний аналіз, аналіз законодавчого процесу, дослідження бюрократії, мережевий аналіз еліт	Network Analysis + Machine Learning (мережевий аналіз + машинне навчання)	Формальні vs неформальні правила, латентні структури	Оптимізація інституційного дизайну, виявлення неефективних процедур, аналіз мереж впливу, прогнозування інституційних змін
Політичні конфлікти	Прогнозування, аналіз причин, моделювання врегулювання	Моделі прогнозування конфліктів, аналіз мови ворожнечі, моделювання ескалації, аналітика миробудування	Machine Learning + NLP (машинне навчання + обробка природної мови)	Етичні дилеми прогнозування, self-fulfilling prophecies	Раннє попередження про конфлікти, кращі стратегії врегулювання, аналіз ефективності миротворчих ініціатив, моделювання post-conflict recovery
Цифрова політика	Вивчення впливу технологій на політику та демократію	Аудит алгоритмів, аналіз цифрових кампаній, дослідження дезінформації, аналіз управління платформами	Machine Learning + NLP + Computer Vision (машинне навчання + обробка природної мови + комп'ютерний зір)	Швидкість технологічних змін, доступ до даних платформ	Розробка етичних стандартів для ШІ, виявлення алгоритмічних упереджень, боротьба з дезінформацією, оптимізація цифрового управління

*Складено автором.

Запропонована модель застосування ШІ в політичній науці становлять один з можливих підходів до систематизації потенційних застосувань різних технологій у тематичних галузях політології, не претендуючи на вичерпність чи остаточність. Найперспективнішими видаються сфери з великими масивами структурованих даних: електоральні дослідження та аналіз громадської думки, де машинне навчання може обробляти дані соцмереж та опитувань. Можливості також убачаються в аналізі політичного дискурсу та політичній теорії через обробку природної мови. Державне управління та політична економія можуть використовувати ШІ для оптимізації рішень та виявлення взаємозв'язків.

Концептуальні виклики очікуються для сфер з високою непередбачуваністю: міжнародні відносини та політичні конфлікти через складність людського фактору. У більшості областей передбачаються спільні проблеми репрезентативності даних, культурних особливостей та етичних питань автоматизації.

Перспективними напрямками розглядається інтеграція різних типів ШІ та розвиток етичних стандартів. Звісно, запропонована схема є робочою моделлю, що потребує емпіричної перевірки та може бути доповнена на основі практичного досвіду.

Одним із найскладніших аспектів методологічної гібридності є *поєднання людського судження з машинною логікою*. Ця проблема виходить за межі технічних питань і торкається фундаментальних епістемологічних питань про природу наукового пізнання.

На що спирається наше судження у політичній науці? На інтуїцію, досвід, теоретичні знання та здатність розуміти контекст. Це дозволяє дослідникам виявляти аномалії, ставити критичні питання, генерувати творчі гіпотези та інтерпретувати результати у ширшому теоретичному контексті. Машинна логіка характеризується послідовністю, масштабністю та здатністю знаходити складні паттерни у великих масивах даних. Ключовий виклик – розробка механізмів, які ефективно поєднують ці різні типи "інтелекту". Один з підходів – створення "human-in-the-loop systems" (систем з участю людини), де дослідник регулярно втручається в алгоритмічний процес, корегуючи його напрямки та інтерпретуючи проміжні результати. Але реалізація таких підходів породжує практичні проблеми: як визначити оптимальні точки людського втручання? Як зрозуміти, що людське судження покращує, а не погіршує якість алгоритмічного аналізу? Як уникнути когнітивних упереджень під час інтерпретації результатів?

Цікавим є дослідження "Інтерпретовані моделі класифікації для прогнозування рецидивів" [Zeng et al., 2017]. Автори досліджують як створити прогностичні моделі рецидивізму, які є точними, прозорими та зрозумілими для прийняття рішень. Стаття стала важливою віхою в дискусії про те, що зрозумілість має бути вбудованою в моделі з самого початку, а не додаватися пізніше через пояснення "чорних скриньок". З погляду політології вона має фундаментальне значення.

