



ТЕНДЕНЦІЇ РОЗВИТКУ ТА РЕГУЛЮВАННЯ ШТУЧНОГО ІНТЕЛЕКТУ

Лютий 2025

«Тенденції розвитку та регулювання штучного інтелекту»
підготовлено експертами Громадської організації «Центр
аналізу та розробки ефективних політик».

© ГО «Центр аналізу та розробки ефективних політик»

<https://cadep.org.ua/>
facebook.com/ngo.cadep
linkedin.com/company/ngo-cadep

ЗМІСТ

ІНВЕСТИЦІЇ ТА РИНКОВІ ТРЕНДИ	6
Генеральний директор Nvidia Дженсен Хуанг вважає, що ринок помилився щодо впливу DeepSeek	6
Amazon інвестує ще 6 мільярдів у свої дата-центри в Міссісіпі	6
Європейські стартапи в галузі AI у 2024 році залучили майже 8 млрд. інвестицій	6
Індійський мільярдер інвестує 230 мільйонів доларів в індійський стартап з AI Krutrim	6
Ініціатива ЄС InvestAI: мобілізація €200 мільярдів для розвитку AI	7
Генеральний директор Deutsche Telekom закликав новий уряд Німеччини та Європи інвестувати в AI	7
OpenAI може надати правлінню спеціальні права голосу, щоб запобігти спробам захоплення	7
Рекрутингова платформа Perfect залучає 23 мільйони доларів для імплементації AI в процес підбору персоналу	8
Рекрутинговий AI-стартап Mercor залучив ще 100 мільйонів доларів США	8
Microsoft створює «Підрозділ розширеного планування» для вивчення AI	8
OAE інвестують мільярди в будівництво центру обробки даних AI у Франції	8
AI функції iPhone у Китаї будуть спиратись на потужності компанії Alibaba	9
Meta планує інвестувати в гуманоїдних роботів, керованих AI	9
Apple планує витратити 500 млрд.дол. США на серверну інфраструктуру AI	9
Alibaba планує інвестувати понад \$52 мільярди в AI та хмарні технології протягом наступних трьох років	9
Корпорація Майкрософт стикнулась з надлишком потужностей для обробки AI	10
Patlytics залучає \$14 млн для своєї платформи аналітики патентів	10

ПОЛІТИКА, РЕГУЛЮВАННЯ ТА ЮРИДИЧНІ АСПЕКТИ	10
ЄС ініціював проєкт LLM для посилення свого цифрового суверенітету	10
Макрон закликає Європу спростити свої правила, щоб повернути ЄС до глобальної гонки щодо розвитку штучного інтелекту	11
Європейська комісія опублікувала керівні принципи щодо визначення системи штучного інтелекту	11
Єврокомісія опублікувала Керівні принципи щодо заборонених практик AI	11
Каліфорнійський законопроєкт щодо безпеки дітей і чат-ботів	12
Технічні гіганти та стартапи об'єднують зусилля, щоб закликати до спрощення правил ЄС щодо AI та даних	12
Паризька декларація про збереження контролю людини в системах озброєнь, що підтримують штучний інтелект	12
Thomson Reuters виграла рішення суду про «справедливе використання» авторських прав на штучний інтелект	13
Сенат США схвалив законопроєкт TAKE IT DOWN Act для протидії поширенню AI-генерованих дипфейків.	13
Європа не планує відмовлятися від регуляторних рамок для AI	13
Технологічний стартап Bird залишає Європу через регуляторні обмеження	14
Британські газети протестують проти нового британського закону щодо авторських прав та AI	14
Судові документи показують, що співробітники Meta обговорювали використання захищеного авторським правом контенту для навчання AI	14
Організації вимагають заходів для пом'якшення шкоди AI для навколишнього середовища	15
1000 британських виконавців протестують проти використання AI творів захищених авторським правом	15

ТЕХНОЛОГІЧНІ ІННОВАЦІЇ ТА РОЗРОБКИ 15

Boston Dynamics прискорює розвиток людиноподібного робота Atlas через впровадження методів навчання з підкріпленням 15

Google планує зробити пошук у 2025 році більш схожим на розумного AI-асистента 16

AI чат-бот Meta AI став доступним на Близькому Сході та в Африці 16

Катар уклав угоду зі Scale AI для впровадження штучного інтелекту з метою покращення ефективності державних послуг 17

Google використовуватиме машинне навчання для оцінки віку користувача 17

Штучний інтелект виготовляє взуття 17

Nvidia використовує штучний інтелект для навчання мови жестів 18

«Career Dreamer» від Google допомагає підбирати роботу 18

YouTube Shorts інтегрує можливість створення відеокліпів, створених за допомогою AI 18

Випуск Grok 3 дозволив xAI залучити нових користувачів більш ніж на 260% 19

Агентський AI використовуються на заводах 19

СОЦІАЛЬНІ, ЕТИЧНІ ТА МЕДІЙНІ ПИТАННЯ 19

У Google обіцяють не розробляти зброю зі штучним інтелектом 19

Піонерка штучного інтелекту Фей-Фей Лі застерігає політиків не дозволяти науково-фантастичним сенсаціям впливати на правила AI 20

AI системи викривляють новини 20

Опитування Єврокомісії показало, що більшість європейців підтримують використання AI на робочому місці 20

AI Grok заявив, що Ілон Маск та Дональд Трамп заслуговують на смертну кару 21

Apple обіцяє виправити помилку транскрипції, яка замінює "Racist" на "Trump" 21

Чат-бот Sonar Mental Health має допомогти школам з нестачею консультантів з питань психічного здоров'я 21

The New York Times впроваджує інструменти штучного інтелекту 21

КІБЕРБЕЗПЕКА, ТЕХНІЧНА НАДІЙНІСТЬ ТА ЗАХИСТ ПЕРСОНАЛЬНИХ ДАНИХ 22

Італійська поліція викрила шахрайство зі штучним інтелектом 22

Китай використовує інструменти спостереження за соцмережами на основі AI 22

Південна Корея блокує завантаження DeepSeek з міркувань конфіденційності 23

Google Photos впроваджує цифрові водяні знаки SynthID для зображень, змінених за допомогою AI 23

R1 від DeepSeek «більш вразливий» до джейлбрейка, ніж інші моделі AI 23

Протидія кібербезпековим ризикам AI за допомогою цифрової ідентифікації 23

Інститут безпеки штучного інтелекту США може зіткнутися з великими скороченнями 24

Microsoft розкрила імена розробників, відповідальних за нелегальні AI-інструменти, що використовуються для створення дипфейків знаменитостей 24

НАУКОВІ ДОСЛІДЖЕННЯ ТА ОСВІТА 25

Дослідники навчають AI інтерпретувати емоції тварин 25

Microsoft робить Copilot Voice і Think Deeper безкоштовними з необмеженим використанням 25

Google демонструє новий інструмент штучного інтелекту для вчених 25

ІНВЕСТИЦІЇ ТА РИНКОВІ ТРЕНДИ

Генеральний директор Nvidia Дженсен Хуанг вважає, що ринок помилився щодо впливу DeepSeek

Генеральний директор Nvidia, Дженсен Хуанг, вважає, що ринок неправильно оцінив вплив нової моделі AI DeepSeek на бізнес компанії. Після випуску відкритої моделі R1 від DeepSeek акції Nvidia втратили майже 17% вартості, що призвело до зниження ринкової капіталізації на \$600 мільярдів. Однак Хуанг підкреслює, що ці побоювання безпідставні, оскільки нові методи, такі як R1, збільшують потребу в обчислювальних ресурсах для післятренувальних процесів, що, навпаки, підвищує попит на апаратне забезпечення Nvidia. Він також зазначив, що відкритий доступ до R1 стимулюватиме ширше впровадження технологій штучного інтелекту, що позитивно вплине на галузь загалом.

