

Narasta problem wykorzystania AI do pisania prac naukowych

29 kwietnia 2025

Coraz lepiej widać, jak duża jest skala wątpliwości dotyczących wykorzystania AI w tworzeniu publikacji naukowych – powiedział PAP Artur Strzelecki, prof. Uniwersytetu Ekonomicznego w Katowicach. Kolejne przykłady nadużyć w tym zakresie przedstawiło czasopismo „Nature”, m.in. związane z „cichą korektą” tekstu.

0 nowych przypadkach ukrytego wykorzystania narzędzia ChatGPT w pisaniu prac naukowych informuje „Nature”. Autorka artykułu Diana Kwon pisze o wykryciu charakterystycznych zwrotów chatbotów w setkach publikacji, wstawionych bez ujawnienia tego faktu (bez deklaracji o korzystaniu z ChatuGPT) – co podważa wiarygodność procesu recenzji.

„Nature” powołuje się m.in. na ustalenia dr. hab. Artura Strzeleckiego, prof. Uniwersytetu Ekonomicznego w Katowicach (UEK). W grudniu 2024 r. na łamach „Learned Publishing” prof. Strzelecki opisał wyniki swoich analiz zasobów Google Scholar – jednej z największych baz publikacji naukowych. Wykazał on, że nieoznaczone fragmenty AI pojawiają się nawet w czasopismach z najwyższą rangą cytowań, a część z nich zdążyła zyskać kolejne odniesienia, co podkreśla systemowy charakter problemu. „Coraz bardziej też widać, że trzeba robić wewnętrzne szkolenia, warsztaty; uczyć, że to jest narzędzie” – prof. Strzelecki

Badacz sprawdzał wówczas, czy „nienaturalne” sformułowania charakterystyczne dla Chatu GPT (np. „Mam nadzieję, że to pomoże!”, „Oczywiście! Tutaj jest/są ...”, „Cieszę się, że mogę pomóc!” itp.) znajdują się w anglojęzycznych publikacjach naukowych. Przeanalizował tylko te artykuły, w których nie deklarowano korzystania z programu Chat GPT w pisaniu pracy.

Wśród fraz, które w publikacjach zaczęły się pojawiać za sprawą użycia ChatuGPT, były m.in. takie zwroty jak: „Według mojej wiedzy na dzień ostatniej aktualizacji...” („as of my last knowledge update”), „Jako model językowy AI” („as an AI language model”) czy „I don’t have access to real-time...” („nie mam dostępu do aktualnych...”). W publikacjach naukowych nierzadko kopiowano też treść interfejsu: pod każdą odpowiedzią ChatuGPT był dawniej przycisk „regenerate response”, czyli polecenie, by chat ponownie opracował swoją odpowiedź. Całkiem wielu nieuwważnych naukowców skopioowało odpowiedź AI wraz z jej „stopką” i nie przeczytało do końca tekstu wysyłanego do wydawnictwa.

Bezmyślne przekopiowanie do artykułu naukowego fragmentu wypowiedzi ChatuGPT nie zostało też zauważone przez współautorów, recenzentów ani redaktorów czasopisma. Biorąc pod uwagę tylko bardziej prestiżowe, recenzowane czasopisma naukowe z bazy Scopus (obecne w pierwszym i drugim kwartylu wskaźnika CiteScore), frazy ChatuGPT prof. Strzelecki odnalazł w 89 artykułach.

Obecnie temat bezrefleksyjnego wykorzystania AI wrócił dzięki „Nature”. Redakcja ta zwróciła na przykład uwagę na zjawisko „cichej korekty” – usuwanie przez wydawców chatbotowych fraz z już opublikowanych artykułów bez oficjalnych errat.

„Niektóre czasopisma same się zaczęły orientować, że przepuściły poważne błędy, że nikt należycie nie sprawdził tekstu: na etapie redakcji, recenzji, copyeditingu. »To my to po cichu poprawimy i błędu nie będzie«. A jednak kilka osób – w tym ja – dostrzegło te błędy. Dość szybko znajdowałem oryginały publikacji je zawierających i archiwizowałem. Potem mogłem zobaczyć, że wydawnictwo faktycznie wprowadziło w tekście zmianę bez żadnego informowania, że zmiana zaszła i na czym polegała” – opowiada prof. Strzelecki, komentując zjawisko opisane w „Nature”.

