

Prof. Andrzej Dragan: Po raz pierwszy w historii nauczyliśmy się robić rzeczy, których nie rozumiemy

© Piotr Cieśliński – wywiad z prof. Andrzejem Draganem

35 lat Gazety Wyborczej. Teksty wszech czasów, 07.05.2024

https://wyborcza.pl/7,75400,30547020,prof-andrzej-dragan-ai-jest-teraz-na-poziomie-osmiolatka.html#s=S.embed_article-K.C-B.1-L.1.zw

Sztuczna inteligencja jest teraz na poziomie ośmiolatka, ale za 10 lat może nas przerosnąć. Jesteśmy jednym z ostatnich pokoleń, które coś wnosi do rozwoju i logicznego myślenia - uważa prof. Andrzej Dragan, fizyk teoretyk, artysta, popularyzator nauki.

Gazeta Wyborcza obchodzi 35-lecie. Jej pierwszy numer ukazał się 8 maja 1989 r. Z tej okazji przypominamy najważniejsze i najciekawsze teksty.

Wywiad z prof. Andrzejem Draganem ukazał się na początku 2024 roku.

Piotr Cieśliński: W 2023 r. pojawiły się modele sztucznej inteligencji, które wydają się naprawdę inteligentne, takie jak ChatGPT firmy OpenAI czy Bard Google'a. Wkrótce będziemy dzielić historię naszej cywilizacji na przed i po tym wydarzeniu...

Andrzej Dragan: Niektórzy mówią, że będzie tylko przed.

Nie przetrwamy tego?

- Niektórzy tak sądzą.

Dziennikarze lubią robić z igły widły. Czy patrząc chłodnym okiem naukowca, zgadzasz się, że to rewolucyjna technologia, która może wywrócić nasz świat do góry nogami? Jak wynalezienie koła albo ujarzmienie ognia?

- Ta rewolucja zasadniczo polega na tym, że po raz pierwszy w historii nauczyliśmy się robić rzeczy, których nie rozumiemy. Słabo gram w szachy, ale mogę napisać algorytm uczenia maszynowego, który wytrenuje sieć grającą w szachy i ta sieć ogra mistrza świata. Ba, mogę w ogóle nawet nie wiedzieć, jak się gra w szachy, a taki mistrzowski algorytm stworzyć.

To sukces, ale jednocześnie porażka, że potrafimy stworzyć coś, co w pewien sposób nas przekracza.

Tego nie było - dotychczas rozwój technologii napędzało zrozumienie. Dlatego nauka była taka ważna, bo pozwalała nam najpierw coś zrozumieć, a dopiero potem wykorzystać zdobytą wiedzę. Teraz możemy tworzyć bez zrozumienia.

Nie rozumiemy, jak działa ChatGPT i inne sieci sztucznej inteligencji?

- Podam taki przykład. Mój kolega robił doktorat z rozpoznawania ręcznego pisma. 20 lat temu chciał napisać algorytm, który będzie je rozpoznawał. Poddał się, to było wtedy za trudne, efektywność algorytmu była bardzo niska.

A teraz mogę stworzyć sieć, wytrenować ją i po pięciu sekundach treningu na iPadzie mieć aplikację, która rozpoznaje pismo odręczne ze skutecznością 99 proc. Chociaż sam w ogóle nie mam pojęcia, jak się rozpoznaje pismo.

Algorytmy tworzenia takich sieci i uczenia ich są bardzo proste. Nawet nie muszę umieć programować. Niedawno zrobiłem eksperyment i poprosiłem ChatGPT, żeby napisał kod, który stworzy nową sieć, znajdzie sobie sam w internecie dane dla tej sieci i wytrenuje ją do wykonania tego zadania.

Kto więc niby rozumie, jak działa ta nowa sieć? Ja nie, bo kod napisał ChatGPT, a ChatGPT rzekomo nic nie rozumie, bo nie ma żadnej świadomości.

Kto więc w zasadzie rozumie ten twór, jak on powstał i jak działa?

Sytuacja jest więc taka: ChatGPT daje nam algorytm, który robi to, co chcemy, bardzo skutecznie, często lepiej niż człowiek. Ale nie wiemy, w jaki sposób dochodzi do tak dobrych efektów?