По-перше, "human-in-the-loop" системи – це спроба вирішити кризу демократичної легітимності алгоритмічних рішень. Це стосується принципу підзвітності: як забезпечити, щоб автоматизовані рішення залишалися під демократичним контролем? Чи можуть алгоритми адекватно представляти інтереси громадян без людського втручання?

По-друге, приклад з рецидивізмом ілюструє перерозподіл влади між судовою системою як традиційним інститутом, технократами-розробниками алгоритмів та громадянами. Це створює нові форми технополітичного управління, де влада концентрується у тих, хто контролює алгоритми.

У таблиці 2 ми спробували відобразити багатовимірну структуру цієї найскладнішої сучасної проблеми політичної науки – як поєднати інтуїтивне, контекстуальне людське мислення з точністю та масштабованістю машинних алгоритмів у політичних дослідженнях та управлінні.

Таблиця 2.

Поєднання людського судження з машинною логікою в політології*

Порівняльні характеристики

Аспект	Людське судження	Машинна логіка
Основа функціонування	Інтуїція, досвід, теоретичне знання	Алгоритми, статистичний аналіз
Ключові переваги	Контекстуальне розуміння, здатність виявляти аномалії, творчі гіпотези	Консистентність, масштабованість, виявлення складних паттернів
Сфера застосування	Інтерпретація у ширшому контексті, критичні запитання	Обробка великих масивів даних
Обмеження	Когнітивні упередженості, суб'єктивність	Відсутність контекстуального розуміння, "чорна скринька"

Механізми інтеграції

Підхід	Опис	Переваги	Виклики
Human-in-the-loop системи	Регулярне втручання людини в алгоритмічний процес	Поєднання переваг обох типів "інтелекту"	Визначення оптимальних точок втручання
Інтерпретовані моделі	Вбудована прозорість з самого початку	Зрозумілість процесу прийняття рішень	Баланс між точністю та інтерпретованістю

Політологічні виміри проблеми

Напрямок	Ключові питання	Наслідки для політології
Демократична легітимність	Як забезпечити підзвітність алгоритмічних рішень?	Криза легітимності автоматизованих рішень
Перерозподіл влади	Хто контролює алгоритми і як це впливає на владні відносини?	Виникнення технополітичного управління
Гібридизація політики	Як адаптуються традиційні інститути?	Нові форми участі експертів, механізми контролю
Епістемологія політичного знання	Чи можуть алгоритми замінювати політичне судження?	Поєднання статистичного та контекстуального розуміння
Майбутнє демократії	Чи сумісні автоматизовані рішення з демократичними принципами?	Баланс між ефективністю та демократичністю

Практичні проблеми реалізації

Проблема	Опис	Можливі рішення
Точки втручання	Коли людина має втручатися в алгоритмічний процес?	Розробка критеріїв для визначення ключових моментів
Якість аналізу	Як забезпечити покращення, а не погіршення результатів?	Системи валідації людського внеску
Когнітивні упередженості	Як уникнути спотворень під час інтерпретації результатів?	Навчання дослідників, стандартизація процедур
Прозорість процесу	Як зробити рішення зрозумілими для всіх зацікавлених сторін?	Розвиток інтерпретованих моделей

Фундаментальні виклики

Рівень	Виклик	Значення для політології
Епістемологічний	Природа наукового пізнання	Переосмислення методів політичного дослідження
Технічний	Комбінування різних типів інтелекту	Розвиток нових методологічних підходів
Політичний	ШІ як новий актор владних відносин	Трансформація політичного порядку
Етичний	Збереження людського фактору в політиці	Переосмислення демократичних принципів

*Складено автором.