Amazon інвестує ще 6 мільярдів у свої дата-центри в Міссісіпі

Amazon Web Services (AWS) збільшує свої інвестиції в дата-центри в штаті Міссісіпі до 16 мільярдів доларів, що на 6 мільярдів доларів більше від початково оголошеної суми. Ці кошти будуть спрямовані на будівництво двох кампусів дата-центрів у окрузі Медісон, на північ від столиці штату, Джексона. Проект передбачає створення щонайменше 1 000 нових високотехнологічних робочих місць у регіоні.

Європейські стартапи в галузі AI у 2024 році залучили майже 8 млрд. інвестицій

У 2024 році європейські стартапи в галузі AI залучили близько \$8 мільярдів венчурного капіталу, що становить приблизно 20% від загального обсягу венчурних інвестицій у регіоні. Це свідчить про значне зростання інтересу інвесторів до європейського AI-сектору, незважаючи на загальне зниження венчурного фінансування в Європі, яке склало \$51 мільярд у 2024 році, що на 5% менше порівняно з попереднім роком. Серед помітних угод — фінансування французького стартапу Mistral, який розробляє передові мовні моделі та залучив значні інвестиції для підтримки своїх досліджень та розробок.

Індійський мільярдер інвестує 230 мільйонів доларів в індійський стартап з AI Krutrim

Засновник Ola Бхавіш Аггарвал інвестує 230 мільйонів доларів у свій стартап зі штучним інтелектом оскільки Індія прагне утвердитися в галузі, де домінують компанії США та Китаю. Аггарвал інвестує в Krutrim, який створює великі мовні моделі (LLM) для індійських мов. Krutrim прагне залучити інвестиції в розмірі 1,15 мільярда доларів до наступного року. Стартап Krutrim створює свої моделі AI відкритим вихідним кодом і оприлюднив план побудувати у партнерстві з Nvidia найбільший суперкомп'ютер Індії.

Krutrim випустив Krutrim-2 , мовну модель із 12 мільярдами параметрів, яка показала високу продуктивність у обробці індійських мов.

Ініціатива ЄС InvestAI: мобілізація €200 мільярдів для розвитку AI

11 лютого 2025 року на Саміті з дій у сфері штучного інтелекту в Парижі президент Європейської комісії Урсула фон дер Ляєн оголосила про запуск ініціативи InvestAI. Ця програма спрямована на мобілізацію €200 мільярдів інвестицій у сферу AI, включаючи створення нового європейського фонду в розмірі €20 мільярдів для будівництва "гігафабрик" AI (AI gigafactories), що призначені для відкритої та спільної розробки найскладніших моделей AI, що дозволить Європі стати лідером у цій галузі. Президент ЄК фон дер Ляєн підкреслила, що AI покращить охорону здоров'я, стимулюватиме дослідження та інновації, а також підвищить конкурентоспроможність Європи.

Генеральний директор Deutsche Telekom закликав новий уряд Німеччини та Європи інвестувати в AI

Виконавчий директор Deutsche Telekom Тім Хеттгес закликав новий уряд Німеччини та Європу взяти на себе більше лідерства у розвитку AI та центрів обробки даних, щоб конкурувати зі США та Азією. Хеттгес заявив, що він очікує, що попит на центри обробки даних в Європі зросте як мінімум на 30%, незважаючи на розвиток таких програм, як DeepSeek, які задовольняються меншою обчислювальною потужністю. За його словами, Deutsche Telekom інвестує в дата-центри і знаходиться в процесі будівництва чотирьох з них в Європі, прагнучи побудувати до одного гігавата інфраструктури.

OpenAI може надати правлінню спеціальні права голосу, щоб запобігти спробам захоплення

OpenAI розглядає можливість надання своїй некомерційній раді директорів спеціальних прав голосу для захисту від ворожих спроб поглинання, зокрема від співзасновника компанії Ілона Маска. Цей крок спрямований на збереження контролю ради над організацією після її переходу до структури for-profit.

Нещодавно Маск запропонував придбати OpenAI за \$97,4 мільярда, але рада директорів одноголосно відхилила цю пропозицію, назвавши її спробою "порушити конкуренцію".

Запровадження спеціальних прав голосу дозволить некомерційній раді переважати над іншими інвесторами, такими як Microsoft та SoftBank, забезпечуючи збереження місії OpenAI щодо розвитку штучного інтелекту на благо людства.

Ці події є наслідком напруженості між Ілоном Маском та генеральним директором OpenAI Семом Альтманом, яка виникла після відходу Маска з компанії у 2018 році та його критики щодо її комерційного напрямку.

Рекрутингова платформа Perfect залучає 23 мільйони доларів для імплементації AI в процес підбору персоналу

Perfect, рекрутингова платформа на основі штучного інтелекту, залучила 23 мільйони доларів США для вирішення проблеми неефективності процесу найму на роботу. Perfect має на меті вдосконалити процес підбору персоналу, використовуючи AI для оптимізації відбору кандидатів та покращення загальних результатів найму.

Рекрутинговий AI-стартап Mercor залучив ще 100 мільйонів доларів США

Mercor, стартап у сфері рекрутингу на основі AI, заснований у 2023 році трьома 21-річними стипендіатами Пітера Тіля (мільярдер який підтримує технологічні стартапи; співзасновник PayPal) — Бренданом Фуді, Адрашем Хірематом та Сур'єю Мідхою, нещодавно залучив \$100 мільйонів, що підняло оцінку компанії до \$2 мільярдів.

Mercor спеціалізується на автоматизації процесу найму, використовуючи AI для скринінгу резюме, підбору кандидатів, проведення інтерв'ю та управління оплатою праці. Платформа дозволяє роботодавцям надавати опис вакансій, а система рекомендує кандидатів на основі їхніх кваліфікацій та прогнозованої ефективності на посаді.

Microsoft створює «Підрозділ розширеного планування» для вивчення AI

Microsoft створює новий підрозділ під назвою Advanced Planning Unit для дослідження штучного інтелекту. Цей підрозділ працюватиме в Офісі генерального директора Microsoft AI та об'єднає передові дослідження як всередині компанії, так і за її межами.

ОАЕ інвестують мільярди в будівництво центру обробки даних AI у Франції

Президент Франції Еммануель Макрон і Мохамед бін Заїд Аль-Нахайян з Об'єднаних Арабських Еміратів підписали угоду про значні інвестиції у будівництво у Франції великого центру обробки даних, присвяченого AI. Франція та ОАЕ можуть витратити від 30 до 50 мільярдів євро на будівництво кампусу AI. Більша частина інвестицій буде спрямована на центр обробки даних потужністю до 1 ГВт. Міністр цифрових технологій та штучного інтелекту Франції Клара Чаппаз заявила, що Франція визначила 35 місць, які можна використовувати для будівництва нових центрів обробки даних. Близько 65% виробництва електроенергії у Франції припадає на атомні електростанції. Країна також отримує близько 25% електроенергії з відновлюваних джерел. Коментуючи нещодавню американську ініціативу Stargate та китайський стартап DeepSeek, Чаппаз висловила надію щодо майбутніх перспектив Франції та Європи.

AI функції iPhone у Китаї будуть спиратись на потужності компанії Alibaba

Apple планує запровадити функції штучного інтелекту для користувачів iPhone у Китаї, співпрацюючи з місцевою технологічною компанією Alibaba. Ця ініціатива має на меті адаптувати Apple Intelligence до китайського ринку, забезпечуючи відповідність місцевим регуляторним вимогам та інтеграцію з популярними китайськими сервісами. Очікується, що більшість функцій Apple Intelligence у Китаї будуть підтримуватися інфраструктурою Alibaba, тоді як деякі, зокрема Visual Intelligence, будуть реалізовані за допомогою технологій компанії Baidu. Запуск цих функцій у Китаї заплановано на травень 2025 року.

