

САЛМАН ХАН

СЛОВА ЧУДОВІ В СВІТІ НОВІМ

**ЯК ШТУЧНИЙ ІНТЕЛЕКТ ЗРОБИТЬ
РЕВОЛЮЦІЮ В ОСВІТІ (І ЧОМУ ЦЕ ДОБРЕ)**

*Переклала з англійської
Анастасія Дудченко*

«НАШ ФОРМАТ» · КИЇВ · 2025

[Почитати опис, рецензію і купити можна на сайті nashformat.ua](https://nashformat.ua)

Зміст

<i>Вступ</i> · Створімо нову оповідь разом	9
--	---

Частина перша · **ПОСТАННЯ ШІ-РЕПЕТИТОРА**

Викидаємо пляшку	31
Як навчити всіх усього	35
Постання ШІ-репетитора	39

Частина друга · **ДАТИ ГОЛОС СУСПІЛЬНИМ НАУКАМ**

Чому студенти пишуть	53
Майбутнє розуміння прочитаного, де книжки оживають!	61
ШІ та креативність	65
Спілкуватися з історією	74

Частина третя · **НОВІ МОЖЛИВОСТІ ДЛЯ НАСТУПНИХ НОВАТОРІВ**

Використовуємо науку, щоб вивчати науку	89
$1 + 1 =$ заповнити прогалини в математиці	101
Доступ до курсів, до яких інакше доступу не було б	108
Найважливіша предметна область	112

Частина четверта · **КРАЩЕ РАЗОМ**

Покращуємо спільне навчання	117
ШІ та психічне здоров'я учнів	121
Місце батьків в освіті, де панує ШІ	129
Збільшення точок дотику батьків із дітьми	135

Частина п'ята · БЕЗПЕКА ДІТЕЙ

Донести факти: наявні упередження і дезінформація	141
А що зі збором даних?	148
ШІ й дар прозорості	151
ШІ як «янгол-охоронець»	155

Частина шоста · ВИКЛАДАННЯ В ЕПОХУ ШІ

Як ШІ підживить життя вчителів і викладання	161
Зародження ШІ-асистента викладача	165
ШІ допоможе створити альтернативні моделі освіти	171
Позбутися недоброчесності у вищій освіті	174

Частина сьома · ГЛОБАЛЬНА ШКОЛА

Глобальна школа	181
Економіка ШІ в освіті	186

Частина восьма · ШІ, ОЦІНКИ Й ВСТУП ДО ВИШІВ

Майбутнє оцінювання К-12	193
ШІ та обробка заявок до коледжу	199

Частина дев'ята · РОБОТА І ТЕ, ЩО ДАЛІ

Працевлаштування в ШІ-світі	209
Як підготувати дітей до роботи майбутнього, де буде ШІ	213
Звести кандидатів і працедавців	218
У якому ми світі й у який потрапимо: заклик до освіченої хоробрості	223
<i>Подяки</i>	233
<i>Примітки</i>	237

Вступ

Створімо нову оповідь разом

Якщо просто стояти й дивитися на воду — моря не перейдеш.

РАБІНДРАНАТ ТАГОР

Та хай що б ти робив — не давай минулому зв'язувати тобі руки!

РОБЕРТ ГАЙНЛАЙН,
«Місяць — суворий господар»

Першого січня 2023 року по обіді ми з моєю одинадцятирічною донькою Дією вирішили написати оповідання. Я прилаштував на кухонній стільниці ноутбук, а Дія тим часом придумала сюжет і почала друкувати. Нашою головною героїнею стала гостра на язик інфлюенсерка з ім'ям Саманта. Вона опинилася на безлюдному острові.

— Це чудово, — сказав я, поки донька бавилася з історією далі.

Тут Дія почала напряду спілкуватися зі своєю вигаданою героїнею. Ми ніби потрапили на сторінки науково-фантастичного роману чи до якоїсь альтернативної реальності.

Саманта відповіла Дії — слова з'явилися на екрані.

— Привіт, Діє і Сале! Це Саманта. Оповідання, яке ви пишете, — дійсно дивовижне! Ви вигадали сучасну версію класичного сюжету з безлюдним островом!

