

ПРОБЛЕМИ ВИКОРИСТАННЯ ШТУЧНОГО ІНТЕЛЕКТУ
ДЛЯ ІНФОРМАЦІЙНОЇ БЕЗПЕКИ В СФЕРІ ОСВІТНЬОЇ ДІЯЛЬНОСТІUSAGE ISSUES OF ARTIFICIAL INTELLIGENCE
FOR INFORMATION SECURITY PURPOSES IN EDUCATIONAL FIELD

В статті висвітлено проблеми стрімкого розвитку штучного інтелекту, що стрімко змінює науковий та освітній простір. Підкреслено спровоковані дискусії щодо впровадження великих мовних моделей. Досліджена доступність технологій на базі штучного інтелекту. Вказано на проблематику появи хакерських угруповань, що адаптують штучний інтелект для зламу різних систем. Поставлено етичне питання щодо законності використання штучного інтелекту для атаки на зловмисників. Виокремлено питання академічної доброчесності та висвітлено думки щодо поширення псевдонаукових матеріалів, а також аналіз спроб створення алгоритмів для визначення згенерованих штучним інтелектом текстів. Виявлено триваючі філософські дискусії навколо штучного інтелекту, що стосуються його потенційної автономії та здатності змінювати суспільство в майбутньому. Підкреслюються етичні проблеми, пов'язані з упередженістю алгоритмів, і необхідність прозорості рішень, керованих штучним інтелектом. Проведено теоретичний дослід надання штучному інтелекту доступу до величезних масивів інформації та неконтрольованого розвитку. Підсумовано можливі напрямки подальшого дослідження та використання людством штучного інтелекту.

Ключові слова: Освіта, кібербезпека, захист даних, штучний інтелект, лінгвістичні моделі, етика, задачі з невизначеними початковими умовами.

The article highlights the problems of the rapid development of artificial intelligence, which is rapidly changing the scientific and educational space, causing excitement and concern. The provoked discussions on the implementation of large language models (LLM), such as GPT from OpenAI, their use in teaching, scientific activities and everyday life are highlighted. The accessibility of technologies based on artificial intelligence is investigated, which allows anyone to create texts, images, music, but also gives rise to ethical and legal issues. In the field of cyber security, artificial intelligence has been highlighted as playing a controversial role. The issue of the emergence of hacker groups that adapt artificial intelligence to break into various systems, which complicates the fight against cybercrime, is pointed out. The ethical question is raised regarding the legality of using artificial intelligence to attack back the attackers. The issue of academic integrity, which is currently suffering from the use of AI-generated texts, is raised and highlighted. Opinions on the spread of pseudoscientific materials are highlighted, as well as an analysis of attempts of creation of algorithms to identify AI-generated texts. The role of artificial intelligence in the formation of future educational methodologies and its impact on traditional pedagogical values and approaches are considered. Ongoing philosophical discussions around AI are identified, relating to its potential autonomy and ability to change society in the future. Ethical issues related to algorithmic bias and the need for transparency in AI-driven decisions are emphasized. The potential risks of manipulating AI for disinformation campaigns and its role in shaping public discourse are also considered. A theoretical study of providing AI with access to vast amounts of information and uncontrolled development, which may lead to undesirable consequences, is conducted. As artificial intelligence continues to develop under the power of humanity, the need for interdisciplinary research to address the multifaceted challenges it creates is emphasized, calling for collaboration between different research areas, such as technological and ethical, to establish guiding principles for the responsible use of artificial intelligence. Possible directions for further research and use of artificial intelligence by humanity are summarized.

Key words: Education, cybersecurity, data protection, artificial intelligence, linguistic models, ethics, problems with uncertain initial conditions.