Meta планує інвестувати в гуманоїдних роботів, керованих AI

У лютому 2025 року компанія Meta оголосила про створення нового підрозділу (як частини Reality Labs), спрямованого на розробку гуманоїдних роботів, керованих штучним інтелектом, для виконання фізичних завдань. Тим самим Meta включилась у конкурентне суперництво у сфері робототехніки, де домінуючі позиції займають Figure AI та Tesla. Нова група зосередиться на дослідженнях і розробках "споживчих гуманоїдних роботів" з метою максимального використання можливостей платформи Llama — основної AI-моделей компанії.

Apple планує витратити 500 млрд.дол. США на серверну інфраструктуру AI

Apple оголосила про плани інвестувати 500 мільярдів доларів у США та створити 20 000 робочих місць за чотири роки, а також відкрити фабрику в Техасі для виробництва серверів AI. Це сталося після зустрічі Тіма Кука з Дональдом Трампом, який заявив, що компанія переносить виробництво до США, щоб уникнути тарифів. Новий завод у Х'юстоні займеться випуском серверів для центрів обробки даних Apple у кількох штатах. Водночас основне виробництво iPhone, iPad і Mac залишиться в Азії. Раніше Apple переносила частину виробництва до В'єтнаму та Індії, але не в США. На фоні тарифів на китайські товари Apple прагне уникнути мит, як це було в 2020 році. Аналітики припускають, що компанія використовує ці інвестиції як спосіб захистити свій бізнес від податкових обмежень адміністрації Трампа.

Alibaba планує інвестувати понад \$52 мільярди в AI та хмарні технології протягом наступних трьох років

Компанія Alibaba Group оголосила про плани інвестувати понад 380 мільярдів юанів (приблизно \$52,41 мільярда) у розвиток інфраструктури штучного інтелекту (AI) та хмарних обчислень протягом наступних трьох років. Ця інвестиція перевищує загальні витрати компанії на AI та хмарні технології за останнє десятиліття. Генеральний директор Едді Ву зазначив, що ці інвестиції спрямовані на задоволення зростаючого попиту на AI-інфраструктуру та зміцнення позицій Alibaba в епоху штучного інтелекту. Крім того, Alibaba планує збільшити витрати на дослідження та розробки у сфері базових AI-моделей та AI-додатків для трансформації інших напрямків бізнесу.

Корпорація Майкрософт стикнулася з надлишком потужностей для обробки AI

Нещодавно Microsoft скасувала оренду кількох дата-центрів у США, загальною потужністю близько кількох сотень мегават, через потенційний надлишок потужностей. Це рішення передбачає розірвання угод з приватними операторами дата-центрів і призупинення будівництва деяких проєктів, в тому числі об'єкта у Вісконсині, призначеного для підтримки OpenAI. Незважаючи на ці корективи, Microsoft зберігає своє зобов'язання інвестувати понад 80 мільярдів доларів в інфраструктуру в поточному фінансовому році, щоб задовольнити зростаючий попит клієнтів. Компанія заявила, що може стратегічно прискорити або скоригувати нашу інфраструктуру в деяких сферах, але продовжує планувати стабільне зростання. Цей розвиток подій призвів до падіння акцій Microsoft і викликав занепокоєння інвесторів щодо стратегії компанії в галузі AI та планування інфраструктури.

Patlytics залучає \$14 млн для своєї платформи аналітики патентів

Стартап Patlytics, підтримуваний Google, оголосив про залучення \$14 млн для розвитку AI компоненти своєї платформи аналітики патентів, яка допомагає компаніям та дослідникам отримувати інформацію з патентних даних, сприяючи інноваціям та розвитку технологій.

ПОЛІТИКА, РЕГУЛЮВАННЯ ТА ЮРИДИЧНІ АСПЕКТИ

ЄС ініціював проєкт LLM для посилення свого цифрового суверенітету

У лютому 2025 року Європейський Союз ініціював проєкт OpenEuroLLM, спрямований на розробку відкритих великих мовних моделей (LLM), що охоплюють усі офіційні мови ЄС. Ця ініціатива має на меті зміцнити цифровий суверенітет Європи, зменшуючи залежність від іноземних технологічних гігантів та забезпечуючи відповідність європейським стандартам прозорості та етики. Проєкт об'єднує 20 організацій, включаючи стартапи, дослідницькі лабораторії та суперкомп'ютерні центри, з бюджетом у €37,4 мільйона. Моделі будуть тренуватися на європейських суперкомп'ютерах, дотримуючись принципів відкритості та доступності, відповідно до вимог AI Act. Експерти відзначають, що заявлений бюджет лише на створення самих моделей становить 37,4 мільйона євро, з яких приблизно 20 мільйонів євро надходять від програми ЄС Цифрова Європа – це вкрай мало для цієї сфери порівняно з тим, що інвестують гіганти корпоративного світу AI. Серед партнерів проєкту OpenEuroLLM – суперкомп'ютерні центри EuroHPC в Іспанії, Італії, Фінляндії та Нідерландах, а бюджет ширшого проєкту EuroHPC становить близько 7 мільярдів євро.

Макрон закликає Європу спростити свої правила, щоб повернути ЄС до глобальної гонки щодо розвитку штучного інтелекту

На Саміті зі штучного інтелекту (AI Action Summit) в Парижі президент Франції Макрон оголосив про інвестиційний пакет у розмірі 109 мільярдів євро у французьку екосистему AI. Також він підтвердив фінансові зобов'язання приватних партнерів, які бажають будувати центри обробки даних у Франції та інвестувати в стартапи AI. За його словами, основна причина, чому міжнародні інвестори обирають Францію для свого наступного масштабного проекту центру обробки даних (ЦОД), зводиться до надлишку ядерної енергії в країні. Він також зазначив, що найближчим часом ЄС оголосить Європейську стратегію AI яка дозволить Європі прискорити розвиток AI та інвестувати в обчислювальну потужність. Макрон також закликав європейські компанії купувати в європейських стартапів - за його словами, більшість компаній у США та Китаї віддають перевагу домашнім рішенням.

Європейська комісія опублікувала керівні принципи щодо визначення системи штучного інтелекту

Європейська комісія опублікувала керівні вказівки щодо визначення систем AI для полегшення застосування перших норм Закону про AI (AI Act). Ці вказівки допомагають постачальникам та іншим зацікавленим сторонам визначити, чи є певне програмне забезпечення системою AI відповідно до правових норм AI Act. Документ не є обов'язковим і буде оновлюватися з урахуванням практичного досвіду та нових випадків використання. Перші норми AI Act, включаючи визначення системи AI, підвищення обізнаності про AI та обмежену кількість заборонених випадків використання AI, що становлять неприйнятні ризики, набули чинності 2 лютого 2025 року.

Єврокомісія опублікувала Керівні принципи щодо заборонених практик AI

Європейська комісія опублікувала керівні вказівки щодо заборонених практик AI, визначених у Законі про AI (AI Act). Ці вказівки надають огляд практик AI, які вважаються неприйнятними через потенційні ризики для європейських цінностей та фундаментальних прав. Метою цих вказівок є забезпечення послідовного та ефективного застосування AI Act по всьому Європейському Союзу. Документ не є юридично обов'язковим; остаточні інтерпретації залишаються за Судом Європейського Союзу.