"Human-in-the-loop" підходи відображають гібридизацію політики: традиційні політичні інститути адаптуються до технологічних змін, з'являються нові форми участі експертів у прийнятті рішень та формуються нові механізми стримувань і противаг. Це також торкається епістемології політичного знання – питань про те, як поєднати статистичне та контекстуальне розуміння політичних явищ, чи можуть алгоритми замінювати політичне судження, як зберегти герменевтичну традицію політичної науки. Розвиток таких систем ставить екзистенційні питання для майбутнього демократії: чи сумісні автоматизовані рішення з демократичними принципами? Як зберегти людський фактор у політичному процесі? Яким буде баланс між ефективністю та демократичністю?

Тобто це не просто технічна проблема, а фундаментальна трансформація політичного порядку, де ШІ стає новим актором у владних відносинах. Ефективне використання ШІ-технологій вимагає не лише технічних знань, але й розуміння специфіки політичних явищ, методологічних принципів політичної науки та етичних аспектів дослідження суспільних процесів.

Візьмемо на себе сміливість зауважити, що це вимагає формування нового типу дослідника — "цифрового політолога", який поєднає глибоке розуміння політичних процесів з технічними навичками роботи з ШІ-системами. Підготовка таких спеціалістів потребує кардинальної трансформації освітніх програм з політичної науки, введення курсів з машинного навчання, статистики та програмування.

3. Проблема каузальності. Питання каузальності є одним із найфундаментальніших та складних методологічних викликів під час застосування штучного інтелекту у політичній науці. *ШІ виявляє кореляції, але чи може встановлювати причинно-наслідкові зв'язки у політиці?*

Більшість сучасних ШІ-методів, особливо машинного навчання, орієнтовані на виявлення статистичних закономірностей та паттернів у даних. Вони можуть з високою точністю передбачати певні політичні явища – електоральні результати чи громадські настрої, – але це не означає, що вони розуміють причинні механізми, що лежать в основі цих явищ. Алгоритм може виявити, що зростання активності в соцмережах корелює зі зміною електоральних преференцій, але це заважає дійти висновку про те, що перше є причиною другого. Можливо, обидва явища є наслідками третього фактору, або зв'язок між ними є випадковим [Pearl & Mackenzie, 2018]. У політичному контексті ця проблема набуває особливої гостроти, оскільки політичні процеси характеризуються складними взаємодіями між множинними факторами, наявністю зворотних зв'язків та можливістю взаємної каузації. Традиційні підходи до каузального аналізу – рандомізовані експерименти чи природні експерименти – часто неможливі або неетичні у політичному контексті. Дослідники намагаються адаптувати методи каузального аналізу для машинного навчання, розробляючи алгоритми, які можуть виявляти каузальні зв'язки у даних. Методи каузального виведення дозволяють оцінити ефект певних "утручань" на основі спостережних даних [Hernán & Robins, 2020].

Проте ці підходи потребують сильних припущень про структуру каузальних зв'язків, які часто важко обґрунтувати у політичному контексті.

Як відрізнити справжні закономірності від статистичних артефактів?

Ця проблема особливо актуальна під час роботи з великими масивами політичних даних, де алгоритми можуть виявити тисячі статистично значущих кореляцій для аналізу електоральної поведінки, соцмедіа або опитувань, більшість з яких є випадковими. Феномен множинного тестування означає, що під час перевірки великої кількості гіпотез деякі неминуче виявляться статистично значущими навіть за відсутності справжніх зв'язків.

У політичній науці проблема ускладнюється тим, що політичні дані часто мають складну багаторівневу структуру – ієрархічну, темпоральну або мережеву. Стандартні статистичні процедури можуть не враховувати цю складність, що призводить до хибних висновків про ефективність кампаній, вплив реформ або причини електоральних зсувів. Одним з підходів є використання методів крос-валідації та тестування на незалежних наборах даних. Проте у політичному контексті отримання справді незалежних даних проблематичне, адже політичні процеси в різних країнах або періодах мають унікальні інституційні та культурні особливості, що ускладнює генералізацію результатів.