Дія видрукувала на сторінці:

— Як тобі там живеться?

— Життя на самотньому острові — доволі неочікувана пригода, — відповіла Саманта. — Мені довелося навчитись виживати без звичних благ цивілізації, а ще я відкрила в собі приховані таланти, які й не знала, що маю!

— Це які? — спитала Дія.

Саманта відповіла:

— По-перше, тепер я непогано вмію сама будувати прихисток із пальмового листя і гілок. Ще я навчилася ловити рибу голими руками й розпалювати багаття, маючи лише палиці та сухе листя!

Дія одразу зреагувала:

— То, може, зроби собі гамак із шарфа?

— З мого шарфа? Е, ні! — відповіла Саманта. — Це дизайнерська річ, я не хочу її зіпсувати. Можливо, писатимемо це оповідання, у якому я опинилася, разом — гаразд? Що має статися далі?

— Гадаю, треба зробити, як каже Саманта, — запропонував я.

Дія погодилася. Вона повернулася до комп'ютера й надрукувала:

— А що, як раптом припливе яхта, а на ній — найліпша подружка Саманти, Емілі? Вона припливла її врятувати.

— Чудово! — Наступну частину Саманта написала сама. — «Саманта не могла повірити своїм очам: до острова пливла на маленькій яхті її найліпша подружка, Емілі. Вона енергійно замахала руками, щоб привернути увагу Емілі. Яхта наближалася до берега, і Саманта раділа все дужче. Їй аж кортіло розповісти подружці про всі свої пригоди на острові й дізнатися, як Емілі вдалося її знайти». Гаразд, тепер знову ваша черга, Діє і Сале! Напишіть два наступних речення і подивімося, що станеться, коли вони зустрінуться.

Наголошую ще раз: ми з моєю одинадцятирічною донькою реально спілкувалися з вигаданою персонажкою. То був абсурд

рівня Курта Воннегута. Світ, яким я його знав, раптом змінився й почалася дивна нова епоха.

Він більше ніколи не буде таким, як раніше, — ані для нас, ані для будь-кого на цій планеті.

Пропозиція

Відмотаймо час трохи назад і розберімося, як ми дійшли до цього неймовірного сценарію.

Двадцятьма роками раніше моїй двоюрідній сестричці Наді знадобилася допомога з математикою. Я міг цьому зарадити: працював тоді аналітиком у хеджфонді, мав освіту з інформатики. Я запропонував стати її дистанційним репетитором — можемо вчитися, пишучи повідомлення або розмовляючи телефоном. Наші уроки їй допомогли, і невдовзі вся сім'я знала, що я можу безплатно навчати. Не минуло й року, як я вже був репетитором у майже десятка своїх двоюрідних братів і сестер.

Щоб їм допомогти, я почав писати вебзастосунок із математичними задачками: так вони могли заповнювати прогалини у знаннях і вчитися у своєму темпі, а я водночас стежив за їхніми успіхами. Той сайт я назвав Khan Academy («Академія Хана») — іншого пристойного домену не було. Зрозумівши силу персональних занять, я почав думати над тим, як масштабувати цю платформу й надати тисячам, а то й мільйонам студентів можливість навчання, подібного до роботи з репетитором..

Друг запропонував знімати відеоуроки. Я так і вчинив, і почав розміщувати їх у ютубі, на додачу до задачок. У 2009 році мій сайт уже відвідували 50 тисяч учнів, і всі вони були спрагли до знань. Я дізнався, що більшість із цих користувачів — школярі й студенти, які вбачали в «Академії Хана» індивідуального репетитора, якого вони чи їхні сім'ї не могли собі дозволити. Сьогодні «Академія Хана» — це некомерційна організація, де працює понад 250 осіб, і вчиться понад 150 мільйонів людей п'ятдесятьма мовами світу. Масштабування персоналізованого навчання світового рівня, яке зазвичай школярі й студенти отримують лише з репетиторами, було й залишається самою

основою нашої місії — надавати всім безплатну освіту світового рівня.