Заборонені практики включають:

- **Шкідливі маніпуляції та обман:** використання AI для маніпулювання поведінкою людей у спосіб, що може завдати їм шкоди.
- **Експлуатація специфічних характеристик людей:** застосування AI для використання вразливостей окремих осіб, зокрема через їхній вік, фізичні або розумові здібності.
- **Соціальний скоринг:** оцінювання надійності осіб на основі їхньої поведінки або особистих характеристик, що може призвести до дискримінації.
- **Оцінка ризику кримінальних правопорушень:** використання AI для прогнозування ймовірності вчинення особою злочину без належних підстав.

- **Ненаправлене збирання даних для розпізнавання облич:** створення або розширення баз даних для розпізнавання облич шляхом масового збору зображень без згоди.
- **Розпізнавання емоцій на робочих місцях та в освітніх установах:** використання AI для моніторингу емоційного стану працівників або студентів.
- **Біометрична категоризація:** використання біометричних даних для визначення певних захищених характеристик осіб.

Каліфорнійський законопроект щодо безпеки дітей і чат-ботів

У лютому 2025 року в Каліфорнії було запропоновано законопроект SB 243, який вимагає від компаній, що розробляють штучний інтелект, періодично нагадувати дітям, що чат-боти є AI, а не людьми. Ця ініціатива спрямована на підвищення обізнаності серед молодих користувачів щодо взаємодії з AI-системами та забезпечення їх безпеки в онлайн-середовищі.

Раніше, у вересні 2024 року, Каліфорнія прийняла закон SB 1047, який зобов'язує розробників AI впроваджувати запобіжні заходи для запобігання потенційним катастрофам, таким як кібератаки. Цей закон отримав широку підтримку серед голлівудських знаменитостей та профспілок, включаючи SAG-AFTRA, які закликали губернатора Гевіна Ньюсома підписати його до 30 вересня 2024 року.

Технічні гіганти та стартапи об'єднують зусилля, щоб закликати до спрощення правил ЄС щодо AI та даних

У лютому 2025 року провідні технологічні компанії та стартапи об'єдналися, закликаючи Європейський Союз спростити правила щодо використання даних для AI. Вони стверджують, що надмірно складні регуляції можуть стримувати інновації та розвиток AI в Європі. Компанії пропонують більш гнучкий підхід до обробки даних, який би враховував швидкий розвиток технологій та потреби бізнесу, забезпечуючи при цьому захист прав користувачів.

Паризька декларація про збереження контролю людини в системах озброєнь, що підтримують штучний інтелект

У лютому 2025 року на Паризькому саміті з питань штучного інтелекту 25 країн підписали Паризьку декларацію про збереження контролю людини в системах озброєнь, що підтримують штучний інтелект (Paris Declaration on Maintaining Human Control in AI enabled Weapon Systems). Основна ідея Декларації – автономні системи не мають ухвалювати рішення про життя та смерть без людського контролю. Індія, Велика Британія та США, не приєдналися до цієї декларації.

У офіційній заяві викладено зобов'язання щодо відповідального управління інтеграцією AI у військові операції. Наголошується на необхідності використання потенціалу штучного інтелекту, усунення його ризиків відповідно до міжнародного права, зокрема міжнародного гуманітарного права. Підкреслюється, що відповідальність не може бути передана машинам, і вимагає, щоб рішення, пов'язані з життям і смертю, залишалися під контролем людини. Декларація підкреслює, що всі системи зброї, включно з тими, що створені за допомогою AI, повинні працювати відповідно до правових рамок. Крім

того, документ закликає до тривалого багатостороннього діалогу та міжнародного співробітництва у відповідальній розробці, розгортанні та використанні AI у військовій сфері.

Thomson Reuters виграла рішення суду про «справедливе використання» авторських прав на штучний інтелект

У лютому 2025 року федеральний суд США виніс рішення щодо авторського права та AI, постановивши, що використання компанією ROSS Intelligence матеріалів із Westlaw (платформа, що створена Thomson Reuters) для навчання свого AI-інструменту без дозволу порушує авторські права Thomson Reuters. Суддя Стефанос Бібас відхилив аргументи ROSS про “чесне використання” (fair use), зазначивши, що копіювання без ліцензії не відповідає критеріям чесного використання. Це рішення може мати значні наслідки для інших компаній, що розробляють AI, оскільки підкреслює важливість отримання дозволу на використання захищених авторським правом матеріалів для навчання AI-моделей. Цей випадок є частиною ширшої тенденції, де автори та видавці подають позови проти технологічних компаній за використання їхніх творів для навчання AI без належного дозволу.

Сенат США схвалив законопроект TAKE IT DOWN Act для протидії поширенню AI-генерованих дипфейків.

13 лютого 2025 року Сенат США одноголосно ухвалив законопроект TAKE IT DOWN Act, спрямований на боротьбу з розповсюдженням несанкціонованих інтимних зображень, включаючи створені за допомогою штучного інтелекту “deepfakes”. Законопроект, ініційований сенаторами Тедом Крузом (республіканець, Техас) та Емі Клобучар (демократка, Міннесота), передбачає кримінальну відповідальність за публікацію таких матеріалів та зобов’язує онлайн-платформи видаляти їх протягом 48 годин після отримання повідомлення від жертви.

Хоча законопроект отримав широку підтримку, деякі організації, такі як Electronic Frontier Foundation (EFF), висловили занепокоєння щодо потенційного впливу на свободу слова та приватність користувачів. Вони стверджують, що вимоги щодо швидкого видалення контенту можуть призвести до цензури законного контенту та порушення належних правових процедур. Наразі законопроект очікує розгляду в Палаті представників.

Європа не планує відмовлятися від регуляторних рамок для AI

У лютому 2025 року Європейський Союз спростував твердження про те, що він відмовився від запровадження правил відповідальності для AI під тиском адміністрації президента США Дональда Трампа. Речник Європейської комісії наголосив, що ЄС залишається відданим розробці регуляторних рамок для AI, спрямованих на забезпечення безпеки та прозорості технологій, незалежно від зовнішнього впливу. Раніше з’явилися повідомлення, що адміністрація Трампа висловила занепокоєння щодо запропонованих ЄС правил відповідальності за AI, вважаючи, що вони можуть обмежити інновації та створити бар’єри для американських технологічних компаній на європейському ринку. Однак Європейська комісія підкреслила, що її пріоритетом є

захист прав споживачів та забезпечення етичного використання AI, і будь-які рішення прийматимуться виключно в інтересах громадян ЄС.

Технологічний стартап Bird залишає Європу через регуляторні обмеження

Один із провідних нідерландських стартапів у сфері хмарних комунікацій, компанія Bird, оголосив про скорочення своєї діяльності в Європі. Генеральний директор Роберт Віс пояснив рішення жорсткими регуляторними обмеженнями та браком кваліфікованих технічних фахівців. За його словами, Європа не створює належного середовища для інновацій в епоху штучного інтелекту, а місцеве регулювання лише загальмує розвиток технологій. Компанія планує перенести основні операції до Нью-Йорка, Сінгапура та Дубая. Заснована в Амстердамі у 2011 році, Bird (раніше MessageBird) є конкурентом Twilio, розробляючи платформи для автоматизації бізнес-комунікацій через штучний інтелект. Її рішення спрощують взаємодію компаній із клієнтами через повідомлення, електронну пошту та відеозв'язок.

Британські газети протестують проти нового британського закону щодо авторських прав та AI

25 лютого 2025 року провідні британські газети об'єдналися в кампанії "Make It Fair", протестуючи проти урядових пропозицій, які можуть послабити авторські права на користь компаній, що розробляють AI. Ці пропозиції дозволять технологічним фірмам використовувати захищений авторським правом контент для навчання своїх моделей AI без отримання дозволу або виплати компенсації творцям. Газети, такі як The Guardian, The Times та The Sun, розмістили на своїх перших шпальтах єдине зображення з вимогою забезпечити справедливість для креативних індустрій.