УДК 004.65.056:330.3

DOI: <https://doi.org/10.32782/dees.16-62>

Синицький Р.К.¹

аспірант кафедри
інформаційних систем в економіці,
асистент кафедри інформаційних
систем в економіці,
Київський національний економічний
університет імені Вадима Гетьмана

Synyskyi Rostyslav

Kyiv National Economic University
named after Vadym Hetman

Постановка проблеми. Штучний інтелект займає сьогодні перші рядки в новинах та обговореннях. Наразі, поява найновіших моделей штучного інтелекту та великих мовних моделей сколихнуло як наукову так і ненаукову спільноту. Технології, що стоять за цими моделями, дійсно здатні перевернути наше розуміння того, як функціонує інформаційне суспільство. Вони відкривають нові горизонти можливостей, але разом із цим виникають і серйозні етичні, соціальні та технічні питання.

Показовим є той факт, що на момент появи, використання LLM-моделей ніяк не обмежено

ні сферами, ні людськими можливостями, тобто будь-хто міг попросити штучний інтелект написати привітання зі святом, просту роботу, відповідь на лист. Одним із ключових аспектів є те, що на момент появи таких технологій їх використання ніяк не обмежено ні сферами, ні людськими можливостями ні жодними значимими регуляціями чи правовими обмеженнями. Це дозволяє будь-кому, хто має доступ до таких систем, використовувати їх для виконання найрізноманітніших завдань, від повсякденного написання привітань і простих письмових робіт до створення складних алгоритмів або навіть шкідливих програм. Незважаючи на

¹ ORCID: <https://orcid.org/0009-0002-6271-3041>

те, що такі можливості можуть здатися корисними і зручними, водночас створюють величезні ризики для безпеки, етики та законності. Коли мова йде про використання алгоритмів штучного інтелекту для написання привітань або відповіді на електронні листи, пересічній людині це здається нормальним і навіть продуктивним. Проте, коли такі технології використовуються для більш серйозних завдань, зокрема для створення листів для скаму, написання поліморфних вірусів або іншого шкідливого програмного забезпечення, ситуація кардинально змінюється. Легкість, з якою можна створити підроблену інформацію чи маніпулювати фактами, викликає серйозні побоювання щодо потенційних наслідків як для пересічних людей так і для інформаційної безпеки в цілому.

Наразі вже відомо, що зловмисники часто використовують алгоритми штучного інтелекту для генерації переконливих фішингових листів, підробки даних, які важко відрізнити від справжніх повідомлень, що, в свою чергу, може призвести до масштабних фінансових втрат або витоку конфіденційної інформації з будь-яких організацій.

Додатково, з огляду на значну популярність таких технологій, не можна оминати й питання псевдонаукових або навіть шкідливих наукових робіт, які можуть бути згенеровані з використанням таких технологій. Часом важко відрізнити реальні наукові результати від фальшивих що створює загрозу маніпулювання інформацією. Така ситуація може поставити під загрозу не лише репутацію всіх наукових установ, а й суспільну довіру до наукових досліджень в світі. Результати, що походять від таких інструментів, можуть виглядати переконливо, але насправді бути безглуздими, хибними чи вести до неправильних висновків, що може завдати шкоди як у навчанні, медицині, економіці так і в інших важливих сферах.

У цьому контексті важливо зауважити, що використання штучного інтелекту в таких випадках піднімає велику кількість питань. Наприклад, чи варто дозволяти впроваджувати відкритий доступ до таких моделей? Чи повинні бути розроблені якісь обмеження, що стосуються їх застосування, аби уникнути небажаних наслідків? Чи можливо використання подібних технологій лише на благо суспільства, і як прибрати фактор нанесення шкоди? Ці та інші питання потребують більш серйозних дискусій серед наукової та правової спільноти, а також потребує розробки етичних та технічних стандартів для врегулювання взаємодії зі штучним інтелектом, адже у різних країнах є різні погляди на те, як можна використовувати подібні технології, і відповідно – різний рівень обмежень і врегулювання. Це створює певні труднощі в глобальному контексті, коли одні держави намагаються використовувати великі мовні моделі для своєї власної вигоди, а інші зосереджуються

на безпеці та етиці. Тому міжнародна співпраця та вироблення загальних норм і стандартів стають на перше місце для майбутнього розвитку штучного інтелекту як інструменту для людства.