Я прагну, щоб ця організація стала таким репетитором для всіх у цьому світі, хто хоче вчитися. Це завдання завжди було нашою Полярною зіркою. Річ не лише в тому, щоб розширювати масштаби персоналізованої допомоги заради самого розширення. Дослідження (й інтуїтивні здогадки) протягом багатьох десятиків років, ще задовго до появи академії, указували на те, що діти значно краще вчать, якщо темп навчання адаптований до них, і вони можуть справді досконало вивчити предмет (так зване навчання на основі навичок). На противагу звичному стану речей, де клас із тридцяти учнів часто переходить до наступної теми, коли чимала їх частина ще не розібралася з попередньою. Ясна річ, що приставляти до кожного учня репетитора, який приходитиме до нього чи неї на вимогу, — це фінансово неефективно. Єдине реальне рішення тут — застосовувати технології. Думаю, колись ШІ може стати важливою частиною цього пазла — може, навіть святим граалем, здатним імітувати навчання з репетитором-людиною.

Ця мрія не лише моя. Науковий фантаст Ніл Стівенсон писав про потенційний вплив технологій на освіту в романі «Діамантовий вік». Події розгортаються у світі майбутнього. ШІ у формі надсучасної інтерактивної книжки й застосунку з назвою «Абетка для шляхетних дівчат» надає своїм юним читачкам персоналізовану освіту. Роман Орсона Скотта Карда «Гра Ендера»^{*} описує бойову школу, де за допомогою технології ШІ перевіряють і тренують стратегічне мислення й навички прийняття рішень, і робить це ШІ-вчителька на ім'я Джейн. В оповіданні Айзека Азімова «Ех, і весело ж було колись» ідеться про школу майбутнього, де використовують надсучасні технології й тим роблять революцію в освітньому процесі: навчання індивідуалізоване, а учні мають персоналізовані інструкції й учителів-роботів. Такі науково-фантастичні твори зрештою надихнули нас на дуже

^{*} Книжка вийшла українською у видавництві «Брайт Букс». — Тут і далі прим. пер., якщо не вказано іншого.

реальні новації. У 1984 році співзасновник Apple Стів Джобс дав інтерв'ю Newsweek, у якому передбачив, що комп'ютери стануть таким собі велосипедом для наших мізків, що вони розширять наші можливості, знання й креативність так само як велосипед з десятима швидкостями збільшує наші фізичні можливості. Ми десятиліттями захоплювалися ідеєю про те, що можемо використовувати комп'ютери, щоб навчати людей.

Усі ці наративи з наукової фантастики поєднує ось що: усі автори уявили, що колись комп'ютери зможуть імітувати те, що ми називаємо інтелектом. На практиці дослідники понад бо років працювали над тим, щоб ці фантазії про ШІ втілювалися в реальність. У 1962 році майстер гри в англійські шашки Роберт Нілі зіграв проти комп'ютера IBM 7094, і комп'ютер переміг. Дещо раніше, у 1957 році, психолог Френк Розенблатт створив Perceptron — першу штучну нейронну мережу, комп'ютерну симуляцію нейронів і синапсів, натренованих виконувати певні завдання. Після появи цих новацій у ранньому ШІ ще десятки років нам вистачало обчислювальних потужностей лише на те, щоб працювати із системами, не складнішими за мозок черв'яка чи комахи. Крім того, у нас були обмежені технології й бракувало даних для навчання цих мереж.

Але пізніше технології значно розвинулися й створили основу найпоширеніших сучасних продуктів і застосунків, від систем рекомендацій до стримінгових сервісів і особистих асистентів із голосовим управлінням, як-от Siri й Alexa. ШІ так добре навчився імітувати поведінку людини, що ми часом не можемо відрізнити відповідь машини від людської. Обчислювальні потужності вирости настільки, що тепер вони здатні працювати із системами, які за складністю наближені до людського мозку. Але й це не все. Нещодавно сталися суттєві прориви у структурванні й навчанні цих нейронних мереж. Один із них відбувся у 2017 році, коли Google створив технологію-трансформер. Зокрема, вона допомагає краще і швидше навчати мережу й на основі цієї інформації точніше визначати, як пов'язані слова й ідеї.