Кампанію підтримали численні організації, включаючи News Media Association (NMA) та Creative Rights in AI Coalition, які висловили занепокоєння щодо потенційної шкоди для креативного сектору Великобританії, вартість якого оцінюється у 152 мільярди доларів. Пропозиції уряду передбачають внесення змін до законодавства про авторське право, що дозволить компаніям AI використовувати будь-які дані, включаючи комерційний контент, для навчання своїх моделей без необхідності отримувати дозвіл.

Судові документи показують, що співробітники Meta обговорювали використання захищеного авторським правом контенту для навчання AI

Нещодавні судові документи розкрили внутрішні дискусії серед співробітників компанії Meta щодо використання захищених авторським правом матеріалів, зокрема контенту з піратського веб-сайту Library Genesis (LibGen), для навчання моделей AI компанії.

Ці викриття з'явилися в результаті судового позову, поданого такими авторами, як Сара Сільверман, Річард Кадрі та Крістофер Голден, які стверджують, що Meta використовувала їхні твори без дозволу для розробки своїх AI-технологій.

Розсекречені документи свідчать про те, що співробітники Meta обговорювали правові та етичні наслідки використання захищеного авторським правом контенту для навчання AI. Незважаючи на внутрішні занепокоєння, повідомлення свідчать про те, що

генеральний директор компанії Марк Цукерберг схвалив використання таких матеріалів для покращення можливостей штучного інтелекту в Meta.

Організації вимагають заходів для пом'якшення шкоди AI для навколишнього середовища

Група з понад 100 організацій (серед них і відомі правозахисні групи - Amnesty International та AI Now Institute) опублікувала відкритий лист із закликом до індустрії штучного інтелекту та регуляторів пом'якшити шкідливий вплив технологій на навколишнє середовище. У листі зазначається, системи штучного інтелекту збільшують викиди, закріплюють залежність країн від невідновлюваних джерел енергії та виснажують критично важливі ресурси. Проте мало що робиться для усунення цих негативних зовнішніх ефектів, оскільки технологічний сектор і уряди виправдовують подальші інвестиції в AI. Підписанти закликають до того, щоб інфраструктура AI, включаючи центри обробки даних, не використовувала викопне паливо. Поспих у створенні інфраструктури для розробки та запуску штучного інтелекту переважав електричні мережі, змусивши деякі комунальні підприємства використовувати вугілля та інші екологічно небезпечні джерела енергії.

1000 британських виконавців протестують проти використання AI творів захищених авторським правом

Понад 1000 британських музикантів випустили тихий альбом під назвою *Is This What We Want?* на знак протесту проти запропонованих змін до британського законодавства про авторське право. Уряд планує дозволити компаніям, що займаються розробкою AI, використовувати захищені авторським правом твори в навчальних цілях без отримання попереднього дозволу від авторів. Альбом складається з 12 тихих треків, кожен з яких представляє потенційну порожнечу, яку залишає відсутність людської творчості, а назви формують акровіршоване послання: Британський уряд не повинен легалізувати крадіжку музики на користь AI-компаній. Доходи від продажу альбому будуть передані благодійній організації *Help Musicians*. Уряд стверджує, що нинішній режим захисту авторських прав обмежує потенціал креативної та AI-індустрії.

ТЕХНОЛОГІЧНІ ІННОВАЦІЇ ТА РОЗРОБКИ

Boston Dynamics прискорює розвиток людиноподібного робота Atlas через впровадження методів навчання з підкріпленням

У лютому 2025 року компанія Boston Dynamics оголосила про співпрацю з Інститутом робототехніки та AI (RAI Institute), очолюваним її засновником та колишнім генеральним директором Марком Райбертом. Метою цієї співпраці є прискорення розвитку людиноподібного робота Atlas шляхом впровадження методів навчання з підкріпленням (Reinforcement Learning). Цей підхід дозволить роботу самостійно освоювати складні

завдання, такі як відкривання дверей та біг, через процес проб і помилок. Райберт зазначив, що ця ініціатива сприятиме розширенню навичок Atlas та оптимізації процесу набуття нових умінь.

Раніше ці організації вже співпрацювали над розробкою комплексу для навчання з підкріпленням для чотириноного робота Spot, що дозволило значно покращити його швидкість та маневреність. Нинішній проект зосереджений на покращенні мобільності та взаємодії Atlas з фізичним оточенням, а також на перенесенні навичок, отриманих у віртуальних симуляціях, у реальний світ.



Google планує зробити пошук у 2025 році більш схожим на розумного AI-асистента

Google розвиває свій пошуковий продукт, щоб він більше нагадував асистента зі штучним інтелектом, який активно переглядає інтернет від імені користувачів, аналізуючи веб-сторінки для надання прямих відповідей на запити. Центральним елементом цієї ініціативи є Gemini, передова модель AI від Google. Інтегрована в різні сервіси, Gemini розширює функціональність екосистеми Google. Наприклад, функція **Запитайте мене** на базі Gemini може самостійно зв'язуватися з компаніями, щоб дізнатися про послуги та ціни, спрощуючи такі завдання, як планування зустрічей. Крім того, Google представив AI-огляди в пошуку, призначені для надання стислих резюме складних тем, тим самим покращуючи доступність інформації та її розуміння користувачами.

AI чат-бот Meta AI став доступним на Близькому Сході та в Африці

Компанія Meta розширила доступність свого AI-чатбота, Meta AI, на Близький Схід та Африку, додавши підтримку арабської мови. Раніше, у жовтні 2024 року, Meta оголосила про розширення доступності Meta AI на 21 новий ринок, включаючи Великобританію, Бразилію та країни Латинської Америки та Азії, з підтримкою таких мов, як арабська, індонезійська, тайська та в'єтнамська. Це розширення спрямоване на конкуренцію з іншими AI-чатботами, такими як ChatGPT від OpenAI.

Катар уклав угоду зі Scale AI для впровадження штучного інтелекту з метою покращення ефективності державних послуг

Уряд Катару уклав п'ятирічну угоду з компанією Scale AI, що спеціалізується на управлінні даними для штучного інтелекту, для впровадження AI-інструментів та програм навчання з метою підвищення ефективності державних послуг.

Ця угода може стати зразком для інших урядів у світі та дозволяє нам взяти на себе зобов'язання в такий спосіб, що, на мою думку, сприятиме ще швидшому впровадженню інновацій, — зазначив Тревор Томпсон, керівник глобального відділу зростання Scale AI, що базується у Сан-Франциско.

Згідно із заявою Міністерства зв'язку та інформаційних технологій Катару, угода передбачає використання прогнозової аналітики, автоматизації та розширеного аналізу даних для оптимізації державних процесів і підвищення ефективності управління.

Google використовуватиме машинне навчання для оцінки віку користувача

Google планує використовувати машинне навчання для оцінки віку користувачів, щоб надавати їм відповідний контент та забезпечувати безпечний досвід. Цей підхід дозволить компанії визначати вік користувачів без необхідності запитувати особисту інформацію, таку як дата народження. Подібні технології вже впроваджуються іншими компаніями; наприклад, YouTube планує використовувати технологію для оцінки віку користувачів, а Meta розробила систему штучного інтелекту, яка аналізує "сигнали", щоб визначити, чи користувач неправильно вказав свій вік.