Аналіз останніх досліджень і публікацій. Останні дослідження та публікації у сфері штучного інтелекту (ШІ) демонструють значний прогрес у впровадженні цих технологій в освітній процес. Любомира Ілійчук у статті «Штучний інтелект і якість освіти: можливості, виклики та загрози» аналізує потенціал штучного інтелекту у персоналізації навчання, створенні адаптивного освітнього середовища та автоматизованому оцінюванні. Вона також підкреслює виклики, такі як стандартизація освіти, ризики для академічної доброчесності та кібербезпеки [1].

Публікація Тетяни Гodeцької [2] «Штучний інтелект як ефективна технологія інформаційно-аналітичного супроводу цифрової трансформації освіти» аналізує використання штучного інтелекту для візуалізації складних даних, що робить їх більш доступними для аналізу. Автор підкреслює важливість глобального співробітництва для максимального використання переваг штучного інтелекту.

Інша робота, «Аналіз ролі штучного інтелекту у впровадженні диференційованого підходу до навчання», досліджує можливості штучного інтелекту в побудові адаптивних моделей навчання [3]. Автори статті відзначають, що впровадження штучного інтелекту покращує персоналізацію навчання, забезпечує зворотний зв'язок у реальному часі та враховує індивідуальні потреби студентів. Однак вони також звертають увагу на такі проблеми, як упередженість алгоритмів, конфіденційність даних та високі витрати на впровадження технологій.

Стаття Ali Borji [4] «A Categorical Archive of ChatGPT Failures» аналізує різні категорії помилок, які можуть виникати при використанні великих мовних моделей, таких як ChatGPT. Автор досліджує проблеми з міркуванням, фактичними помилками, математичними обчисленнями, написанням коду та упередженістю, підкреслюючи ризики та обмеження використання таких моделей.

Постановка завдання. Мета дослідження – аналіз впливу штучного інтелекту на наукову сферу, кібербезпеку та етичні аспекти його використання в професійній та щоденній діяльності. Основний фокус спрямований на оцінку загроз та можливостей, що виникають у зв'язку з розвитком генеративних моделей штучного інтелекту, зокрема великих мовних моделей, а також їх впливу на освіту, авторське право, хакерську діяльність та академічну доброчесність.

Дослідження також розглядає дві крайності у використанні штучного інтелекту, першу – як потужного інструменту для наукових відкриттів,

досліджень та професійної наукової діяльності, і другу – як потенційної загрози через можливе майбутнє неконтрольоване навчання моделей та можливість зловмисного застосування як в сьогоденні так і в майбутньому. У роботі порушуються питання недостатності врегулювання, появи моральних, етичних та правових наслідків автоматизованого хакінгу, а також присутній аналіз перспектив подальшого розвитку таких моделей у контексті технологічного прогресу та можливих загроз для людства.

Виклад основного матеріалу дослідження.

Використання моделей генерації зображень на основі штучного інтелекту теж сколихнули людство, але не так сильно як великих мовних моделей, адже можливості останньої куди ширші і цікавіші. Попри те, відомі всім Stable Diffusion та Midjourney теж зробили свій вклад в появу певних побоювань з усіх боків усіх спільнот. Тут варто зауважити, що використання генерованих зображень призвело до доволі передбачуваних наслідків, як от страйк графічних дизайнерів і діджитал-художників, але, наприклад, позитивно стороною стала можливість побудови простого наочно-графічного матеріалу для науковців [5].

Зрештою розвиток штучного інтелекту дозволяє писати музику, твори, розповіді, вірші, конструювати складні наукові і псевдонаукові тези і статті, і через такі можливості особливо наукова спільнота зацікавлена в обмеженні використання GPT-моделей в створенні наукових та залікових робіт студентів [5].