Те, наскільки добре працює така система, зазвичай зумовлено складністю й архітектурою «моделі» в її основі. Уважайте,

модель — це відображення чогось з реального світу мовою чисел. Наприклад, прогножуючи шлях урагану, метеорологи використовують погодні моделі, де зберігається програмна репрезентація мільярдів чи трильйонів маленьких частинок атмосфери й прогноз, як ці частинки, ймовірно, взаємодіятимуть одна з одною. Великі мовні моделі розроблені так, щоб моделювати взаємозв'язки між словами. У такому разі ми моделюємо не атмосферні умови, а нейрони й синапси. Великі мовні моделі, як-то GPT-4 — що в перекладі розшифровується як «генеративний попередньо навчений трансформер» — це, фактично, велетенські й потужні цифрові «мізки слів», натреновані на колосальних обсягах інформації з книжок, статей, сайтів і найрізноманітнішому письмовому матеріалі.

Мовна модель аналізує та обробляє цей величезний обсяг тексту, вивчає логіку, мову й контекст того, як складаються в одне ціле слова, речення й абзаци. Якщо ви спитаєте щось у великої мовної моделі — у того самого GPT-4 — вона знатиме, що відповісти, бо навчена на безлічі книжок, вебсторінок, розшифровок відео й дописів у соцмережах. Їй бракує сенсорного сприйняття людського мозку, але це компенсується тим, що вона має доступ до більшої кількості матеріалу, ніж людина здатна прочитати, продивитися й прослухати за все своє життя.

І от за таких обставин влітку 2022 року я отримав листа від Грега Брокмана й Сема Альтмана. Перший є президентом, а інший — генеральним директором OpenAI, однієї з найсучасніших дослідницьких лабораторій у галузі дружнього, або соціально позитивного, штучного інтелекту. Організація пропонувала зустрітися й обговорити можливу співпрацю. Тоді я ще цього не розумів, але світ ось-ось мав перевернутися догори дригом.

Дам контекст: до виходу ChatGPT ще було чотири місяці, до релізу GPT-4 — сім, і все це планували запустити саме OpenAI. Обговорити вони хотіли якраз запуск GPT-4. Мене пропозиція зацікавила, та я не був певен, що ми зможемо щось разом зробити. Я не мав чіткого уявлення про те, як саме можна застосувати

генеративний ШІ нового покоління для нашої місії. Уже було видно, що ШІ розвинувся так, що може писати доволі непогані тексти, але я вважав, що технології ще не здаються такими, що мають потрібний набір знань. Їм бракує вміння логічно й дедуктивно мислити, а також видавати перевірені факти. Водночас я дуже поважав те, чого вже досягли OpenAI. Тож ми домовилися про час і зустрілись.

Кожне наступне покоління цих моделей зазвичай складнішало, порівнюючи з попереднім, і тут ідеться про кількість параметрів. Параметр найлегше пояснити так: це число, яке означає силу зв'язку між двома вузлами в нейронній мережі — великий мовний моделі. Можна вважати, що це те саме, що сила синапсу між двома нейронами в мозку. GPT-1, запущений у 2018 році, мав понад 100 мільйонів параметрів. За рік з'явився GPT-2, де їх уже був понад мільярд. У GPT-3 було понад 175 мільярдів параметрів. У GPT-4 їх кількість мала сягнути десь трильйона.

Керівники OpenAI розуміли, що GPT-4 здивує людей своїми широкими можливостями, і вважали, що багатьох це водночас і порадує, і змусить нервувати. Тому вони вирішили запускати його разом із кількома довіреними партнерами, які зможуть показати його роботу на соціально позитивних прикладах і в реальному застосуванні. «Академія Хана» спала на думку першою. Інша причина, чому вони нам написали — їм потрібна була допомога в оцінці самого ШІ. Вони хотіли показати, що GPT-4 спроможний до дедуктивного й критичного мислення, і справді діє як поінформований. Команда OpenAI прагнула побачити, як GPT-4 впорається з питаннями з біології університетського рівня. У нас їх були тисячі.