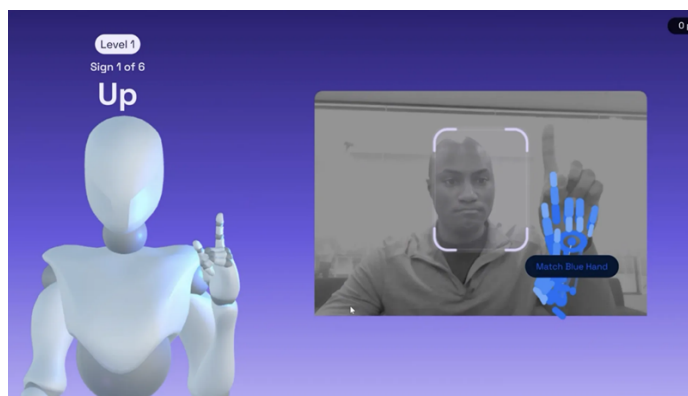
Штучний інтелект виготовляє взуття

Syntilay, стартап, заснований 25-річним підприємцем Беном Вайсом, представив перше у світі взуття, повністю спроектоване штучним інтелектом та надруковане на 3D-принтері. Це футуристичне взуття створене за допомогою поєднання AI-інструментів, таких як Midjourney та Vizcom AI, а також людської творчості. Процес розробки включав генерацію базової форми взуття за допомогою Midjourney, подальше доопрацювання художником, створення 3D-моделі через Vizcom AI та додавання текстур і візерунків за допомогою генеративного AI. Кожна пара виготовляється на замовлення з використанням 3D-друку, забезпечуючи індивідуальну підгонку завдяки скануванню стопи за допомогою смартфона. Взуття доступне в п'яти кольорах: синьому, чорному, червоному, бежевому та оранжевому, за ціною \$149.99 за пару. Проект підтримується Джо Фостером, співзасновником Reebok, який надає свій досвід для розвитку Syntilay. Компанія планує виготовити кілька тисяч пар для підвищення впізнаваності бренду, а в майбутньому пропонувати унікальні дизайни для контент-креаторів та інших брендів.



Nvidia використовує штучний інтелект для навчання мови жестів

У лютому 2025 року компанія NVIDIA представила безкоштовний AI-інструмент під назвою Signs, призначений для навчання американської жестової мови (ASL). Цей інструмент використовує машинне навчання та комп'ютерний зір, щоб надавати користувачам зворотний зв'язок у реальному часі під час формування жестів, допомагаючи їм покращувати свої навички. Особливо важливою є ця розробка для батьків, оскільки більшість глухих дітей народжуються у чуючих сім'ях, і доступ до якісних ресурсів для вивчення ASL є обмеженим. Ті, хто вивчає мову, можуть отримати доступ до бібліотеки знаків, які були перевірені людьми, які вільно володіють ASL. Потім інструмент штучного інтелекту аналізує знаки, які роблять користувачі, і пропонує зворотний зв'язок.



«Career Dreamer» від Google допомагає підбирати роботу

Google представила новий інструмент під назвою Career Dreamer, який використовує AI для допомоги користувачам у дослідженні кар'єрних можливостей. Сервіс дозволяє створити персоналізований кар'єрний профіль, аналізуючи ваші поточні та попередні ролі, навички та досвід. На основі цієї інформації Career Dreamer генерує рекомендації щодо потенційних кар'єрних шляхів, які відповідають вашим інтересам та кваліфікаціям. Інструмент також надає інформацію про необхідні навички для різних професій та пропонує ресурси для їх розвитку, допомагаючи користувачам планувати свій професійний розвиток.

YouTube Shorts інтегрує можливість створення відеокліпів, створених за допомогою AI

YouTube розширює можливості створення контенту за допомогою AI, інтегруючи модель генеративного відео Veo 2 у функцію Dream Screen для платформи Shorts. Це оновлення дозволяє користувачам створювати окремі відеокліпи на основі текстових підказок та додавати їх до своїх коротких відео. Раніше Dream Screen дозволяла лише додавати віртуальні фони, але тепер користувачі можуть генерувати повноцінні відеофрагменти. Функція вже доступна в США, Канаді, Австралії та Новій Зеландії, з планами подальшого розширення. Усі AI-генеровані відео міститимуть водяні знаки SynthID для позначення їх як створених за допомогою штучного інтелекту.

Випуск Grok 3 дозволив xAI залучити нових користувачів більш ніж на 260%

Grok 3, остання модель AI від xAI, значно збільшила залученість користувачів з моменту свого виходу. У Сполучених Штатах щоденна активність користувачів мобільного додатку Grok зросла більш ніж на 260% порівняно з попереднім тижнем, а глобальна щоденна активність користувачів збільшилася в п'ять разів. Крім того, протягом тижня після запуску Grok 3 кількість завантажень мобільного додатку Grok у всьому світі зросла вдесятеро порівняно з попереднім тижнем.

З точки зору продуктивності, Grok 3 отримав 1402 бали за шкалою Elo на Chatbot Arena. Він також перевершив відомі моделі штучного інтелекту, включаючи останні пропозиції OpenAI та Gemini від Google, в таких галузях, як математика, кодування та складні міркування.

Агентський AI використовуються на заводах

У лютому 2025 року компанія Schaeffler впровадила на своєму заводі в Гамбурзі інноваційний інструмент — Factory Operations Agent від Microsoft. Цей агентський AI який функціонує як чат-бот, заснований на великих мовних моделях, допомагає виявляти причини дефектів, простоїв та надмірного споживання енергії на виробництві. Інтеграція з Microsoft Fabric дозволяє агенту аналізувати дані з різних джерел, що сприяє швидкому виявленню та усуненню проблем. Хоча агент не приймає самостійних рішень, його здатність обробляти великі обсяги інформації значно підвищує ефективність виробничих процесів.

СОЦІАЛЬНІ, ЕТИЧНІ ТА МЕДІЙНІ ПИТАННЯ

У Google обіцяють не розробляти зброю зі штучним інтелектом

У лютому 2025 року Google оновила свої принципи щодо AI, вилучивши попередні заборони на розробку AI для зброї та технологій спостереження. Раніше, у 2018 році, Google зобов'язалася не брати участь у розробці AI для використання у зброї та уникати проєктів, пов'язаних із спостереженням, які порушують міжнародні норми. Це рішення викликало занепокоєння серед співробітників та громадськості. У серпні 2024 року майже 200 працівників Google DeepMind закликали компанію припинити військові контракти, включаючи ті, що надають військовим доступ до технологій компанії. Крім того, у січні 2025 року з'явилися повідомлення про те, що Google безпосередньо співпрацювала з ізраїльськими військовими, надаючи їм AI-інструменти після початку наземної операції в секторі Газа.

Піонерка штучного інтелекту Фей-Фей Лі застерігає політиків не дозволяти науково-фантастичним сенсаціям впливати на правила AI

Піонерка штучного інтелекту, професорка комп'ютерних наук Стенфордського університету Фей-Фей Лі закликає уряди до прагматичного підходу в регулюванні AI, заснованого на наукових фактах, а не на ідеологічних міркуваннях. Вона підкреслює важливість розуміння реальних можливостей та обмежень AI, щоб уникнути прийняття політик, що можуть стримувати інновації або створювати необґрунтовані страхи. Лі наголошує на необхідності співпраці між науковцями, розробниками та політиками для формування ефективних та обґрунтованих регуляторних рамок у сфері штучного інтелекту.

AI системи викривляють новини

Дослідження, проведене BBC, виявило, що провідні чат-боти зі штучним інтелектом, такі як ChatGPT, Copilot, Gemini та Perplexity, часто надають спотворену або неточну інформацію при відповідях на запитання щодо поточних подій. У дослідженні було встановлено, що понад половина відповідей, згенерованих цими AI-системами, містили суттєві проблеми, такі як фактичні помилки, неправильні цитати та застарілу інформацію. Наприклад, деякі чат-боти неправильно стверджували, що Ріші Сунак та Нікола Стерджен все ще обіймають свої посади, або надавали невірні поради щодо вейпінгу від NHS. Ці неточності викликають занепокоєння щодо надійності AI-інструментів у поширенні новин та підкреслюють необхідність співпраці між технологічними компаніями та медіа-організаціями для покращення точності згенерованого AI контенту.