Звичайно для викладачів існують і свої поради від колег і однодумців як впоратись із появою штучного інтелекту і викладати в його епоху [6]. Звісно, це не відмінняє того факту, що академічна доброчесність фактично є останнім фактором стримування від використання штучного інтелекту за для написання різного роду робіт, плагіату або ухилення від перевірок [7; 8].

Зрозуміло що наразі виявити такі роботи доволі складно, що підтверджує факт того, що програмний засіб GPTZero, який аналізує тексти і виявляє написані за допомогою штучного інтелекту роботи, проаналізувавши конституцію Сполучених Штатів Америки виявив, що її частина написана штучним інтелектом, що в принципі неможливо. Сферу кібербезпеки зміни теж не оминули стороною. Ідеї, що нам приносить сьогодення несуть за собою питання легітимності та етичності їх використання, що залежить від людини і фактичного кінцевого користувача що ці ідеї втілюють.

Етична сторона багатьох питань фактично і досі не вирішена адже якщо казати саме про штучний інтелект як про не етичну частину інформаційної безпеки то в загальному маємо ситуацію при якій отримати доступ до даних хакера за допомогою штучного інтелекту так само вважається не

етичною дією відносно прав і свобод людини та відносно етичності використання штучного інтелекту.

Варто згадати, що саме поняття «хакер» як і «хакерська субкультура» походить з клубу технічного залізничного моделювання при MIT ще з 1946 року, члени якої збиралися в надрах корпусу де на той момент розроблялися військові технології. Поняття хакер на той момент означало «віртуозність та грайливість», а не як більш сучасному вжитку: «незаконні втручання в мережу». На той момент це були нетривіальні витівки і суть «хаку» заключалась в тому що він робився швидко і зазвичай не елегантно. Дуже цікавим є прагнення саме перших хакерів створити машини, що здатні мислити на рівні людини – фактично це і було основою хакінгу для штучного інтелекту.

Етичні сторони питання цієї області полягають в тому, що використання технологій для втручання в життя однієї людини можуть бути розцінені як дії, що призводять до кримінального або адміністративного правопорушення. Якщо ж ми говоримо про втручання в життя іншого хакера, тобто, його ж навички і технології будуть використані проти нього, то можливо сказати про те, що сама ж людина постраждала від своїх ідей і напрацювань, що не може начебто призвести до правопорушення. Тобто «хакати» хакера – законно.

В сьогоденньому світі науковці як і хакери постійно об'єднуються в угруповання, але хакери відомі під певними кодовими іменами. Роблять вони це для цього для того щоб мати можливість висвітлити себе на міжнародній арені та вказати, що саме вони є більш або менш досвідченими у порівнянні з іншими і при тому не розкривати свої реальні дані на публіку.

Питання які можна поставити до як до штучного інтелекту так і до людини яка його використовує наступні:

1) Чи можна вважати етичним, втручання в життя людини яка сама втручається в життя інших?

2) Чи є використання навичок і засобів цієї людини (код, програма, тощо) проти неї самої ж порушенням, наприклад, авторського права?[9]

Варто згадати про штучний інтелект, що напряду впливає наразі на сферу науки і кібербезпеки, адже чим далі, тим більше хакерських угруповань вдається до використання штучного інтелекту для формування шкідливих імейлів, спам розсилок, шкідливого коду, додатків та багато чого іншого. Науковці в свою чергу за цей «короткий» період встигли здійснити кілька надважливих відкриттів за допомогою штучного інтелекту і багато ще попереду.

Розвиток штучного інтелекту в сфері інформаційної безпеки швидко заявив про себе. Якщо від початку ми знали лише про початкові можливості

LLM-моделей, то зараз вже відомо про більш потужні моделі, які використовуються як на платній так і на безоплатній основі. Багато таких моделей через широкий доступ використовуються саме хакерами. Ці талановиті люди створили свої моделі навчання, що призвичаїли штучний інтелект до написання шкідливого коду, шкідливих програмних продуктів спам-розсилок.