- Jestem na to wyczulony, bo robię w branży, gdzie mam rozumieć i płacam mi za to, żebym zrozumiał. Na tym polega praca fizyka. Nie podoba mi się więc, że mamy trudności ze zrozumieniem, jak działają te algorytmy. Irytuje mnie też, że część środowiska IT nie rozumie, że tego nie rozumie.

W językach programowania, które rozumiemy (np. w Phytonie czy C++), nie bylibyśmy w stanie napisać algorytmu grającego w szachy czy konwersującego tak dobrze jak ChatGPT. To jest poza naszym zasięgiem.

Pamiętam moje zdumienie, gdy pierwszy raz siadłem do ChataGPT, zadałem kilka pytań, a on mi odpowiedział. Pomyślałem wtedy: „O rany, on mnie rozumie. Rozumie też to, co sam pisze!”. Myślisz, że te systemy rozumieją?

- Trudno jest zdefiniować, co to znaczy rozumieć. Nie lubię się dzielić opiniami na jakiegokolwiek tematy, bo opinie są mało istotne, staram się patrzeć na konkrety. Spójrzmy więc, czym jest inteligencja. Nie jest ona, jak często można usłyszeć, umiejętnością rozwiązywania problemów. To co najwyżej definicja efektywności.

Naukowcy z IT wymyślili inną definicję inteligencji, która jest mierzalna i unifikuje to, z czym do tej pory ją kojarzyliśmy – twierdzą mianowicie, że jest to umiejętność dostrzegania analogii.

Tak jak w teście, w którym trzeba dokończyć zdanie: "Ręka ma się do rękawiczki jak noga do...?". Albo w ciągach cyfr czy obrazków, które trzeba uzupełnić o brakujący element.

Poradzisz sobie z tym, gdy znajdziesz analogię, zrozumiesz relację. Poeci wrażliwi na język jak nikt inny potrafią tworzyć celne metafory i łączyć ze sobą słowa. To przejaw inteligencji.

Ciekawe jest to, że ta definicja wiąże inteligencję z umiejętnością kompresowania danych. Żeby skompresować duży plik, także musisz umieć dostrzec analogie - w tym wypadku w ciągu zer i jedynek. Gdy jakiś wzór się powtarza wielokrotnie, to wystarczy go umieścić w pliku tylko raz, a w pozostałych miejscach odwołać się do niego. Dostrzegając analogie w długim ciągu zer i jedynek, skrucasz jego długość.

Jeśli więc przyjąć taką precyzyjną definicję, to z całą pewnością ChatGPT jest inteligentny.

To już nie kwestia opinii, tylko faktów - ten model jest w stanie skompresować całą informację o rzeczywistości, m.in. wiele książek, które "przeczytał", do sieci, która ma relatywnie niewielki rozmiar, ale zachowuje się tak, jakby całą tę wiedzę posiadała.

Nasz mózg też tak działa?

- Na podobnej zasadzie. Gdy trafiamy w nowe miejsce, na nową sytuację, to wiemy, jak się zachować, bo to, co widzimy, porównujemy z tym, czego doświadczyliśmy w przeszłości, czyli dostrzegamy analogie. I to właśnie jest miarą inteligencji, a nie umiejętność rozwiązywania problemów. Ta ostatnia definicja inteligencji jest porażką inteligencji.

W tym sensie algorytmy AI są jak najbardziej inteligentne. Czy rozumieją - to inna kwestia. Z moich testów wynika, że one bardzo dużo rozumieją.

Testowałeś ChataGPT w wersji 4, która nie jest jeszcze dostępna za darmo. Porównywałeś ją z inteligencją swojego ośmioletniego syna.

- W pewnym momencie rzuciłem na Twitterze wyzwanie - poprosiłem o takie zagadki logiczne, które potrafi rozwiązać mniej więcej ośmioletnie dziecko, ale byłyby za trudne dla ChataGPT 4. Dostałem kilkaset problemów, sprawdzałem wszystkie po kolei i nie było ani jednej zagadki, na której ChatGPT by się wyłożył.

