

Катерина Колєватова,
студентка філологічного факультету,
Національний університет
«Чернігівський колегіум» імені Т.Г. Шевченка,
м. Чернігів, Україна

Науковий керівник:
Марина Ольховик
кандидат філософських наук, доцент,
доцент кафедри філософії та культурології,
Національний університет
«Чернігівський колегіум» імені Т.Г. Шевченка,
м. Чернігів, Україна

ФІЛОСОФСЬКІ ПРОБЛЕМИ ШТУЧНОГО ІНТЕЛЕКТУ

Штучний інтелект змінює уявлення про розум, осмислення та етичні межі технологій. Завдяки стрімкому розвитку технологій штучного інтелекту перед людством постають глибокі філософські виклики. Одним із головних викликів філософії штучного інтелекту є питання свідомості: Чи здатна машина володіти усвідомленням, чи це лише імітація? Штучний інтелект вплинув не лише на технологічний процес, але й на такі філософські категорії як «буття», «свідомість», «розум», «розвиток» та інше.

Проблема штучного інтелекту охопила широке коло досліджень в науці. Серед головних напрямів – створення аналогу інтелекту людини, розробка «суперінтелекту», проєктування структури психіки, аналіз впливу штучного інтелекту на суспільство. Процес втілення в життя нових результатів досліджень в галузі психології, неврології, кібернетики потребує подальшого самоаналізу наявних знань, розуміння стану проблеми та формування прогнозів на майбутнє.

Історія розвитку наукового напрямку вбачає вирішення ряду технічних і філософських проблем. До обговорення цього питання долучалися не тільки представники технічних і природничих наук, а також філософи, психологи та соціологи. Переважна частина філософів і дослідників, таких як Дж. Серль, Д. Деннет, Дж. Лукас, Х. Дрейфус, Р. Пенроуз дотримувалися критики теоретичної бази дисципліни та наголошували на обмеженнях, які відповідають створенню штучного інтелекту. Але переважна більшість науковців та дослідників, як-от: С.А. Клейн, Д. Чалмерс, А. Тюрінг, спростовують обґрунтування колег.

Уперше у 1955 році визначення «штучний інтелект» запровадив доцент математики Датмутського коледжу Джон МакКарті. Вчений прагнув відокремити цю галузь досліджень від кібернетики, піднявши питання під час семінару влітку 1956 року. Саме на цьому заході було вперше використано визначення «штучний інтелект». Джон МакКарті шукав новий термін, який міг би об'єднати окремі дослідження в єдине, де винайдення будуть зосереджені на «розумних машинах», які мають відтворювати людський інтелект.

У статті в газеті «The Guardian» британський фізик-теоретик Девід Дойч наголошує на тому, що філософський підхід може допомогти програмістам і нейрофізіологам в створенні штучного інтелекту [4]. Професор вважає, що штучний інтелект несе ряд фундаментальних помилок. Одна із них, недооцінення штучного інтелекту. Дехто вважає, що він не буде «розумнішим» за програмне забезпечення. До того ж суспільство завжди хоче бути найідеальнішим в світі [3].

Вважається, що основною філософською проблемою в області штучного інтелекту являється можливість або неможливість моделювання мислення людини. Чи має бути штучний інтелект схожий на людський? Відповідь на це питання отримано з точки зору нейрофізіології. Створення двійника людського мозку вказує на те, що суттєві відмінності існують. Головне в результатах функціонування штучного інтелекту – він повинен мати певні обмеження, правила близькі до людських. Одне із головних правил – штучний інтелект повинен бути безпечним для людей, без перешкод надавати можливість усунути помилки. Не повинен обмежувати приватність людини, свободу волі. Свою думку про небезпеку розумних машин висловлював британський фізик-теоретик Стівен Гокінг в одному із інтерв'ю BBC. Він наголошував, що із розвитком штучного інтелекту може прийти кінець людської раси [2].

Штучний інтелект не має використовуватись на користь озброєння, підривати громадянські процеси, здоров'я людини. Дослідники відмічають надмірну довіру штучному інтелекту серед людства, вважаючи його надто розумним. Перекладення своїх обов'язків на штучні системи, можуть бути більш небезпечними для майбутнього покоління. Тому, щоб не допустити таких помилок з боку штучного інтелекту, над ним потрібен контроль [3]. Але спроба контролю людини над штучним інтелектом може сприйматися ним негативно, можливий процес відчуження. Людська бездумність та легковажність стають причиною небезпеки, яка походить від штучного інтелекту.

Слід відмітити, що штучний інтелект не спроможний робити винаходи, змінювати, покращувати, він спроможний приймати рішення на підставі прикладів, які були засвоєні під час навчання. Він не спроможний розуміти, адже спирається лише на тексти та мовні програми. На відміну від людини, він не має контакту з навколишнім середовищем, не має змоги отримати інформацію з різних джерел. Тому перед вченими постає нове завдання, яке буде полягати у створенні пристроїв, що будуть виявляти причину-наслідкові зв'язки, та приймати сенсорну та лінгвістичну інформацію [1].

Отже, дослідження штучного інтелекту показало, що це міждисциплінарний напрям, який охопив зв'язки філософії, комп'ютерних наук, психології, фізіології. Суспільству слід пам'ятати, що докладаючи багато зусиль до розвитку штучного інтелекту, необхідно застерегти себе від нелюдських вчинків та мислення.

Література

1. Бондар С.В. Обмеження та перспективи розвитку штучного інтелекту. Філософський аналіз. URL: <https://ir.kneu.edu.ua/server/api/core/bitstreams/1770a446-cdd6-4e29-86df-d0c29a59741e/content>
2. Головащенко І.О. Проблема штучного інтелекту в сучасній філософії / [Електронний ресурс] / І.О. Головащенко, Ю.П. Носковенко. Матеріали XL VIII науково-технічної конференції підрозділів ВНТ, Вінниця, 13-15 березня 2019 р. Електрон. Текст. дані. 2019. URL: <https://conferences.vntu.edu.ua/index.php/all-hum/all-hum-2019/paper/view/7596>
3. Зелінська Д.О. Філософія як ключ до створення штучного інтелекту. URL: <https://ir.lib.vntu.edu.ua/bitstream/handle/123456789/20448/4144.pdf?sequence=3&isAllowed=y>
4. Романчук О.К. Штучний інтелект як «експериментальна філософія» чи вимога реальності. URL: <https://universum.lviv.ua/magazines/universum/2016/6/shtuch-int.html>