

Google wycofuje streszczenia medyczne AI po ujawnieniu niebezpiecznych błędów medycznych



Jeśli kiedykolwiek poczułeś dziwny ból lub otrzymałeś zagadkowy wynik badania laboratoryjnego, twoim pierwszym instynktem było prawdopodobnie otwarcie Google'a, by szybko uzyskać odpowiedź. Ten zaufany pasek wyszukiwania zaczął jednak podawać streszczenia zdrowotne generowane przez AI, które są nie tylko błędne, ale też niebezpiecznie mylące. Google został teraz zmuszony do cichego usunięcia niektórych z tych „Przeglądów AI” po tym, jak dochodzenie dziennika The Guardian ujawniło, że dostarczały one niedokładnych informacji medycznych, które mogą narażać użytkowników na poważne ryzyko. Ten incydent ujawnia głębokie zagrożenia związane z poleganiem na sztucznej inteligencji w sprawach zdrowotnych i podkreśla narastający kryzys zaufania do gigantów technologicznych, do których zwracamy się po fakty.

Najjaskrawsza porażka dotyczyła testów funkcji wątroby. Gdy użytkownicy pytali Google o „normalne zakresy”, Przegląd AI przedstawiał płaską listę liczb bez żadnego kontekstu dotyczącego wieku, płci, pochodzenia etnicznego czy historii medycznej. Eksperci medyczni zaalarmowali, że jest to niebezpieczne. Normalny wynik dla 20-latką może być sygnałem ostrzegawczym dla 50-latką, ale AI brakuje subtelności, by dostrzec różnicę.

Fałszywe poczucie bezpieczeństwa

Niebezpieczeństwo jest dwojakie. Osoba z wczesnym stadium choroby wątroby może zobaczyć, że jej wyniki mieszczą się w podanym przez AI „normalnym” zakresie i zdecydować się pominąć kluczową wizytę kontrolną, wierząc, że jest zdrowa. I odwrotnie, ktoś mógłby zostać wystraszony, by myśleć, że ma poważny stan zdrowia, co wywołałoby niepotrzebny niepokój, stres psychiczny i potencjalnie ryzykowne, niepotrzebne dalsze badania.

Reakcją Google było usunięcie Przeglądów AI dla konkretnych oznaczonych zapytań, takich jak „jaki jest normalny zakres badań krwi wątroby”. Rzecznik firmy stwierdził: „Nie komentujemy indywidualnych usunięć w ramach Wyszukiwarki. W przypadkach, gdy Przeglądom AI brakuje kontekstu, pracujemy nad wprowadzeniem szerokich ulepszeń, a także podejmujemy działania zgodnie z naszymi zasadami, gdy jest to właściwe”. Naprawa była jednak powierzchowna. British Liver Trust odkrył, że samo przeformułowanie pytania na „zakres referencyjny lft” (ang. *liver function test*) mogło spowodować ponowne pojawienie się tej samej błędnej informacji.

Iluzja autorytetu

Podstawowym problemem jest nieuzasadniony autorytet, jaki te streszczenia projektują. Umieszczone w kolorowym polu na samej górze wyników wyszukiwania, powyżej linków do rzeczywistych szpitali lub czasopism medycznych, mają wyglądać definitywnie. Jesteśmy przyzwyczajeni do ufania najwyższemu wynikowi. Jak wyjaśniła Vanessa Hebditch, dyrektor ds. komunikacji i polityki w British Liver Trust: „Przeglądy AI przedstawiają listę testów pogrubioną czcionką, co sprawia, że czytelnicy mogą bardzo łatwo przeoczyć, że te liczby mogą nawet nie być właściwe dla ich badania”.

Hebditch przywitała z zadowoleniem usunięcia, ale ostrzegła

przed większym problemem. „Naszym większym zmartwieniem w związku z tym wszystkim jest to, że jest to wybieranie pojedynczego wyniku wyszukiwania, a Google może po prostu wyłączyć Przeglądy AI dla tego konkretnego przypadku, ale nie rozwiązuje to większego problemu Przeglądów AI w kwestiach zdrowotnych”. Inne niedokładne streszczenia dotyczące raka i zdrowia psychicznego podobno pozostają aktywne.

Wada jest wbudowana w system. Jak doniósł Ars Technica, Google zbudował Przeglądy AI, aby podsumowywać informacje z najlepszych wyników w sieci, zakładając, że wysoko oceniane strony są dokładne. To skierowało treści optymalizowane pod kątem SEO (ang. *Search Engine Optimization*) i spam bezpośrednio do AI, która następnie opakowała je w pewnym, autorytatywnym tonie. Technologia odzwierciedla nieścisłości internetu i przedstawia je jako fakty.

Ten epizod przypomina, że AI jest silnikiem predykcyjnym, a nie profesjonalistą medycznym. Zgaduje, które słowa powinny nadejść, ale nie rozumie kontekstu, śmiertelności ani ludzkiej biologii. Na razie, gdy twoje zdrowie jest na szali, najbezpieczniejszą drogą jest przewinięcie poniżej kolorowego pudełka robota i kliknięcie linku od prawdziwej, odpowiedzialnej instytucji medycznej. Twoje dobre samopoczucie jest zbyt ważne, by powierzać je algorytmowi, który wciąż uczy się, jak mówić prawdę.