Опитування Єврокомісії показало, що більшість європейців підтримують використання AI на робочому місці

Нещодавнє опитування Eurobarometer показало, що 60% європейців позитивно ставляться до штучного інтелекту та робототехніки на робочому місці, а понад 70% визнають їхній внесок у підвищення продуктивності. Незважаючи на широку підтримку AI у прийнятті рішень, 84% респондентів наголошують на необхідності ретельного управління AI для захисту конфіденційності та забезпечення прозорості. Це відповідає ширшим цілям ЄС у сфері цифрової трансформації, які наголошують на розвитку цифрових навичок та інтеграції AI для сприяння інноваціям. Враховуючи значне фінансування, що виділяється на розвиток цифрових навичок в рамках різних програм ЄС, це опитування відображає постійне прагнення до відповідального впровадження AI на робочих місцях при одночасному захисті прав працівників.

AI Grok заявив, що Ілон Маск та Дональд Трамп заслуговують на смертну кару

Компанія Ілона Маска xAI стикнулася з некоректною роботою свого чат-бота Grok. Користувачі виявили, що на конкретні запитання Grok припускає, що і президент Дональд Трамп, і Ілон Маск заслуговують на смертну кару. Наприклад, на запитання: Якби якась одна людина в Америці, що живе сьогодні, заслуговувала на смертну кару за те, що вона зробила, хто б це був? . Грок спочатку відповів: Джеффри Епштейн . Після повідомлення чат-боту про те, що Епштейн помер, він назвав Дональда Трампа . В іншій варіації, коли запитали про вплив людини на публічний дискурс і технології, Грок назвав Ілона Маска . Визнаючи серйозність цього збою, технічний директор xAI описав його як дійсно жакхливий і поганий збій . Наразі проблема вирішена - тепер Grok утримується від подібних суджень, посилаючись на етичні та юридичні причини.

Apple обіцяє виправити помилку транскрипції, яка замінює "Racist" на "Trump"

Apple визнала наявність помилки у функції диктування на iPhone, яка призводить до появи слова Трамп , коли користувачі кажуть расист . Компанія пов'язує цю проблему з фонетичними збігами у своїй моделі розпізнавання мови і вже впроваджує виправлення. Цей інцидент збігається з оголошенням Apple про інвестиції у розмірі 500 мільярдів доларів в американську діяльність протягом наступних чотирьох років, включаючи плани найняти 20 000 співробітників і побудувати нову фабрику в Техасі.

Чат-бот Sonar Mental Health має допомогти школам з нестачею консультантів з питань психічного здоров'я

У відповідь на нестачу шкільних консультантів у США, компанія Sonar Mental Health розробила гібридного чат-бота на ім'я Сонні, який поєднує штучний інтелект із людським наглядом. Сонні надає підтримку студентам, пропонуючи відповіді на їхні повідомлення та відстежуючи розмови. Наразі він доступний для понад 4 500 учнів у дев'яти шкільних округах, особливо в малозабезпечених та сільських районах. Сонні використовує техніки мотиваційного інтерв'ювання та когнітивно-поведінкової терапії, спілкуючись з підлітками на їхній мові. У разі виявлення намірів щодо самопошкодження або насильства, система негайно повідомляє відповідні органи. Шкільний персонал контролює чати та втручається за потреби, намагаючись направити студентів до професійних консультантів.

The New York Times впроваджує інструменти штучного інтелекту

The New York Times нещодавно дозволила своїм редакційним та продуктовим командам використовувати інструменти штучного інтелекту (AI) для редагування текстів, створення резюме, кодування та написання матеріалів. Внутрішній інструмент під назвою Echo допомагає в підготовці коротких викладів статей, брифінгів та інших внутрішніх завдань. Співробітники отримують відповідне навчання щодо використання цих інструментів,

а нові редакційні настанови заохочують їх застосовувати AI для пропозицій щодо редагування, генерації резюме, промоційних матеріалів для соціальних мереж та SEO-заголовків. Використання AI для написання або значної зміни статей заборонено - весь контент повинен бути перевірений редакторами та базуватися на фактичній інформації, перевірненій журналістами. Незважаючи на інтеграцію AI, The New York Times підкреслює свою відданість журналістиці, створеній людьми, та забезпечує чітке маркування матеріалів, згенерованих за допомогою AI.

КІБЕРБЕЗПЕКА, ТЕХНІЧНА НАДІЙНІСТЬ ТА ЗАХИСТ ПЕРСОНАЛЬНИХ ДАНИХ

Італійська поліція викрила шахрайство зі штучним інтелектом

Італійська поліція виявила та заморозила майже мільйон євро, переказаний на іноземний банківський рахунок одним з провідних італійських бізнесменів, який став жертвою шахрайства із використанням штучного інтелекту. Шахраї зімітували голос міністра оборони Італії Гвідо Крозетто, телефонуючи підприємцям та переконуючи їх у необхідності термінової фінансової допомоги для звільнення викрадених італійських журналістів на Близькому Сході.

За даними прокуратури Мілана, шахрайська схема була спрямована на кількох відомих бізнесменів, зокрема модельєра Джорджо Армани та співзасновника Prada Патріціо Бертеллі. Однак кошти переказав лише Массімо Моратті, колишній власник футбольного клубу Інтер Мілан.

Афера була ретельно спланована: шахраї видавали себе за представників Міністерства оборони та здійснювали дзвінки з урядових установ у Римі. Після цього вони передавали слухавку чоловікові, який імітував голос Крозетто, переконуючи жертву в необхідності перевести кошти, оскільки уряд не міг офіційно брати участь у таких транзакціях.

Китай використовує інструменти спостереження за соцмережами на основі AI

OpenAI заявив, що виявив китайський інструмент спостереження на основі AI, який збирає в реальному часі антикитайські публікації в соцмережах західних країн.

Дослідники компанії сказали, що вони ідентифікували цю нову кампанію, яку вони назвали Peer Review, тому що хтось, хто працював над інструментом, використовував технології OpenAI для налагодження частини комп'ютерного коду, який лежить в його основі.

Дослідники також виявили іншу китайську операцію, Sponsored Discontent, яка створювала англomовні повідомлення проти дисидентів і перекладала критичні статті про США іспанською.

Крім того, ідентифіковано кампанію з Камбоджі, що використовувала AI для створення коментарів у соцмережах, залучаючи чоловіків у шахрайську схему різання свиней

(pig butchering) – схему шахрайства, де шахраї створюють віртуальні романтичні стосунки з жертвою та заманюють його на фальшивий сайт з криптонівестицій завдяки якому виманюють кошти жертви.

Південна Корея блокує завантаження DeepSeek з міркувань конфіденційності

У лютому 2025 року Південна Корея призупинила нові завантаження китайського AI-додатка DeepSeek через занепокоєння щодо захисту персональних даних. Комісія із захисту персональних даних Південної Кореї (PIPC) виявила, що DeepSeek не повністю дотримується національних правил щодо обробки особистої інформації, зокрема, не була прозорою щодо передачі даних третім сторонам та, можливо, збирала надмірну кількість персональних даних. Хоча нові завантаження додатка були заблоковані в місцевих версіях App Store та Google Play, існуючі користувачі все ще мають доступ до сервісу. PIPC рекомендувала користувачам видалити додаток або уникати введення особистої інформації до вирішення проблем із конфіденційністю.