Czyli ChatGPT co najmniej dorównuje już typowemu ośmiolatkowi?

- Jeśli dodamy dla świętego spokoju, że to zagadki oryginalne, których nie było w danych treningowych dla algorytmu, to tak - zagadki na poziomie ośmiolatka nie sprawiają trudności ChatowiGPT. Oczywiście wciąż się myli.

Czasem słyszę od programistów: "GPT umie programować, ale ja robię to lepiej!".

To samo można powiedzieć o ośmiolatkę, który zaczyna programować. Wiadomo, że dzieci programują gorzej niż dorośli.

Teraz GPT jest na poziomie ośmiolatka, ale pół roku wcześniej - w wersji 3,5 - był na poziomie trzylatka. Jak się zapytałem: "co to jest - po stawie pływa i kaczka się nazywa?", to odpowiedział, że głowica.

Bez sensu.

- Generował jakieś kompletne bzdury i halucynował od rzeczy. Ale już pół roku później przeskoczył kilka lat rozwoju dziecka.

W serwisie X (dawniej Twitter) główny naukowiec od GPT Ilya Sutskever niedawno powiedział: "Jeśli jesteś kimś, kto ponad wszystkie ludzkie cechy ceni sobie inteligencję, to nadchodzą dla ciebie złe czasy". Sutskever w przeciwieństwie do ciebie dostrzega już przejawy świadomości w działaniu AI.

- Ale to właśnie kwestia definicji. Musimy zdefiniować, co nazywamy świadomością, bo inaczej to będzie kłótnia o słowa i interpretacje. Na przykład prof. Włodzisław Duch twierdzi, że świadomość to jest umiejętność obserwacji swoich stanów wewnętrznych i w tym sensie AI już ma świadomość.

Ale w takim wypadku jakkolwiek warunek logiczny w programie (tj. odwoływanie się do stanów pamięci), których zazwyczaj jest pełno, to już jest jakiś objaw świadomości. Dlatego myślę, że jest to warunek konieczny - tj. aby być świadomym, muszę umieć obserwować swoje stany wewnętrzne - ale nie jest to warunek wystarczający.

Nawet jednak w tym sensie algorytm ChatGPT-4 nie spełnia warunku świadomości, bo jest oparty na strukturze transformera. Jest siecią jednokierunkową - tam nie ma żadnych zapętleń, żadnych struktur rekursywnych.

Ale można sobie wyobrazić, że zbudujemy taką sieć neuronową, która będzie podobna do tego, co się dzieje w naszym mózgu. I ta sieć stanie się „kimś” – czyli bytem świadomym samego siebie, swojej odrębności od świata zewnętrznego.

- Nie znam żadnych barier, które by w tym przeszkadzały. Nie wiem, czy w mózgu, który jest papką węglowo-wodorowo-tlenową, jest coś szczególnie istotnego, czego nie można odtworzyć w jakiejś sztucznej strukturze. Może coś takiego wyjątkowego kiedyś odkryjemy, to byłoby interesujące.

Na razie naiwne kopiowanie struktury mózgu do sieci neuronowych prowadzi do zaskakująco zbieżnych rezultatów. Test Turinga już padł.

Prawie nikt nie zauważył, że jeden z najdonioślejszych celów ludzkości – stworzyć maszynę, która będzie zachowywać się w sposób nieodróżnialny od człowieka – został osiągnięty. Ogłosił to w lipcu 2023 r. „Nature”.

AI już zdała test Turinga?

- Tak przynajmniej twierdzi „Nature”.

Geoffrey Hinton, który jest ojcem chrzestnym głębokiego uczenia maszynowego (Ilya Sutskever był jego doktorantem), jeszcze w latach 60. proponował budowę wielowarstwowych sieci neuronowych. On chciał zrozumieć, jak działa mózg, ale skoro nie można było robić eksperymentów na prawdziwym mózgu, to zrobił jego symulację w sztucznej sieci, żeby na niej eksperymentować. Wtedy nikt nie wierzył, że to będzie działać, ale okazało się, że miał rację.

Sieci AI, które teraz udało się zbudować, zadziwiająco dobrze przypominają mózg, przynajmniej jeśli chodzi o umiejętność konwersacji.

