

8 największych zalet i wad sztucznej inteligencji



Sztuczna inteligencja przeżywała boom w ciągu ostatnich kilku lat i jest to dobre i złe z wielu powodów. Przyszłość nie jest tym, czym „była kiedyś” i to jest pewne. Kto mógł sobie to wyobrazić? Czy dojdzie do przejęcia „Sky-Net”, jak w filmie „Terminator”? Czy sztuczna inteligencja będzie rewolucyjna i pomoże miliardom ludzi żyć bezpieczniejszej, zdrowiej i wydajniej? Oto 8 zalet i wad tej oszałamiającej ery technologicznej, w której wszyscy teraz żyjemy.

1. Sztuczna inteligencja (AI) może pomóc w badaniach i pisaniu, ale może również tworzyć fałszywe wiadomości i krytyczne błędy, w tym nielogiczne wnioski i halucynacje.
2. Sztuczna inteligencja może (pomóc) kontrolować drony, pojazdy i broń wojskową, ale może zostać zhakowana lub przechytrzyć użytkowników i zwrócić się przeciwko ludzkości.
3. Sztuczna inteligencja może tworzyć obrazy i filmy, łącząc informacje i wizualizacje w celu pobudzenia wyobraźni i w celach rozrywkowych, ale może też oszukać ludzi, by uwierzyli w fałszywe koncepcje, takie jak inwazje kosmitów lub przemówienia polityczne, które nie miały miejsca.
4. Sztuczna inteligencja może być wykorzystywana do wykonywania wielu zadań wydajniej i spójniej niż ludzie, takich jak praca w fabrykach, ale może to oznaczać koniec milionów miejsc pracy, które nie wymagają krytycznego myślenia, kreatywności ani interakcji międzyludzkich.

5. Sztuczna inteligencja może być wykorzystywana przez roboty do ratowania ludzkiego życia, na przykład w walce, ale wtedy ci żołnierze i psy-roboty mogą stać się agresywne lub popełniać krytyczne błędy, które ranią lub zabijają ludzi.

6. Sztuczna inteligencja może być pomocna w domu, pomagając w prostych zadaniach, wyszukiwaniu informacji lub rozrywce, ale to eliminuje wiele interakcji międzyludzkich, czyniąc życie mniej społecznym, serdecznym, satysfakcjonującym i uduchowionym.

7. Sztuczna inteligencja dostarcza informacji bez dramatyzmu, nastawienia czy ego, ale informacje te mogą być cenzurowane, aby celowo dostarczać nielogicznych dezinformacji i dezinformacji na najważniejsze tematy, takie jak zdrowie i bezpieczeństwo.

8. Sztuczna inteligencja może zmienić przyszłość dzięki technologii, ale może też zmienić przeszłość, pisząc historię na nowo.

Nowy model języka wizyjnego LLaVA-01 opracowany w Chinach poprawia zdolności rozumowania, ale zmagają się ze złożonymi zadaniami wymagającymi logicznego rozumowania.

Naukowcy z wielu uniwersytetów w Chinach zaprezentowali LLaVA-01, nowy model języka wizji, który znacznie poprawia zdolności rozumowania dzięki zastosowaniu systematycznego i ustrukturyzowanego podejścia.

Tradycyjne modele języka wizji (VLM) o otwartym kodzie źródłowym często zmagają się ze złożonymi zadaniami

wymagającymi logicznego rozumowania. Zazwyczaj wykorzystują one metodę bezpośredniego przewidywania, w której generują odpowiedzi bez rozbijania problemu lub nakreślania kroków niezbędnych do jego rozwiązania. Takie podejście często prowadzi do błędów, a nawet halucynacji.

Aby zaradzić tym ograniczeniom, badacze stojący za LLaVA-o1 zainspirowali się modelem o1 OpenAI. Model o1 OpenAI pokazał, że wykorzystanie większej mocy obliczeniowej podczas procesu wnioskowania może zwiększyć umiejętności rozumowania modelu językowego. Jednak zamiast po prostu zwiększać moc obliczeniową, LLaVA-o1 wykorzystuje unikalną metodę, która dzieli rozumowanie na ustrukturyzowane etapy.

LLaVA-o1 działa w czterech odrębnych etapach:

1. Zaczyna od podsumowania pytania, identyfikując główny problem.
2. Jeśli obraz jest obecny, skupia się na istotnych częściach i opisuje je.
3. Następnie przeprowadza logiczne rozumowanie w celu uzyskania wstępnej odpowiedzi.
4. Na koniec przedstawia zwięzłe podsumowanie odpowiedzi.

Ten etapowy proces rozumowania jest niewidoczny dla użytkownika, co pozwala modelowi zarządzać własnym procesem myślowym i skuteczniej dostosowywać się do złożonych zadań. Aby jeszcze bardziej poprawić możliwości rozumowania modelu, LLaVA-o1 wykorzystuje technikę zwaną „wyszukiwaniem wiązki na poziomie etapu”. Metoda ta generuje wiele kandydujących wyników na każdym etapie rozumowania, wybierając najlepszego kandydata do kontynuowania procesu. Podejście to jest bardziej elastyczne i wydajne niż tradycyjne metody, w których model generuje kompletne odpowiedzi przed wybraniem najlepszej.

Naukowcy uważają, że to ustrukturyzowane podejście i wykorzystanie etapowego wyszukiwania wiązki sprawi, że LLaVA-

o1 będzie potężnym narzędziem do rozwiązywania złożonych zadań rozumowania, co czyni go znaczącym postępem w dziedzinie sztucznej inteligencji. Rozwój ten może potencjalnie zrewolucjonizować sposób interakcji ze sztuczną inteligencją, szczególnie w dziedzinach wymagających złożonego rozumowania, takich jak opieka zdrowotna, finanse i usługi prawne. LLaVA-o1 stanowi obiecujący krok w kierunku bardziej inteligentnych i elastycznych systemów sztucznej inteligencji.

Artykuł przetłumaczono przy pomocy AI □