

ДОБРОЧЕСНІСТЬ ВИКОРИСТАННЯ ШІ НЕВІДДІЛЬНА ВІД ДОБРОЧЕСНОСТІ ЗАПИТАНЬ



У блозі Міжнародного центру академічної доброчесності (ICAI) опублікована стаття Ерл Абрахамсон «Генеративний ШІ як сучасний Сократ: чому чесність починається з більш правильних питань».

У ній автор пропонує нетривіальну точку зору на проблему використання ШІ: можливо, найбільша цінність генеративного штучного інтелекту не в тому, що він може відповісти на запитання, а в тому, що він виявляє, наскільки поганими були поставлені йому запитання? Наразі значна більшість дискусій щодо ГШІ зосереджена на питаннях ефективності, швидкості отримання результатів, оперативних рішеннях та експертизі. У контексті принципу «Цілісність має значення» ця дискусія є неповною. Необхідно говорити про відповідальність та судження. Коли люди поспішають з відповідями, вони часто ігнорують ці базові цінності. ГШІ гостро висвітлює проблему, відповідаючи, навіть коли запитання погано сформульовані, етично некоректні, поверхневі або концептуально слабкі. Запитання не є нейтральними: вони несуть припущення, визначають пріоритети та цінності. Коли питання вузьке, воно може спотворити

відповідь, коли його формулюють поспіхом, воно може нести шкоду, коли воно зосереджене лише на оптимізації, воно може ігнорувати наслідки. Тому цілісність у використанні ШІ починається не з результатів, а з аналізу того, що насправді ми запитуємо. Розглядаючи, кому вигідно, а кому може зашкодити те, як сформульоване запитання, Ерл говорить у першу чергу про освіту, управління дослідженнями, журналістику та державну політику, де питання формують результати з реальними наслідками. Генеративний штучний інтелект може допомогти у дослідженні, але лише за умови, що дослідження є цілісним. Вивчення співпраці людини та ШІ засвідчує, що найбільші переваги виникають тоді, коли штучний інтелект використовується для переосмислення проблем, а не для їх вирішення. Цілісність покращується, коли питання перевіряються на те, які відповіді вважаються достовірними. Такий підхід протистоїть упередженості автоматизації та підтримує людське судження, а не замінює його. Головний етичний ризик ГШІ полягає не в тому, що він замінить людське мислення, а в тому, що він замаскує його відсутність. Коли користувачі покладаються на ШІ для генерування відповідей, не ставлячи під сумнів питання, вони послаблюють власне судження. Коли вони використовують його для перевірки припущень, дослідження наслідків та оскарження формулювань, вони його посилюють. Тому чесність у використанні ШІ невіддільна від чесності постановки запитань. Насправді цінним запитанням є не те, на яке ШІ дає швидку і чітку відповідь, а те, яке покращує якість людських міркувань, прийняття рішень та етичну усвідомленість. Дискомфорт, невизначеність та складність – це не недоліки. Вони є сигналами того, що до чесності ставляться серйозно. ГШІ не має сприйматись як авторитет, – він віддзеркалює якість дослідження: поверхневі відповіді є результатом поверхневих питань, які їх породили.

Детальніше: https://www.academicintegrity.org/aws/ICAI/pt/sd/news_article/612514/_PARENT/layout_details/false

Фото: скріншот

#НРАТ_Усі_новини #НРАТ_ШтучнийІнтелект #НРАТ_АкадемДоброчесність
#НРАТ_Науковцям_новини #НРАТ_Освітянам_новини

2026-02-25

Інформація з офіційного вебпорталу Національного репозитарію академічних текстів