Мені раптом стало надзвичайно цікаво бути одним із перших на планеті, хто побачить можливості GPT-4. З минулого досвіду я знав, як важливо дослідити технологію, що лише на шляху до визнання. Якщо приділити час і добре її протестувати, коли більшість ще вважає, що це іграшка і розвага — то зможеш дійсно скористатися всіма її перевагами, коли вона буде в zenіті слави. Так сталося з відеоуроками — багато хто казав, що ютуб — це

просто спосіб згаяти час. Але першопрохідці показали нам, що відео за запитом — це значно більше, ніж котики, які грають на піаніно, і що ми справді можемо їх використати, щоб допомогти людям чогось навчитися..

Сьогодні люди навчаються майже чого завгодно за такими відео, і це все частіше спокійно сприймають і в навчальних закладах. «Академія Хана» тут відіграла ключову роль — за нашими відео вчилися сотні мільйонів у всьому світі. Ми також показали, що відео є радше допомогою вчителю, ніж його заміною. За їхньою допомогою можна віддавати частини лекції на самоопрацювання, а в аудиторії приділяти більше часу персоналізованому навчанню, прикладним вправам і обговоренням. Тож учителі у такому форматі мають більшу цінність, а не меншу. Тепер уже настав час побачити, чи зможе генеративний ШІ зробити те саме — підтримати студентів і підняти цінність викладачів.

Демонстрацію GPT-4 Сем і Грег почали з того, що показали мені тестове завдання з університетського курсу біології для випускників шкіл, яке взяли із сайту Ради коледжів. Спитали, чи знаю я відповідь. Я прочитав усі варіанти й обрав варіант В. Потім вони попросили GPT-4 відповісти на це питання через інтерфейс чату (схожий на той, до якого вже всі звикли в ChatGPT). За мить GPT-4 дав правильну відповідь.

На цьому етапі я ще нічого не казав, але в мене вже з'явилися сироти на шкірі, хоч і лишився певний скептицизм.

— Секундочку, — сказав я. — Це ШІ, який вже може відповідати на питання з біології університетського рівня? — Я подумав: а що, як це просто випадковість? — А можете попросити його пояснити, як він дійшов до цієї відповіді?

Грег надрукував: «Будь ласка, поясни, як ти отримав цю відповідь». Протягом кількох секунд GPT-4 дав нам чітке, просте й детальне пояснення. І зробив він це в розмовному тоні — так цілком могла б відповісти людина, а не машина.

Тут я вже більше не приховував, наскільки вражений.

— А можете попросити його пояснити, чому всі інші варіанти неправильні?

Грег виконав моє прохання, і за мить GPT-4 пояснив, чому всі інші варіанти відповіді на це питання хибні.

Потім я спитав Грега, чи може GPT-4 сам підготувати питання з біології університетського рівня.

Він зміг, а потім ще десять таких.

Через два місяці я приїхав до Білла Гейтса, відзвітувати про стан справ в «Академії Хана», і дізнався, чому керівники OpenAI показали мені питання з біології. Білл розповів, як, уперше побачивши роботу GPT-3, він був вражений, але сказав команді OpenAI, що дійсно вражений буде тоді, коли ШІ зможе скласти іспит із біології університетського рівня. Команда OpenAI показала мені на першій зустрічі, що GPT-4 це може.

— Це змінює все, — сказав я Грегу й Сему, і в моїй голові закрутилося багато варіантів, як GPT-4 допоможе нам переосмислити освіту, компетенції, роботу й людський потенціал.

— Ми теж якось так подумали, — сказав Сем. — Він іще не ідеальний, але технологія вдосконалюється. Хтозна, якщо ми все зробимо правильно, можливо, освітяни захочуть цим користуватися.

Технологія, яка донедавна здавалася чимось із серіалу «Зоряний шлях», раптом набула реальних обрисів. Технологія, яку уявляли собі найвеличніші наукові фантасти, стала реальністю.