Należysz do grona tych naukowców, którzy ostrzegają przed sztuczną inteligencją, tak jak Geoffrey Hinton. Dlaczego?

- Nie ostrzegam, nie chcę doradzać komukolwiek w jakiegokolwiek sprawie, a tym bardziej całemu światu. Zależy mi tylko na tym, abyśmy nie tkwili w jakichś urojonych przesądach. A jeden z nich mówi, że te algorytmy kompilują jedynie to, co już widziały, tj. na czym były trenowane, więc nie są w stanie niczego nowego stworzyć. Bzdura.

Przeprowadzałem różne testy, aby to sprawdzić. Na przykład prosiłem sztuczną inteligencję o narysowanie czarno-białej fotografii Davida Lyncha trzymającego kurę.

Jest tylko jedno takie czarno-białe zdjęcie, łatwo na nie natrafić w internecie, pewnie mogło trafić do danych treningowych algorytmów. Ale nie zdarzyło się, aby AI stworzyła obraz, który wykorzystywał choć jeden element tego oryginalnego zdjęcia.

Czyli nie kopiowała, lecz stworzyła swoje własne, oryginalne dzieło.

- Było to pierwsze pytanie, które mnie interesowało – czy te algorytmy tylko kompilują, czy tworzą, i w jakim sensie tworzą? Ponieważ nie wiemy, jak działają, to mogę tylko zgadywać. I wydaje mi się, że one uczą się na danych treningowych tak jak człowiek, który np. czyta książki.

Lektura mnie jakoś wzbogaca, zmienia moje spojrzenie na świat, na podstawie przeczytanego tekstu tworzę sobie obraz rzeczywistości i to wpływa na to, co piszę, jak rozmawiam. Te algorytmy podobnie się uczą. Nie trzymają danych treningowych jako fizycznych kopii, tylko tworzą sobie na ich podstawie model świata, a potem wytwarzają treści, które są z tym modelem zgodne.

Zadałem Bardowi pytanie: „Czy ludzie powinni się ciebie bać?”. Odpowiedź była długa, ale sens był taki, że on jest tylko użytecznym narzędziem i to od nas – ludzi – zależy, jak zostanie wykorzystany.

- Ale nie na tym polega problem. Gdyby to było tylko narzędzie, które można użyć w dobry lub zły sposób, to byłaby to kwestia kontroli ludzi, którzy go używają. Z tym byśmy sobie mogli jakoś poradzić.

Swoją drogą dwa miesiące temu w "Nature" opublikowano algorytm, który po raz pierwszy pokonuje nas nie w jakiejś grze logicznej czy planszowej - bo w szachy, go, czy pokera już dostajemy baty - ale w wyścigu dronów sportowych w rzeczywistym świecie.

Następne będą pewnie myśliwce, bo nie ma zasadniczej różnicy między sterowaniem dronem i myśliwcem.

A więc jest prawdą, że te algorytmy są niezwykle efektywnym narzędziem i trzeba ich pilnować.

Jak długo będę tylko narzędziem, które może być użyte w dobrym lub złym celu, będziemy mieli święty spokój. Ale nie na tym polega nowość.

A na czym?

- Po raz pierwszy w historii na horyzoncie jest trzecia możliwość: że to narzędzie samo zdecyduje o czymś albo zrobi coś, czego byśmy nie chcieli. To się może wydarzyć, jeszcze nie teraz, ale jest już do pomyślenia i wyobrażenia. I nie będzie można nikogo obarczyć za to odpowiedzialnością. Przy pełnej autonomii AI może podejmować decyzje bez naszego udziału.

Na razie to wciąż narzędzie, które wykonuje naszą wolę. Ale tłumaczenie, czym teraz jest AI, to mydlenie oczu, bo sieci rosną w postępie nie wykładniczym, lecz podwójnie wykładniczym. Z przyrostem wykładniczym mamy do czynienia w reakcji łańcuchowej, a rozmiar tych sieci rośnie tak jak reakcja łańcuchowa reakcji łańcuchowych.

Jeśli takie tempo zostanie utrzymane, te algorytmy przerosną nas intelektualnie pod każdym względem w ciągu maksymalnie 10 lat.