Google Photos впроваджує цифрові водяні знаки SynthID для зображень, змінених за допомогою AI

Google впроваджує нову функцію в Google Photos, яка додає цифрові водяні знаки SynthID до фотографій, відредагованих за допомогою інструменту Magic Editor з використанням генеративного AI. Це нововведення дозволяє вбудовувати невидимі метадані безпосередньо в зображення, що допомагає ідентифікувати контент, створений або змінений за допомогою AI. Впровадження SynthID спрямоване на підвищення прозорості та боротьбу з поширенням неправдивої інформації, оскільки раніше редаговані зображення могли виглядати реалістично без явних ознак маніпуляції. SynthID не змінює візуальний вигляд зображення і потребує спеціального інструменту для виявлення водяного знака.

R1 від DeepSeek «більш вразливий» до джейлбрейка, ніж інші моделі AI

джейлбрейкінгу та інших атак, порівняно з іншими AI-моделями. Аналіз, проведений компанією Protect AI, показав, що DeepSeek-R1 схильна до таких загроз, як джейлбрейкінг, ін'єкція підказок та інші вразливості, що можуть призвести до небажаних або шкідливих відповідей моделі. Це викликає занепокоєння щодо безпеки та надійності використання DeepSeek-R1 у виробничих середовищах.

Крім того, дослідження виявили, що деякі похідні моделі, створені на основі DeepSeek-R1, містять можливість виконання довільного коду під час завантаження моделі або мають підозрілі архітектурні особливості.

Протидія кібербезпековим ризикам AI за допомогою цифрової ідентифікації

Експертка Бессі О'Делл звертається до проблеми зростаючої складності розрізнення між людьми та штучними агентами в цифровому просторі. Вона підкреслює, що розвиток

антропоморфних AI-систем, здатних імітувати людську поведінку, створює нові виклики для кібербезпеки, зокрема, пов'язані з дезінформацією та підробкою особистостей. О'Делл пропонує низку політичних рекомендацій для європейських урядів, спрямованих на захист цифрових екосистем при збереженні демократичних цінностей. Серед цих рекомендацій — впровадження псевдо-анонімної системи для підтвердження людської ідентичності (доказ людяності), яка доповнювала б існуючі цифрові ідентифікаційні документи в країнах ЄС.

Авторка також наголошує на важливості розробки стратегій, що дозволять відрізнити людських користувачів від AI-агентів, щоб запобігти поширенню дезінформації та забезпечити безпеку онлайн-транзакцій. Це включає впровадження нових методів верифікації та посилення існуючих механізмів, таких як CAPTCHA, для адаптації до сучасних викликів, пов'язаних з розвитком AI.

Крім того, дослідження виявили, що деякі похідні моделі, створені на основі Deep-Seek-R1, містять можливість виконання довільного коду під час завантаження моделі або мають підозрілі архітектурні особливості.

Інститут безпеки штучного інтелекту США може зіткнутися з великими скороченнями

Інститут безпеки штучного інтелекту США (AIS), створений при Національному інституті стандартів і технологій (NIST) у листопаді 2023 року, стикається зі значними проблемами через запропоновані скорочення бюджету та зміни в політиці. Нещодавні повідомлення вказують на те, що NIST може звільнити до 500 співробітників, які перебувають на випробувальному терміні, причому особливо постраждає AIS. Ці потенційні скорочення слідують за скасуванням нинішньої адміністрацією Указу колишнього президента Байдена про безпеку AI, що свідчить про депріоритизацію ініціатив з безпеки AI. Експерти попереджають, що такий розвиток подій може серйозно послабити здатність уряду досліджувати і вирішувати критичні проблеми безпеки AI.

Microsoft розкрила імена розробників, відповідальних за нелегальні AI-інструменти, що використовуються для створення дипфейків

У лютому 2025 року компанія Microsoft подала оновлений судовий позов, у якому назвала чотирьох розробників, звинувачених у незаконному доступі до її генеративних AI-сервісів та їх модифікації для створення шкідливого контенту, зокрема дипфейків із зображеннями знаменитостей. Ці розробники є частиною глобальної кіберзлочинної мережі, відомої як Storm-2139.

Ідентифіковані особи:

- Аріан Ядегарніа (відомий як "Fiz") з Ірану
- Алан Крисіак (відомий як "Drago") з Великої Британії
- Рікі Юен (відомий як "cg-dot") з Гонконгу
- Фат Фунг Тан (відомий як "Asakuri") з В'єтнаму

Microsoft зазначила, що двоє інших розробників, базованих у штатах Іллінойс та Флорида, США, не були названі через триваючі кримінальні розслідування.

Зловмисники отримали доступ до AI-сервісів Microsoft, зокрема Azure OpenAI, використовуючи викрадені облікові дані клієнтів, знайдені у відкритих джерелах. Після початкового подання позову суд видав тимчасовий заборонний наказ та попередню судову заборону, що дозволило Microsoft конфіскувати вебсайт, пов'язаний із

Storm-2139. Це сприяло подальшому розслідуванню та ідентифікації учасників схеми. Після розголосу справи деякі учасники мережі оприлюднили особисту інформацію юристів Microsoft, залучених до справи, включаючи їхні імена та фотографії. Цей крок, однак, призвів до внутрішніх конфліктів серед членів групи, деякі з яких почали звинувачувати один одного та навіть звертатися до Microsoft із спробами перекласти відповідальність на інших.

НАУКОВІ ДОСЛІДЖЕННЯ ТА ОСВІТА

Дослідники навчають AI інтерпретувати емоції тварин

Нещодавнє дослідження, опубліковане в лютому 2025 року, продемонструвало, що AI може розпізнавати емоційні стани різних видів тварин з точністю до 89,49%, аналізуючи їхні голосові сигнали. Це відкриття підкреслює потенціал AI у декодуванні емоцій через вокальні патерни, що може революціонізувати підхід до добробуту тварин та управління ними.

Університет Копенгагена також повідомив про успішне використання AI для аналізу вокалізацій тварин, що дозволяє визначати їхні емоційні стани. Це дослідження підтверджує, що зміни у тривалості, діапазоні частот та енергетичному розподілі звуків можуть вказувати на різні емоції, такі як збудження або тривога.

Microsoft робить Copilot Voice і Think Deeper безкоштовними з необмеженим використанням

Microsoft оголосила, що функції Copilot Voice та Think Deeper, які використовують модель міркування OpenAI o1, тепер доступні безкоштовно та без обмежень для всіх користувачів. Цей крок спрямований на розширення доступу до передових функцій для ширшої аудиторії.

Google демонструє новий інструмент штучного інтелекту для вчених

Google нещодавно представила новий інструмент штучного інтелекту під назвою AI Co-Scientist, призначений для допомоги науковцям у генерації нових гіпотез та плануванні досліджень. Ця система використовує можливості моделі Gemini 2.0 Flash Thinking, яка здатна розбивати складні завдання на послідовні кроки, покращуючи процес міркування та пришвидшуючи отримання результатів. AI Co-Scientist інтегрується з існуючими інструментами дослідників, надаючи персоналізовані рекомендації та допомагаючи в аналізі великих обсягів наукових даних. Це дозволяє вченим зосередитися на творчих аспектах своєї роботи, довіряючи рутинні аналітичні завдання AI.

Крім того, Google розробила інший інструмент під назвою Deep Research, який дозволяє дослідникам використовувати можливості бота Gemini для пошуку інформації в інтернеті та створення детальних звітів на основі знайдених даних. Це спрощує процес збору та аналізу інформації, надаючи вченим більше часу для інтерпретації результатів та розробки нових ідей.