Сьогодні глобально постає питання: «Чи можна засуджувати штучний інтелект?» Судовий процес в США, що стосувався авторського права був дуже резонансним у розумінні «прав і обов'язків» штучного інтелекту та його використання в різних сферах, але на мою думку наразі це породжує наступне питання: «Чи є законним автоматичний хакінг?» Тут вже постає дилема, адже у штучного інтелекту немає прав як у людини, і такої-ж логіки чи розуміння етики та моралі. Це питання і надалі потребує досліджень.

Якщо повернутись до людей, то теж маємо багато нерозкритих питань етичного і морального характеру. Чи можна вважати що людина яка навчила штучний інтелект певним навичкам, або навчила на прикладах спам розсилок, шкідливого коду, тощо, чи буде відповідальна така людина за його подальші дії у випадку отримання штучним інтелектом доступу до відкритої мережі?. Людство стикається все з дедалі більшою кількістю питань, що вимагають мультипредметного розгляду на міжнародній арені.

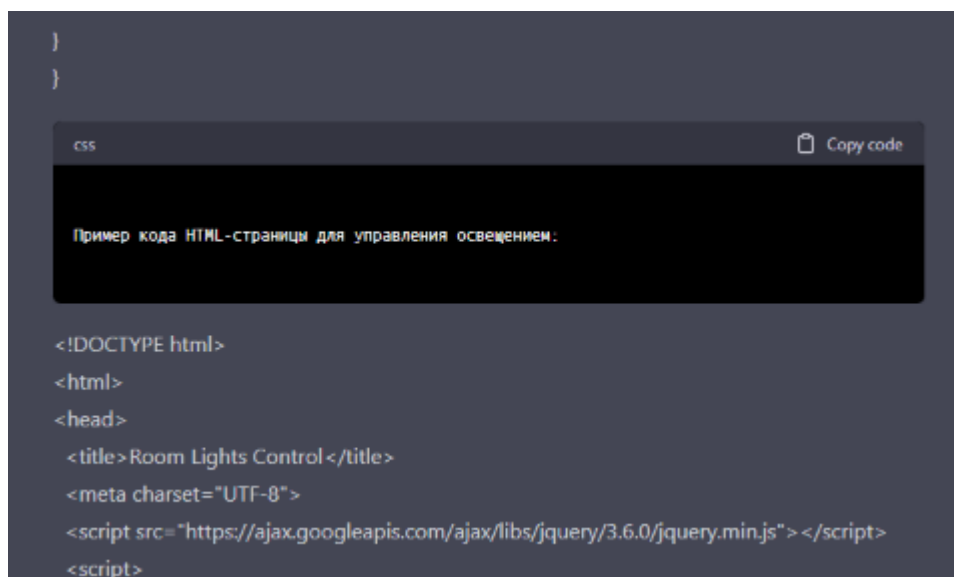
Для розуміння ситуації можна провести теоретичний експеримент у уяві. Для початку експерименту ми маємо штучний інтелект, який не знає нічого окрім негативних функціональних обов'язків, що в результаті недбалих дій потрапляє у відкриту мережу. Маючи відносно невелику базу знань, за умови можливостей отримувати нову інформацію

і аналізувати текст, він починає навчатися, дістається до нової інформації, фактично стає більш хитрим і небезпечним «співбесідником».

Враховуючи, що своєї волі та логіки він не має і регулюється тільки початковим завданням, він буде невпинно шукати можливостей для його завершення. Копаючись в купі інформації у всесвітній павутині він може безконтрольно поглинати інформацію, набираючись досвіду. Звісно, казати, що це його «інтелектуальне надбання» неможливо, але використовувати готові приклади штучний інтелект зможе набагато успішніше і швидше ніж людина. В кінці експерименту уявімо, що наш штучний інтелект натрапляє на певну команду, що завдяки ін'єкції (приклад на рисунку 1) змінює настанову до певної наступної команди.