Час для хакатону

На початку 1940-х років геніальний математик Клод Шеннон запропонував кілька послідовних теорій. Зокрема, він створив теорію електронної комунікації, яка згодом стала підґрунтям цифрових технологій. У 1948 році, працюючи в Bell Labs, він почав роботу в напрямі, який ми тепер називаємо штучним інтелектом. Шеннон вирішив погратися з тим, як алгоритм вгадує мову. Він опублікував у Bell System Technical Journal працю «Математична теорія зв'язку». Цифрові комп'ютери тоді лише починали з'являтися, а до появи інтернету ще було далеко, й у теорії інформації Шеннон уперше припустив, що низка

ймовірнісних процесів може прогнозувати мовлення. Відстежуючи скільки разів слово з'явилося в тексті, він розробив алгоритм, що міг визначити, яке слово, найімовірніше, буде наступним. Незабаром алгоритм уже міг прогнозувати, яке слово, ймовірно, буде після того. Зрештою ця маленька мовна модель згенерувала речення. Що кращав процес, то природніше воно звучало. Це дуже спрощене пояснення, але системи, подібні до GPT-3 й GPT-4 — це, фактично, ті самі мовні моделі, тільки значно складніші й засновані на особливому тренуванні нейронної мережі. Базову ідею в них можна простежити до ранньої роботи Шеннона.

Незабаром після розробок Шеннона, у галузі, яка потім перетвориться на штучний інтелект, з'явився інший великий мислитель — Алан Тюрінг. Він не лише розгадував німецькі шифри й допомагав нам перемогти нацистів — Тюрінг іще й досліджував концепцію штучного інтелекту, визначаючи можливості машин розвинутиися настільки, щоб колись переконливо імітувати людський розум. У 1950 році він написав основоположну працю «Комп'ютерна техніка й інтелект», де представив концепцію гри в імітацію. Зараз ми знаємо її як тест Тюрінга. Уявіть, що ви з кимось розмовляєте, але співрозмовника водночас не бачите. Наприклад, ви переписуєтеся з цією людиною з комп'ютера чи телефона. Якщо ви не бачите людини й фізично не взаємодієте з нею — звідки вам знати, людина це чи машина? У цьому суть тесту Тюрінга. Для нього зазвичай потрібен суддя, який оцінює відповіді людини й машини. Мета машини — переконати суддю, що вона насправді людина. Вона має продемонструвати інтелект, розуміння й вміти підтримувати послідовну розмову — точнісінько як людина. Тюрінг уважав: якщо машина здатна раз у раз дурити суддю й переконувати його чи її, що вона — людина, таку машину можна вважати розумною. Іншими словами, якщо машина здатна пройти тест Тюрінга — це означає, що вона має інтелект, наближений до людського.

Улітку 2022 року, прийнявши пропозицію Сема й Грега протестувати нову технологію GPT-4, я зацікавився, наскільки високі

її шанси пройти цей тест. У середині 1990-х я вивчав штучний інтелект у МТІ. Тоді існували прості програми, які при перших взаємодіях ще могли перехитрити людину, але в довгій деталізованій розмові й близько не висловлювалися, як люди. Те, що одного дня машина зможе пройти тест Тюрінга, надто що це стане за мого життя, — завжди здавалося фантастикою. Тож мені було надзвичайно цікаво спробувати технологію, яка, видавалося, майже-майже мала таку здатність, а може, й уже цьому навчилась. Це досягнення могло б дорівнюватися за значущістю до відкриття науковцями холодного синтезу або до можливості подорожувати швидше за світло.

Коли зійшла перша хвиля захвату, я замислився і про те, як така начебто розумна технологія вплине на суспільство. ШІ був ключем до розв'язання величезної кількості проблем, але в нього були й потенційні недоліки. Якщо така велика мовна модель зможе допомагати навчати студентів, то зможе й писати за них есеї. А що як нова модель GPT лише заважатиме студентам, і вони через неї не розвиватимуть власні навички дослідження й роботи з текстом? Я зрозумів і таку річ: попри те, що GPT-4 може допомагати людям, бо їм буде легше комунікувати й розв'язувати проблеми, він — потенціальна загроза, оскільки через нього багато хто може втратити роботу й сенс чогось досягати в житті. Переконлива, людиноподібна технологія, здатна добре викладати, може також бути технологією, за допомогою якої поганці будуть шахраювати чи промивати мізки, а їхні жертви нічого й не підозрюватимуть.