Na łamach „Nature” potwierdzono też, że w setkach publikacji z

różnych dziedzin pojawiają się charakterystyczne frazy, takie jak „jako model językowy AI” czy wspomniane „regenerate response”, zdradzające udział ChatuGPT w pisaniu (a może raczej generowaniu) tekstu. „Nature” wskazało też na ryzyko kaskady błędów – nieujawnione wtręty AI są już bowiem cytowane w literaturze, co może multiplikować niezweryfikowane informacje.

„Jeśli już powstają prace przy użyciu AI i zawierają wskazujące na to fragmenty – to one są później cytowane przez kolejnych naukowców, którzy to, co zostało napisane, uznają za rzetelne informacje, a niekoniecznie powinni. Sztuczna inteligencja miewa halucynacje i może coś sama od siebie coś dołożyć, wymyślić” – potwierdza prof. Strzelecki.

Cytowany w „Nature” Alex Glynn, ekspert w zakresie umiejętności badawczych i komunikacji z Uniwersytetu Louisville w Kentucky zaznaczył, że wspomniane zmiany pojawiły się w „mniejszości czasopism”. Jednak biorąc pod uwagę, że prawdopodobnie istnieje również wiele przypadków, w których autorzy używali sztucznej inteligencji bez pozostawiania oczywistych śladów – Glynn wyraził zaskoczenie, „jak wiele tego jest”.

Podobnie jak prof. Strzelecki, również Glynn znalazł setki artykułów ze śladami wykorzystania AI. Jego internetowy program do rejestrowania takich przypadków obejmuje już ponad 700 artykułów. Redaktorzy „Nature” skontaktowali się z wydawcami niektórych artykułów, takimi jak Springer Nature, Taylor & Francis, IEEE i inni, którzy zostali wskazani przez Glynna i prof. Strzeleckiego. Redakcje udzieliły informacji, że wszystkie oflagowane artykuły są obecnie sprawdzane. Powoływały się również na swoje zasady dotyczące AI – które w niektórych przypadkach nie wymagają ujawniania użycia AI (np. nie trzeba sygnalizować zmian wprowadzonych z powodów redakcyjno-językowych).

Prof. Strzelecki podkreślił, że w środowisku osób

zainteresowanych rośnie świadomość na temat skali problemu bezrefleksyjnego wykorzystania narzędzi AI w pisaniu publikacji naukowych. Trwa też dyskusja, kiedy i w jakiej formie autorzy oraz recenzenci powinni ujawniać wykorzystanie narzędzi AI. Polityki wydawców są rozbieżne, lecz zgadzają się co do jednego: pełna transparentność jest konieczna, by chronić zaufanie do literatury naukowej.

„Coraz bardziej też widać, że trzeba robić wewnętrzne szkolenia, warsztaty; uczyć, że to jest narzędzie” – mówił prof. Strzelecki. „Tak, należy go używać – ale trzeba wiedzieć, w jaki sposób i do czego. Zasadna jest na przykład pomoc językowa dla autorów, których język ojczysty nie jest angielskim. W ich przypadku wykorzystanie narzędzi AI jako pomoc przy korekcie gramatyki, pisowni, słownictwa itd. wydaje się sensowna. W tego typu przypadkach duże wydawnictwa naukowe idą w takim kierunku, żeby nie trzeba było ujawniać faktu korzystania z tego rodzaju wsparcia. W tym kontekście nie stawiamy więc pytania: czy będę używał sztucznej inteligencji? Tylko: w jaki sposób?” – powiedział.

Profesor UEK zauważył, że obecnie wydawnictwa naukowe idą w kierunku ujawniania faktu, że autor publikacji wykorzystał AI do analizy tekstu lub przetwarzania danych, które samodzielnie zebrał. „Wydawnictwa oczekują, żeby w publikacji zaznaczyć, że zostało użyte narzędzie sztucznej inteligencji; a jeszcze lepiej – wskazać, w jakim dokładnie zakresie było ono zastosowane” – mówił.

Jednocześnie zastrzegł, że nie da się kontrolować sposobu użycia sztucznej inteligencji w pisaniu prac naukowych. „Praktycznie już teraz są to narzędzia powszechnego użytku, które wchodzą do użycia na ogromną skalę. Od niedawna, wyszukując informacji w Google, zalogowani użytkownicy na dzień dobry dostają podsumowanie zrobione przez AI. Nie dostajemy autentycznych wyników wyszukiwania, tylko dane przetworzone przez AI”.

Autorstwo: PAP

Na podstawie: [Nature.com](https://www.nature.com)

Źródło: [NaukawPolsce.pl](https://naukawpolsce.pl)