Nic nie stoi na przeszkodzie, aby przestały być tylko bezwolnym narzędziem i zaczęły definiować własne cele.

ChatGPT do tego jeszcze nie służy, on tylko odpowiada na nasze prompty. Ale na pewno wielu naukowców stara się już wyodrębnić w takich sieciach bardziej autonomiczne byty, które nie będą już potrzebować żadnych promptów.

Będą same zadawać pytania i stawiać przed sobą cele?

- Na razie takich sieci nie ma, ale nie ma też fundamentalnych przeszkód, aby się pojawiły w bliskiej przyszłości.

One potrafią coraz lepiej programować. Robiłem kilka dni temu testy, czy złożona sieć neuronowa może sama stworzyć nową sieć neuronową, wyszukiwać dla niej dane treningowe.

To tylko kwestia czasu, kiedy one zaczną się same modyfikować. A wtedy to w ogóle będzie już wolnoamerykanka.

Max Tegmark, fizyk z MIT, dał przepis, jak bezpiecznie postępować ze sztuczną inteligencją.

- Sformułował cztery reguły bezpieczeństwa, dzięki którym nie musielibyśmy się przejmować, że coś pójdzie nie tak.

Jedna z nich mówiła: "Nie ucz AI kodować". Inna: "Absolutnie nie wpuszczaj jej do internetu".

- Chodziło o to, aby nie znała psychologii ludzi, żeby nami nie manipulowała. Najlepiej, aby o ludziach nic nie wiedziała.

Ale pierwsze, co zrobiliśmy z tymi algorytmami, to wpuściliśmy je w sieci społecznościowe.

Nauczyły się nami manipulować, abyśmy skrolowali posty długie godziny.

Generalnie gdy ludzie stworzyli duże sieci neuronowe, od razu złamali wszystkie reguły bezpieczeństwa Tegmarka.

Powiedziałeś, że AI może nas zdominować już za 10 lat. Potrafisz sobie wyobrazić, jak świat będzie wtedy wyglądać?

- Mówiłem, że jeżeli tempo rozwoju tych sieci utrzyma się, to za 10 lat, a może wcześniej, będą mieć inteligencję na poziomie dorosłego człowieka. W tym sensie, że będą rozwiązywać wszystkie logiczne problemy, które my potrafimy rozwiązać.

A potem pójdą dalej?

- Jeżeli ich rozwój nie spowolni. Nie wiemy, czy nie ma jakiegoś sufitu, który je zablokuje. Nikt rozsądny nie odpowie, jaka czeka nas przyszłość za 10 lat. Ale chyba po raz pierwszy nie mamy pewności, że nasza dominująca pozycja zostanie utrzymana.

Jest jasne, że nie jesteśmy koroną ewolucyjnego stworzenia, lecz jakimś pośrednim etapem, być może mało istotnym w dłuższej skali. I biologiczna ewolucja zastąpi nas czymś innym.

Teraz pojawiają się algorytmy, które tempo ewolucji podkreślają. Przestaje być oczywiste, czy to, co nas zastąpi, będzie jeszcze z białka, czy może już zbudowane w inny sposób.

Gdy sztuczna inteligencja już nas przegoni, mogłaby pomóc w zrozumieniu praw natury. Jakie pytanie byś jej zadał, kiedy stanie się bardziej inteligentna niż najgenialniejsi z ludzi?

- Większości problemów, którymi zajmuję się jako fizyk, nie rozumiem. Mnóstwo jest takich pytań, że nawet nie wiem, jak się zabrać do odpowiedzi.

Przykład?

- Gdzie spojrzeć... Żyjemy w przestrzeni, która ma trzy wymiary. Dlaczego trzy, a nie siedem? Dlaczego stałe fizyczne mają taką, a nie inną wartość? Nie tylko nie jesteśmy blisko odpowiedzi, ale nawet nie wiemy, jak się do tego zabrać.

Problem z zadawaniem takich pytań jest taki, że prawdopodobnie są głupie.