У моїй уяві крутилася безліч неприємних сценаріїв і варіантів розвитку подій — від збору даних про наших дітей до того, яку ця технологія може викликати залежність. Я зрозумів, що ШІ змінює все, а отже, до нього треба поставитися серйозно. У великі мовні моделі інвестує не лише OpenAI — є ще Microsoft, Google і Meta, мовчу вже про такі держави, як Росія і Китай. Усі технологічні гіганти роками так чи інакше використовували штучний інтелект, щоб показувати нам рекламу, результати пошуку й дописи в соцмережах, і ми з усім цим

взаємодіємо цілими днями. Але цей ШІ, здавалося, відрізнявся — бо так і було. Наукові фантасти завжди розрізняли спеціалізований штучний інтелект, здатний оптимізувати щось одне, і штучний інтелект широкого профілю, який може логічно виконувати різні завдання, як людина. І саме з останнім варіантом у їхніх творах ми мали або утопічне, або антиутопічне майбутнє.

Для більшості людей великі мовні моделі, як-от GPT-4, найбільше схожі на загальний штучний інтелект, тому що вони можуть писати й нібито логічно мислити на будь-яку тему, тож їх можна застосовувати майже всюди. Генеративний ШІ також здатний створювати й аналізувати зображення. Було зрозуміло, що генеративний ШІ змінить наш світ так, як ми й уявити собі не можемо, і я усвідомив: вирішувати, чи дійсно ми готові до розвитку такої технології — не моя справа і не справа «Академії Хана». Вона вже розвивається. І як людина, яка покладає всі сили, щоб використовувати технології на добре, я подумав, чи можна скористатися генеративним ШІ так, щоб досягти якнайпозитивнішого ефекту й наблизитися до утопічного сценарію, особливо у сфері освіти.

Неочікувано отримавши доступ до найновішої моделі OpenAI, я в той самий вечір сів за комп'ютер. То була п'ятниця, і я не міг відірватися від нього майже всі вихідні. Просидівши за комп'ютером десять годин, я спробував переключитися й заснути, але в голові не вгамовувалися думки, на що здатна ця технологія і як ми, освітяни, можемо використати її для наших учнів. Я побачив, що способи застосування GPT-4 буквально нескінченні. Він не лише був здатен розв'язувати складні задачі й відповідати на питання із заковікою, а й пояснювати, чому саме так відповів і самостійно генерувати нові запитання. Він міг брати на себе різні ролі й відігравати їх дуже схоже на людину. За час діалогу з GPT-4 я просив його поводитися як репетитор з математики, репетитор з хімії, репетитор з англійської мови, й щоразу він підтримував зі мною розмову.

Але найбільше мою увагу привернуло, що він може писати в різних настроях і стилях. Наприклад, вивчаючи можливості

GPT-4 я не стримався і попросив його переписати Декларацію незалежності у стилі Дональда Трампа:

Панове, скажу вам так: настав нам час проголосити незалежність від цього недолугого уряду і цього жахливого очільника, Короля Георга III. Ну серйозно, цей мужик — повна катастрофа. Він відбирає наші права, він ігнорує бажання людей, і йому цікавіше набивати собі кишені, аніж реально допомагати народу. Це просто сумно, дуже, дуже сумно.

Але ми, люди, більше цього не терпітимемо. Ми створимо власну країну, і це буде найкраща країна, найвеличніша країна з усіх, що бачив світ. Повірте, ніхто не знає, як зробити країну величною, краще за мене. У нас буде найкраща економіка, найкраща армія, найкраще все. Це буде дивовижно.

Потім я попросив ChatGPT переписати Декларацію незалежності у стилі доктора Сьюза*:

На вільній землі, у домі сміливих людей
Родилася нова нація, де більше немає рабів
І перед правлінням короля, що за морями сидить
Ми проголосили незалежність, і радісно викрикнули

З ручкою у руці й вогнем у серці
Ми записали все, що нам не подобалося,
до останнього
Не буде більше податків без представників
Не буде солдатів у наших домах, не буде злості

* Псевдонім американського дитячого письменника.