4. Темпоральність та контекстуальність. Політичні процеси мають історичну специфіку. Кожна політична ситуація є унікальною – вона формується специфічним поєднанням історичних, культурних, економічних та соціальних факторів. Те, що працювало в одному політичному контексті, може виявитися неефективним в іншому. *Алгоритми навчаються на минулих даних, але політика постійно змінюється.* Це одна з найфундаментальніших проблем застосування ШІ у політичній науці. Політичні процеси характеризуються здатністю до якісних змін та трансформації існуючих закономірностей. Спроба передбачення українських парламентських виборів 2019 року на основі даних попередніх виборів це ілюструє. Алгоритм, навчений на даних 2012 і 2014 років, міг виявити стійкі паттерни: високу підтримку "Партії регіонів" на Сході, стабільні електоральні преференції регіонів, кореляцію між економічними показниками та голосуванням.

Але події 2013-Револуція Гідності, окупація Криму, початок війни на Донбасі – кардинально змінили політичний ландшафт. "Партія регіонів" припинила існування, з'явилися нові сили, змінилася електоральна географія. Алгоритм не зміг би передбачити появу "Слуги народу" або кардинальну зміну уподобань у південних та східних регіонах.

Яскравіший приклад – спроба передбачення виборів 2024 року на основі довоєнних даних. Повномасштабна війна настільки трансформувала українське суспільство, що алгоритми, навчені на мирному часі, стали непридатними. З'явилися нові критерії оцінки політиків, змінилися пріоритети виборців, частина електорату фізично недоступна. Це демонструє фундаментальну проблему: машинне навчання припускає стабільність паттернів, тоді як політика – це сфера якісних стрибків та непередбачуваних трансформацій.

Проблема "дрейфу концепцій" у політичному аналізі. "Дрейф концепцій" (*concept drift*) – це явище, коли статистичні властивості цільової змінної змінюються з часом непередбачуваним чином. У політичному контексті це може означати, що саме значення політичних понять та категорій змінюється з часом. Наприклад, поняття "лівих" та "правих" політичних поглядів може мати різні значення у різні історичні періоди. Те, що вважалося "лівими" поглядами у 1970-х роках, може сприйматися як "центристське" або навіть "праві" позиції сьогодні. Для подолання цієї проблеми дослідники розробляють методи адаптивного навчання, які дозволяють алгоритмам оновлювати свої моделі у відповідь на зміни у даних. Проте це потребує постійного моніторингу якості моделей та регулярного перенавчання, що може бути ресурсозатратним.

5. Етичні та нормативні аспекти

Хто відповідний за алгоритмічні рекомендації у політиці?

Коли алгоритмічна система генерує рекомендації щодо політичних рішень, хто має відповідати за наслідки цих рекомендацій? Традиційна модель наукової відповідальності передбачає, що дослідник відповідає за коректність методів та чесність у представленні результатів, але не за те, як ці результати будуть використані. Проте у випадку з ШІ-системами, які можуть безпосередньо впливати на політичні рішення, таке розмежування стає проблематичним. З одного боку, розробники алгоритмів можуть стверджувати, що вони лише створюють технічні інструменти, а відповідальність за їх використання лежить на політичних акторах. З іншого боку, політики можуть посилатися на "об'єктивність" алгоритмічних рекомендацій, перекладаючи відповідальність на технічних експертів. Ця проблема ускладнюється тим, що ШІ-системи часто працюють як "чорні скриньки", що робить важким визначення того, як було прийнято те чи інше рішення. Якщо алгоритм рекомендує певну політику, яка згодом виявляється помилковою або шкідливою, як визначити, чи була проблема у даних, алгоритмі, інтерпретації результатів чи у самому політичному рішенні?

Як зберегти демократичні принципи у разі автоматизації політичного аналізу?

Демократія базується на принципах прозорості, підзвітності, участі громадян у прийнятті рішень та можливості оскарження владних дій. Автоматизація політичного аналізу може загрожувати цим принципам кількома способами.