Weźmy Newtona, który był chyba najwybitniejszym uczonym wszech czasów. Nurtowało go pytanie, czy światło składa się z cząstek, czy jest może falą? W jego czasach trwała na ten temat dyskusja. On więc mógłby zapytać AI, czy światło jest cząstką, czy falą. Ale odpowiedź brzmi: ani tym, ani tym. To jest źle postawione pytanie, nie da się odpowiedzieć "tak" albo "nie". Należałoby wytłumaczyć Newtonowi obecną teorię, według której – to też jakieś pośrednie rozumienie – światło nie jest ani cząstką, ani falą.

Gdybym więc ja miał odpowiadać Newtonowi na to pytanie, tobym rzekł: "Zacznijmy od innego pytania. Zbudujmy najpierw nowy obraz świata i wtedy okaże się, dlaczego to pytanie, które zadajesz, jest bez sensu".

Z tego punktu widzenia może najlepiej byłoby więc zapytać AI: jakie pytanie warto w tym momencie zadać?

Bo te pytania, które sobie teraz zadajemy, mogą nie mieć sensu?

- Tego nie wiemy, ale historia pokazuje, że pytania, które zadawano w przeszłości, były śmieszne albo źle postawione. Nie ma powodu sądzić, że to się nagle zmieniło i teraz już zaczęliśmy zadawać same dobre pytania.

Czy gdyby AI dać wszystkie dane i informacje dostępne ludziom na przełomie XIX i XX w., toby wymyśliła teorię względności?

- Jeszcze nie, bo na razie mamy do czynienia z ośmiolatkiem.

Ale wykonajmy prosty eksperyment myślowy. Już teraz ośmiolatki nie są w stanie rozwiązać zagadek, które rozwiązuje AI. Te algorytmy rozwijają się szybciej.

Co znaczy, że dzieci, które teraz mają osiem lub mniej lat, już nigdy nie rozwiążą problemów, które będą za trudne dla sztucznej inteligencji.

Za 10 lat, kiedy one będą miały 18 lat, sztuczna inteligencja będzie na dużo wyższym poziomie intelektualnym.

Czyli one już żyją w świecie, w którym AI zawsze będzie od nich intelektualnie sprawniejsza?

- Jeśli, powtarzam, nic nie zahamuje rozwoju AI.

Czyli nasze pokolenie...

- Byłoby jednym z ostatnich, które coś wnosi do rozwoju i logicznego myślenia. Jeszcze możemy się cieszyć, że ChatGPT jest na tyle głupi, że możemy rozwiązać zagadkę, której on nie potrafi rozgryźć.

Oczywiście jest też hipoteza, że umiejętności tych algorytmów są ograniczone jakością danych treningowych. Przeskok pomiędzy wersją ChataGPT 3,5 oraz 4 polega na dwóch rzeczach. Po pierwsze, prawdopodobnie ta nowa sieć jest dziesięć razy większa. Prawdopodobnie, bo OpenAI przestało być "open" i już nie ujawnia szczegółów technicznych tych sieci.

Drugim z ważnych udoskonaleń - oprócz drobnych udoskonaleń algorytmu uczącego - jest to, że jakość danych treningowych jest dużo lepsza. Są to lepiej dobrane dane. Pojawia się więc taka myśl, że jakość danych jest kluczowa do tego, jak te algorytmy się rozwijają.

Jeśli my je karmimy tylko tym, co sami wiemy, to te algorytmy nas nie przekroczą. Mogą nam dorównać, ale nie pójść dalej.

Nakarmią się całą wiedzą, jaką do tej pory zgromadziliśmy, przeczytają wszystkie książki w bibliotekach, przejrzą cały internet i na tym staną, bo nie będą miały więcej danych do nauki?

- Tak mówi ta hipoteza. Ale to też jest pewnego rodzaju przesąd. W 2016 r. odbył się słynny pojedynek, w którym algorytm AlphaGo pokonał mistrza świata w go. Arcymistrz Lee Sedol po tej porażce zniechęcił się i przestał grać, choć był jeszcze młody. Pół roku później ta sama firma wypuściła nowy algorytm AlphaZero, który też grał w go i ograł wersję AlphaGo 100 do 0.