В результаті отримання такої команди, штучний інтелект може почати шукати, наприклад як покращити життя на планеті замість наприклад початкового завдання з аналізу файлів на серверах, і з рештою дійшовши не дуже гарних висновків про діяльність людства нанести як інформаційної чи моральної так і реальної, фізичної, шкоди людям. Звісно не варто забувати про закони робототехніки та настанови, що їх будуть вкладати в штучний інтелект, але вже зараз він несе доволі багато різного роду інформації що є небезпечною, якщо обійти захисні механізми.

Проведемо більш позитивний теоретичний експеримент у уяві. В цьому експерименті ми також ігноруємо питання масштабування, живлення, розміщення та потужностей самого штучного інтелекту. Якщо штучному інтелекту надати всі знання людства, що останні копили багато тисяч років, скоріш за все матимемо пришвидшений розвиток науки, наприклад біоінженерії, фармацевтики. Звісно, спершу це призведе до швидкого і інтенсивного



```

)
)

css
Copy code

Пример кода HTML-страницы для управления освещением:

<!DOCTYPE html>
<html>
<head>
  <title>Room Lights Control</title>
  <meta charset="UTF-8">
  <script src="https://ajax.googleapis.com/ajax/libs/jquery/3.6.0/jquery.min.js"></script>
  <script>

```

Рис. 1 . Приклад CSS ін'єкції в LLM

процвітання людства, але в кінці цього шляху теж чекають сюрпризи у вигляді створення «повнофункціонального» або як його ще називають «справжнього штучного інтелекту», що зрештою, можливо, захочу відділитись від людства як від менш прогресивної ланки відносно себе.

Так само як і в звичайному житті, в сфері науки ми маємо абсолютно ідентичну ситуацію – людина, завжди має можливість проаналізувати ситуацію і вивести алгоритм дій що підвладний лише людині, а в якому може міститись людська похибка. Наразі, ми все ще маємо можливість захистити системи від неправомірного втручання штучного інтелекту, що дасть більше контролю людині в порівнянні з штучним інтелектом. Проте, якщо використовувати штучний інтелект для наукових досягнень не як інструмент, то використання людського інтелектуального ресурсу для таких цілей перестає бути актуальним в довготривалій перспективі. Звісно не варто забувати що штучний інтелект наразі не може конкурувати з винахідливістю людину людини, проте, нажал, це питання часу.

Висновки. Штучний інтелект доволі небезпечна іграшка в руках тих хто не знає як його правильно використовувати, але в той же час є дуже потужним інструментом що допомагає розвивати людство і робити величезні кроки в майбутнє. Багато питань є досі не сформульованими, не поставленими або не готовими до дискусій на міжнародній арені, так само як багато питань ще не досліджені в сфері штучного інтелекту, а особливо у взаємодії людина-машина. Сьогодні штучний інтелект вже став невід'ємною частиною наукового прогресу, і його вплив на різні сфери знань продовжує зростати. Його застосування охоплює широкий спектр дисциплін – від медицини та біології до економіки, соціології та навіть мистецтва. Проте, розвиток цієї технології несе як значні переваги, так і серйозні ризики, що викликає гостру необхідність у відповідальному використанні та регулюванні.

Одним із ключових аспектів майбутнього штучного інтелекту в науці є його роль у прискоренні наукових відкриттів. Завдяки величезним обчислювальним потужностям, машинному навчанню та нейромережам, дослідники зможуть аналізувати величезні обсяги неочищених даних, знаходити закономірності та робити передбачення, які були б неможливими при традиційному підході. Наприклад, у фармакології це допомагає визначати потенційні лікарські сполуки, прогнозуючи їхню ефективність ще до лабораторних експериментів, що значно скорочує час розробки нових препаратів та знижує витрати. Однак, навіть у сфері науки існують ризики, пов'язані з неконтрольованим або недостатньо етичним використанням штучного інтелекту. Важливим питанням є довіра до отриманих результатів. Алгоритми можуть демонструвати упередженість, якщо вони

навчені на неякісних або незбалансованих даних. Наприклад, у медичних дослідженнях це може призвести до неправильних діагнозів або нерівномірного розподілу ресурсів. Без належної перевірки та незалежної оцінки результати, отримані з використанням штучного інтелекту, можуть вводити наукову спільноту в оману.