По-перше, непрозорість алгоритмічних систем може підривати принцип відкритості демократичного управління. Якщо політичні рішення базуються на алгоритмічних рекомендаціях, логіка яких є незрозумілою навіть для експертів, громадяни втрачають можливість зрозуміти та критично оцінити ці рішення.

По-друге, алгоритмічні системи можуть посилювати існуючі нерівності та упередженості. Якщо дані, на яких навчається алгоритм, містять історичні упередженості проти певних груп населення, система може відтворити та інституціоналізувати ці упередженості. Це особливо проблематично у контексті розподілу ресурсів чи надання публічних послуг.

По-третє, автоматизація може знижувати рівень людської участі у політичних процесах. Це може призводити до технократизації політики та послаблення демократичної легітимності.

Крім того, існує ризик маніпулювання громадською думкою через алгоритмічні системи. Якщо політичні актори мають доступ до потужних ШІ-інструментів для аналізу та впливу на електоральну поведінку, це може створювати нечесні переваги та підривати принцип рівності політичної конкуренції.

Для подолання цих проблем необхідно розробити етичні стандарти використання ШІ у політичній сфері, які б забезпечували прозорість алгоритмічних рішень, справедливість їх наслідків та збереження демократичних цінностей.

Висновки і подальші перспективи дослідження. Інтеграція ШІ-технологій у політичну науку не є простим технічним удосконаленням існуючих методів, а становить фундаментальний виклик для епістемологічних та методологічних засад дисципліни. Алгоритмічні підходи створюють нову реальність наукового пізнання політичних процесів, яка потребує переосмислення традиційних критеріїв валідності, процедур верифікації та стандартів наукової строгості.

Епістемологічні проблеми ШІ у політичній науці не можуть бути розв'язані виключно технічними засобами. Питання валідації алгоритмічних висновків, "розуміння" політики машинами та інтерпретації результатів у політичному контексті потребують філософського осмислення та розробки нових концептуальних рамок.

Методологічна гібридність є новою нормою політичних досліджень у цифрову епоху. Ефективна інтеграція ШІ вимагає не заміщення традиційних методів алгоритмічними підходами, а їх творчого поєднання. Це передбачає розвиток нового типу дослідницької компетентності, яка поєднує глибоке розуміння політичних процесів з технічними навичками роботи з ШІ-системами.

Проблема каузальності залишається однією з найскладніших методологічних дилем. Здатність алгоритмів виявляти кореляції не означає їхньої спроможності встановлювати справжні причинно-наслідкові зв'язки. Розвиток методів каузального виведення для машинного навчання є перспективним напрямом, проте потребує сильних теоретичних припущень, які часто важко обґрунтувати у політичному контексті.

На особливу увагу заслуговують темпоральні та контекстуальні обмеження алгоритмічного аналізу. Політичні процеси мають принципово історичний характер, що створює фундаментальні проблеми для ШІ-методів, які базуються на припущенні про стаціонарність процесів.

Критично важливі етичні та нормативні аспекти використання ШІ у політичній науці. Питання відповідальності за алгоритмічні рекомендації, збереження демократичних принципів під час автоматизації політичного аналізу та запобігання алгоритмічній упередженості потребують розробки комплексних етичних стандартів.

Пріоритетними напрямками подальших досліджень мають стати:

По-перше, розробка нових валідаційних процедур, адаптованих до потреб політичної науки – критеріїв оцінки політичної релевантності алгоритмічних результатів, методів верифікації теоретичної обґрунтованості ШІ-висновків.

По-друге, дослідження принципів побудови ефективних гібридних методологій. Особлива увага – вивченню оптимальних способів поєднання людського судження з машинним аналізом, розробці "людино-машинних петель" для політичних досліджень.

По-третє, емпіричне дослідження ефективності різних ШІ-методів у конкретних галузях політичної науки. Це передбачає систематичне порівняння алгоритмічних та традиційних підходів, ідентифікацію галузей, де людська експертиза залишається незамінною.