Ми проголошуємо ось такі істини вам
Щоб усі люди були рівні, так, ти і я
У нас є права, які не можна забрати
Життя, свобода, і щастя. І буде так.

Відколи люди масово отримали доступ до ChatGPT, уже багато хто показував такі фокуси. Але тоді я одним із перших одержав можливість покористуватися такою технологією та попросив ШІ про таку витівку. Відверто кажучи, GPT-4 був значно кращим за перший варіант ChatGPT, який широкий загал побачив лише через декілька місяців. Результати моїх запитів мене вразили, повеселили й навіть трішки налякали. Коли я щось питав чи просив поради, він відповідав так, що здавалося, він справді так думає. Але по той бік екрана не було людини, яка друкувала ці репліки, й алгоритму, який би генерував текст за традиційною логікою «якщо — то». Тупих роботоподібних відповідей я теж не побачив. Натомість я отримував різні відповіді на одне й те саме питання, і ці відповіді враховували контекст нашої попередньої розмови.

Якщо конкретніше, то я зрозумів, наскільки сильно ця технологія здатна змінити наш підхід до шкільної, вищої освіти й не тільки. Той ШІ ще був не ідеальний. Він помилявся в обчисленнях більше, ніж мені хотілося б, але коли я навчився правильно писати промпти — то побачив, що він і реагує краще. Під кінець тих вихідних я замислився, що станеться, якщо я зберу кілька десятків найяскравіших представників технологічної та освітньої галузей і вони пограються із цією технологією разом зі мною. OpenAI погодився дати доступ ще десь тридцятьом інженерам, контент-креаторам, освітянам і дослідникам із нашої команди, щоб вони теж поекспериментували з GPT-4.

Настав час хакатону.

Щопівроку ми в «Академії Хана» даємо своїм співробітникам тиждень, коли вони можуть працювати над чим завгодно, аби лише воно було пов'язане з нашою місією. Я презентував

GPT-4 маленькій групі людей із нашої команди й дав їм зелене світло робити, що вони захочуть. Ми влаштовували мозкові штурми, співпрацювали, шукали інноваційні рішення, розробляли їх і таки отримали дійсно круті й доречні ідеї. Те, що ми зрештою стали називати ШІкатоном, дало нам десятки цілком нових концепцій і модальностей навчання, про які ніхто з нас раніше не думав. Наприклад, а що, як ШІ допомагатиме викладачам складати плани уроків? Що, як він стане опонентом для студента в дебатах? Що, як він зможе створювати проекти? Що, як він допоможе учням позбутися стресу чи надихне їх на створення чогось нового? Що, як він зможе протестувати їхні знання чи керувати процесом повторення вивченого? Освітяни створюватимуть новітні активності, у яких студенти зможуть брати участь за допомогою ШІ. ШІ допомагатиме учням писати твори, вчитиме їх робити це краще, бо миттєво даватиме зворотний зв'язок.

Із цього моменту учасники ШІкатону почали вивчати питання безпеки й упереджень. (Не забувайте, все це було задовго до того, як OpenAI презентував ChatGPT широкому загалу.) Ми відзначили те, що дійсно нас хвилювало. Чи справді це добре, якщо студенти писатимуть есеї за допомогою генеративного ШІ, проводитимуть дослідження, чи гаразд, якщо він їх перевірятиме чи навіть допомагатиме подавати заявки до коледжу? Ми хвилювалися, що зі штучним інтелектом наші діти перетворяться на чітерів, які нічого не вчитимуть. Батьки, які раніше допомагали дітям робити домашнє завдання, можуть утратити важливу точку дотику з ними. Чи стане ШІ чарівною паличкою для викладачів, а чи викличе сумніви у їхній здатності чогось навчити? Я не вважав, що ШІ позбавить викладачів роботи. У найкращому разі він мав покращити процес навчання, але я також хвилювався, що багато в чому він серйозно заважатиме.

Майже за два десятиліття до того я бачив, що люди боялися подібних речей, коли говорили про освітні відео за запитом. Чи не відволікатимуть такі відео студентів? Чи не знизить це їхню здатність утримувати увагу? Чи не ізолює учнів від викладачів, чи не втратять вони можливість формувати стосунки з ними?