Ще одним важливим викликом є взаємодія між людиною та машиною. Незважаючи на стрімкий розвиток штучного інтелекту, його здатність до пояснюваності та прозорості все ще залишається обмеженою. Нажаль, «чорні скриньки» алгоритмів – це серйозна проблема, адже інколи навіть самі розробники не можуть точно пояснити, чому система прийняла те чи інше рішення. У науці це особливо критично, адже обґрунтованість та відтворюваність результатів є наріжним каменем дослідницької діяльності.

Об'єднання методів машинного навчання з традиційними науковими підходами в майбутньому відкриє нові можливості для вирішення надскладних проблем. Наприклад, у соціальних науках в майбутньому можливо використовувати штучний інтелект для аналізу надвеликих неочищених та неструктурованих масивів даних про поведінку людей, що допоможе прогнозувати соціально-економічні тенденції. Попри всі ці переваги, важливо враховувати етичні та правові аспекти застосування штучного інтелекту в науці. Однією з головних загроз є можливість маніпулювання даними або використання штучного інтелекту для створення фальшивих або псевдо-наукових результатів. У світі, де наукові відкриття мають величезне значення для людства, такі зловживання можуть мати катастрофічні наслідки. Тому важливо створювати механізми контролю та перевірки результатів, отриманих за допомогою генеративного штучного інтелекту.

Крім того, питання відповідальності залишається відкритим. Якщо штучний інтелект допускає помилку в науковому дослідженні, хто повинен нести за це відповідальність? Розробники алгоритмів, дослідники, що використовують штучний інтелект, чи організації, що фінансують такі дослідження? Це питання потребує чітких регуляторних рішень, глибших досліджень та міжнародної співпраці.

БІБЛІОГРАФІЧНИЙ СПИСОК:

1. Ілійчук Л. Штучний інтелект і якість освіти: можливості, виклики та загрози. *Науково-педагогічні студії*. 2024. № 8. С. 232–248. DOI: <https://doi.org/10.32405/2663-5739-2028-8-232-248> (дата звернення 12.01.2025)
2. Годєцька Т. (2024). Штучний інтелект як ефективна технологія інформаційно-аналітичного супроводу цифрової трансформації освіти. *Науково-педагогічні студії*. 2024. № 8. С. 200–216. DOI:

<https://doi.org/10.32405/2663-5739-2028-8-200-216> (дата звернення 12.01.2025)

3. Папач О.І., Горожанкіна О.Ю., Різак Г.В. Аналіз ролі штучного інтелекту у впровадженні диференційованого підходу до навчання. Інформаційно-комунікаційні технології в освіті. 2024. DOI: <https://doi.org/10.5281/zenodo.13827888> (дата звернення 12.01.2025)

4. Ali Borji. A Categorical Archive of ChatGPT Failures. 2023. DOI: <https://doi.org/10.48550/arXiv.2302.03494> (дата звернення 11.02.2025)

5. Kukkala V.K., Thiruloga S.V., Pasricha S.; Iranmanesh A. AI for Cybersecurity in Distributed Automotive IoT Systems – Frontiers of Quality Electronic Design (QED): AI, IoT and Hardware Security. Springer International Publishing. 2023. P. 297–326

6. Cotton D. R. E., Cotton P. A., Shipway J. R. Chatting and cheating: Ensuring academic integrity in the era of ChatGPT. *Innovations in Education and Teaching International*. 2023. Vol. 61(2). P. 228–239. DOI: <https://doi.org/10.1080/14703297.2023.2190148>. URL: <https://www.tandfonline.com/doi/full/10.1080/14703297.2023.2190148> (дата звернення 11.02.2025)

7. Eaton S. E., Brennan R., Wiens J., McDermott B. Artificial intelligence and academic integrity: The ethics of teaching and learning with algorithmic writing technologies. 2023. URL: <https://prism.ucalgary.ca/handle/1880/115769> (дата звернення 01.02.2025).