По-четверте, розвиток методів каузального аналізу, адаптованих до специфіки політичних даних, особливо у контексті складних політичних систем з множинними взаємодіючими факторами.

По-п'яте, вивчення темпоральних аспектів алгоритмічного аналізу – розробка методів адаптивного навчання, дослідження концептуального дрейфу у політичних даних та створення підходів для врахування історичної специфіки.

По-шосте, дослідження етичних та соціальних наслідків використання ШІ, особливо впливу алгоритмізації на демократичні процеси та розробки механізмів запобігання алгоритмічним упередженням.

Бібліографічні посилання

- Бут, С. (2025). Штучний інтелект у політичній діяльності: основні напрямки використання. *Публічне управління і політика*, 5(9), 1–9.
- Климчук, І. (2025). Роль штучного інтелекту у сучасній дипломатичній практиці. *Наукові праці Міжрегіональної Академії управління персоналом. Політичні науки*, (1), 12–25.
- Костенко, О. В. (2022). Аналіз національних стратегій розвитку штучного інтелекту. *Інформація і право*, 2(41), 58–69.
- Радіонова, І. (2023). Політика у сфері штучного інтелекту: безпековий вимір. *Науково-теоретичний альманах "Грані"*, 26(1), 50–56.
- Сабов, І. (2018). Політичний аспект використання штучного інтелекту. *Міжнародні відносини, суспільні комунікації та регіональні студії*, 1, 15–22.
- Янишівський, М. М. (2024). Штучний інтелект як загальноцільова технологія: виклики та підходи до публічної політики. *Проблеми сучасних трансформацій*. Серія право публічне управління та адміністрування, (14).
- Alvi, Q., Ali, S. F., Ahmed, S. B., Khan, N. A., Javed, M., & Nobanee, H. (2023). On the frontiers of Twitter data and sentiment analysis in election prediction: a review. *PeerJ Computer Science*, 9, e1517. <https://doi.org/10.7717/peerj-cs.1517>
- Beduschi, A., & McAuliffe, M. (2022). Chapter 11: Artificial Intelligence, Migration and Mobility: Implications for Policy and Practice. *World Migration Report 2022*.

- Brito, K., Silva Filho, R. L. C., & Adeodato, P. J. L. (2024). Stop trying to predict elections only with twitter – There are other data sources and technical issues to be improved. *Government Information Quarterly*, 41(1), 101899.
- Caruso, M., & Spadaro, A. (2024). Digital humanities and artificial intelligence: An accelerationist perspective of the future. *Proceedings*, 96(1), 10.
- Chapinal-Heras, D., Díaz-Sánchez, C. (2024). A review of AI applications in human sciences research. *Digital Applications in Archaeology and Cultural Heritage*, 32, e00323. <https://doi.org/10.1016/j.daach.2024.e00323>.
- Dedema, M., & Ma, R. (2024). The collective use and perceptions of generative AI tools in digital humanities research: Survey-based results. *arXiv*.
- Dignum, V. (2019). *Responsible artificial intelligence: How to develop and use AI in a responsible way*. Springer.
- Evans, J. A. (2008). Electronic publication and the narrowing of science and scholarship. *Science*, 321(5887), 395–399.
- Goldfarb, A. (2024). Pause artificial intelligence research? Understanding AI policy challenges. *Canadian Journal of Economics*, 57(2), 363–377.
- Grimmer, J., & Stewart, B. M. (2013). Text as data: The promise and pitfalls of automatic content analysis methods for political texts. *Political Analysis*, 21(3), 267–297.
- Hernán, M. A., & Robins, J. M. (2020). *Causal inference: What if*. Chapman & Hall/CRC.
- Huang, Z. (2022). Introducing Neuro-Symbolic Artificial Intelligence to Humanities and Social Sciences: Why Is It Possible and What Can Be Done? *TEM Journal*, 11(4), 1863-1870.
- Lemke, N., Trein, P., & Varone, F. (2024). Defining artificial intelligence as a policy problem: A discourse network analysis from Germany. *European Policy Analysis*, 10(2), 162–187.
- Lin, Z. (2024). Towards an AI policy framework in scholarly publishing. *Trends in Cognitive Sciences*, 28(2), 85–88.
- Liu, J., Wang, Z., Xie, J., & Pei, L. (2024). From ChatGPT, DALL-E 3 to Sora: How has Generative AI changed digital humanities research and services? *arXiv*.
- Mathkunti, N. M., Ananthanagu, U., & Menon, S. (2023). A review of AI's influence on humanities and social sciences: neuro-symbolic AI and transformative potential. *Tuijin Jishu / Journal of Propulsion Technology*, 44(5).
- O'Neil, C. (2016). *Weapons of math destruction: How big data increases inequality and threatens democracy*. Crown Publishers.
- Pearl, J., & Mackenzie, D. (2018). *The book of why: The new science of cause and effect*. Basic Books.
- Russell, S., & Norvig, P. (2020). *Artificial intelligence: A modern approach* (4th ed.). Pearson.
- Ulnicane, I., & Erkkilä, T. (2023). Politics and policy of Artificial Intelligence. *Review of Policy Research*, 40(5), 612–625.
- Walter, Y. (2024). Managing the race to the moon: Global policy and governance in Artificial Intelligence regulation – A contemporary overview and an analysis of socioeconomic consequences. *Discover Artificial Intelligence*, 4(1), 1–24.
- Watts, D. J. (2011). *Everything is obvious: How common sense fails us*. Crown Business.