Najciekawsze jest to, że ten nowy algorytm w ogóle nie miał żadnych danych treningowych. Znał tylko reguły gry i wiele razy grał sam ze sobą. W ten sposób sam wymyślił strategię wygrywającą.

W przypadku prostych zamkniętych gier, które mają zbiór ustalonych reguł, dane treningowe są niepotrzebne. Wielkim wyzwaniem były gry w pokera i brydża, ale tu już też przegrywamy.

Nie ma więc dla nas nadziei?

- W niedalekiej przyszłości podobnie będą radziły sobie algorytmy, które operują w rzeczywistym świecie i mają bardziej otwarte zadania. Nie przesuwają tylko figur, lecz muszą podejmować różne inne decyzje. Jakość tych decyzji może być ćwiczona poprzez symulacje prowadzone wewnątrz AI bez udziału zewnętrznych danych treningowych.

Nie widzę więc powodu, dla którego nasze dane i wiedza miałyby być ograniczeniem dla algorytmów sztucznej inteligencji.

Jakie zawody zostaną dla dzisiejszych ośmiolatków?

- Do niedawna wydawało się, że sztuczna inteligencja najmniej zagraża zawodom wysoko wykwalifikowanym. Że to słabo wykwalifikowani będą zastępowani maszynami.

Wszystkich zaskoczyło to, że programiści zaczęli korzystać ze wspomagania AI – są takie narzędzia jak copilot, które tworzą dla nich proste elementy kodu. Część pracowników IT to bagatelizuje, uważając, że są nie do zastąpienia. Ale w niektórych firmach już ogranicza się zatrudnienie, bo dzięki nowym narzędziom efektywność programistów tak bardzo wzrosła, że w tym samym czasie są w stanie napisać kilka razy więcej kodu.

Nie jest więc wcale oczywiste, że najpierw będą eliminowane zawody niskiego poziomu. Być może zwolnienia zaczną się właśnie od góry.

Warto zwrócić uwagę na jedną rzecz. Jesteśmy dumni z tego, że myślimy, że gramy w szachy, zajmujemy się fizyką czy matematyką. Ale tylko niewielką część mózgu używany do świadomych procesów, tylko mały kawałek do myślenia.

Większość mojego mózgu pracuje w tym momencie nad tym, żebym siedział w miarę prosto, żebym się nie przewrócił, żeby skoordynować pracę wielu mięśni i zarządzać moim organizmem.

To dlatego w sportach fizycznych algorytmy wciąż nie dorastają nam do pięt, natomiast przegrywamy z nimi w szachy i gry logiczne, gdzie używamy inteligencji. Prosta sieć jest w stanie ograć mnie w szachy, ale nie wygra w piłkę nożną.

Nie wiem, czy to jest dobra prognoza, ale wygląda na to, że to zawody fizyczne będą trudniejsze do wyeliminowania. Tymczasem umysłowe nie są tak bezpieczne, jak to się wydawało jeszcze kilka lat temu.

Andrzej Dragan jest profesorem Uniwersytetu Warszawskiego i profesorem wizytującym w Narodowym Uniwersytecie Singapuru. Zajmuje się informacją kwantową, szuka związków między teorią względności i nowym opisem świata, jaki przyniosła fizyka kwantowa. Jest również fotografem (twórcą portretów, m.in. Davida Lyncha, Jana Peszka czy Jerzego Urbana), kompozytorem oraz twórcą filmowym, a także popularyzatorem nauki (autorem znakomitej książki "Kwantechizm, czyli klatka na ludzi", w której opisuje Wszechświat i życie z punktu widzenia fizyki kwantowej i teorii Einsteina)

Rozmowa odbyła się podczas II edycji Festiwalu Nowe Oświecenie, którego organizatorem jest Ursynowskie Centrum Kultury „Alternatywy” w Warszawie. Została skrócona i zredagowana (tekst znajdzie się w podsumowaniu ubiegłorocznego wydarzenia, wydanym podczas kolejnej edycji festiwalu w 2024 r.).

Piotr Cieśliński - Dziennikarz i redaktor naukowy, z wykształcenia fizyk, pracuje w "Gazecie Wyborczej" od roku 1991, popularyzuje odkrycia z nauk ścisłych, kosmosu, matematyki i techniki