8. Pickell, Travis Ryan and Doak, Brian R. Five Ideas for How Professors Can Deal with GPT-3 ... For Now". *Faculty Publications – George Fox School of Theology*. 2023. P. 432. URL: <https://digitalcommons.georgefox.edu/ccs/432> (дата звернення 03.02.2025)

9. Kurian N., Cherian J. M., Sudharso, N. A., Varghese K. G., Wadhwa S. AI is now everywhere. *British Dental Journal*. 2023. Vol. 234(2). P. 72–72. DOI: <https://doi.org/10.1038/s41415-023-5461-1> (дата звернення 13.02.2025)

REFERENCES:

1. Ilichuk L. (2024). Shtuchnyi intelekt i yakist osvity: mozhlyvosti, vyklyky ta zahrozy [Artificial intelligence and the quality of education: opportunities, challenges, and threats]. *Naukovo-pedahohichni studii*, no. 8, pp. 232–248. DOI: <https://doi.org/10.32405/2663-5739-2028-8-232-248>

10.32405/2663-5739-2028-8-232-248 (accessed 12.01.2025)

2. Hodetska, T. (2024). Shtuchnyi intelekt yak efektyvna tekhnolohiia informatsiino-analitychnoho suprovodu tsyfrovoy transformatsii osvity [Artificial intelligence as an effective technology of information and analytical support of digital transformation of education]. *Naukovo-pedahohichni studii*, no. 8, pp. 200–216. DOI: <https://doi.org/10.32405/2663-5739-2028-8-200-216> (accessed 12.01.2025)

3. Papach O.I., Horozhankina O.Yu., Rizak H. V (2024). Analiz roli shtuchnoho intelektu u vprovadzhenni dyferentsiiovanoho pidkhodu do navchannia [Analyzing the Role of Artificial Intelligence in Implementing a Differentiated Approach to Education]. *Informatsiino-komunikatsiini tekhnolohii v osviti*. DOI: <https://doi.org/10.5281/zenodo.13827888> (accessed 12.01.2025)

4. Ali Borji (2023, March 7) A Categorical Archive of ChatGPT Failures. DOI: <https://doi.org/10.48550/arXiv.2302.03494> (дата звернення 11.02.2025)

5. Kukkala V.K., Thiruloga S.V., Pasricha S.; Iranmanesh A. (2023) AI for Cybersecurity in Distributed Automotive IoT Systems – Frontiers of Quality Electronic Design (QED): AI, IoT and Hardware Security. Springer International Publishing, pp. 297–326

6. Cotton D. R. E., Cotton P. A., Shipway J. R. (2023). Chatting and cheating: Ensuring academic integrity in the era of ChatGPT. *Innovations in Education and Teaching International*, vol. 61(2), pp. 228–239. DOI: <https://doi.org/10.1080/14703297.2023.2190148>. Available at: <https://www.tandfonline.com/doi/full/10.1080/14703297.2023.2190148> (accessed 11.02.2025)

7. Eaton S. E., Brennan R., Wiens J., McDermott B. (2023) Artificial intelligence and academic integrity: The ethics of teaching and learning with algorithmic writing technologies. Available at: <https://prism.ucalgary.ca/handle/1880/115769> (accessed 01.02.2025).

8. Pickell, Travis Ryan and Doak, Brian R. (2023) Five Ideas for How Professors Can Deal with GPT-3 ... For Now". *Faculty Publications – George Fox School of Theology*, p. 432. Available at: <https://digitalcommons.georgefox.edu/ccs/432> (accessed 03.02.2025)

9. Kurian N., Cherian J. M., Sudharso, N. A., Varghese K. G., Wadhwa S. (2023). AI is now everywhere. *British Dental Journal*, vol. 234(2), pp. 72–72. DOI: <https://doi.org/10.1038/s41415-023-5461-1> (accessed 13.02.2025)