- Zeng, J., Ustun, B., & Rudin, C. (2017). Interpretable classification models for recidivism prediction. *Journal of the Royal Statistical Society*, 180(3), 689–722.
- Zhao, H. (2023). Challenges or opportunities: When artificial intelligence is applied to digital humanities. *Tidskrift för ABM*, 8(1), 57–65.

References

- Alvi, Q., Ali, S. F., Ahmed, S. B., Khan, N. A., Javed, M., & Nobanee, H. (2023). On the frontiers of Twitter data and sentiment analysis in election prediction: a review. *PeerJ Computer Science*, 9, e1517. <https://doi.org/10.7717/peerj-cs.1517>
- Beduschi, A., & McAuliffe, M. (2022). Chapter 11: Artificial Intelligence, Migration and Mobility: Implications for Policy and Practice. *World Migration Report 2022*.
- Brito, K., Silva Filho, R. L. C., & Adeodato, P. J. L. (2024). Stop trying to predict elections only with twitter – There are other data sources and technical issues to be improved. *Government Information Quarterly*, 41(1), 101899.
- But, S. (2025). Shtuchnyi intelekt u politychnii diialnosti: osnovni napriamky vykorystannia [Artificial Intelligence in Political Activity: Main Areas Of Use]. *Publichne upravlinnia i polityka*, 5(9), 1–9 [in Ukrainian].
- Caruso, M., & Spadaro, A. (2024). Digital humanities and artificial intelligence: An accelerationist perspective of the future. *Proceedings*, 96(1), 10.
- Chapinal-Heras, D., Díaz-Sánchez, C. (2024). A review of AI applications in human sciences research. *Digital Applications in Archaeology and Cultural Heritage*, 32, e00323. <https://doi.org/10.1016/j.daach.2024.e00323>.
- Dedema, M., & Ma, R. (2024). The collective use and perceptions of generative AI tools in digital humanities research: Survey-based results. *arXiv*.
- Dignum, V. (2019). *Responsible artificial intelligence: How to develop and use AI in a responsible way*. Springer.
- Evans, J. A. (2008). Electronic publication and the narrowing of science and scholarship. *Science*, 321(5887), 395–399.
- Goldfarb, A. (2024). Pause artificial intelligence research? Understanding AI policy challenges. *Canadian Journal of Economics*, 57(2), 363–377.
- Grimmer, J., & Stewart, B. M. (2013). Text as data: The promise and pitfalls of automatic content analysis methods for political texts. *Political Analysis*, 21(3), 267–297.
- Hernán, M. A., & Robins, J. M. (2020). *Causal inference: What if*. Chapman & Hall/CRC.
- Huang, Z. (2022). Introducing Neuro-Symbolic Artificial Intelligence to Humanities and Social Sciences: Why Is It Possible and What Can Be Done? *TEM Journal*, 11(4), 1863-1870.
- Klymchuk, I. (2025). Rol shtuchnoho intelektu u suchasnoi dyplomatychnii praktytsi [The role of Artificial Intelligence in Modern diplomatic practice]. *Naukovi pratsi Mizhrehionalnoi Akademii upravlinnia personalom. Politychni nauky*, (1), 12–25 [in Ukrainian].
- Kostenko, O. V. (2022). Analiz natsionalnykh stratehii rozvytku shtuchnoho intelektu [Analysis of national artificial intelligence development strategies]. *Informatsiia i pravo*, 2(41), 58–69 [in Ukrainian].

- Lemke, N., Trein, P., & Varone, F. (2024). Defining artificial intelligence as a policy problem: A discourse network analysis from Germany. *European Policy Analysis*, 10(2), 162–187.
- Lin, Z. (2024). Towards an AI policy framework in scholarly publishing. *Trends in Cognitive Sciences*, 28(2), 85–88.
- Liu, J., Wang, Z., Xie, J., & Pei, L. (2024). From ChatGPT, DALL-E 3 to Sora: How has Generative AI changed digital humanities research and services? *arXiv*.
- Mathkunti, N. M., Ananthanagu, U., & Menon, S. (2023). A review of AI's influence on humanities and social sciences: neuro-symbolic AI and transformative potential. *Tuijin Jishu / Journal of Propulsion Technology*, 44(5).
- O'Neil, C. (2016). *Weapons of math destruction: How big data increases inequality and threatens democracy*. Crown Publishers.
- Pearl, J., & Mackenzie, D. (2018). *The book of why: The new science of cause and effect*. Basic Books.
- Radionova, I. (2023). Polityka u sferi shtuchnoho intelektu: bezpekovyi vymir [Artificial Intelligence Policy: the Security Dimension]. *Naukovo-teoretychnyi almanakh "Hrani"*, 26(1), 50–56 [in Ukrainian].
- Russell, S., & Norvig, P. (2020). *Artificial intelligence: A modern approach* (4th ed.). Pearson.
- Sabov, I. (2018). Politychnyi aspekt vykorystannia shtuchnoho intelektu [Political Aspect of the Use of Artificial Intelligence]. *Mizhnarodni vidnosyny, suspilni komunikatsii ta rehionalni studii*, 1, 15–22 [in Ukrainian].
- Ulnicane, I., & Erkkilä, T. (2023). Politics and policy of Artificial Intelligence. *Review of Policy Research*, 40(5), 612–625.
- Walter, Y. (2024). Managing the race to the moon: Global policy and governance in Artificial Intelligence regulation – A contemporary overview and an analysis of socioeconomic consequences. *Discover Artificial Intelligence*, 4(1), 1–24.
- Watts, D. J. (2011). *Everything is obvious: How common sense fails us*. Crown Business.
- Yanyshivskyi, M. M. (2024). Shtuchnyi intelekt yak zahalnotsilova tekhnolohiia: vyklyky ta pidkhody do publichnoi polityky [Artificial Intelligence as a General-purpose technology: Challenges and Policy approach]. *Problemy suchasnykh transformatsii. Seriia pravo publichne upravlinnia ta administruvannia*, (14) [in Ukrainian].
- Zeng, J., Ustun, B., & Rudin, C. (2017). Interpretable classification models for recidivism prediction. *Journal of the Royal Statistical Society*, 180(3), 689–722.
- Zhao, H. (2023). Challenges or opportunities: When artificial intelligence is applied to digital humanities. *Tidskrift för ABM*, 8(1), 57–65.