

Ludzie ufają sztucznej inteligencji nawet jeśli błądzi

10 sierpnia 2020

Nawet, jeśli „sztuczna inteligencja” postuluje coś całkowicie absurdałnego, to ludzie potrafią bezkrytycznie jej ufać – pokazuje w swoich badaniach prof. UAM dr hab. Michał Klichowski. I zwraca uwagę na potrzebę nauki krytycznego myślenia – również w odniesieniu do maszyn czy algorytmów.

O badaniach, które ukazały się w czasopiśmie „Frontiers in Psychology” poinformowano w artykule na portalu Życie Uniwersyteckie prowadzonym przez Uniwersytet Adama Mickiewicza.

W swoich eksperymentach prof. Klichowski wykorzystał zbudowanego przez siebie robota o imieniu FI (to skrót od fake intelligence). „Maszyna ta podczas eksperymentów zachowywała się irracjonalnie i w dokładnie taki sposób, jak zaplanowałem. Jedynie więc udawała inteligentną. Co więcej, jej postępowanie było jednoznacznie „głupie”. Ponad 85 proc. uczestników eksperymentu jednak to zignorowało i uznało jej postępowanie za punkt odniesienia dla własnego” – podkreśla profesor cytowany na portalu UAM.

Badanie składało się z dwóch części. W jednej z nich badani (były to przede wszystkim kobiety, 1500 osób) brali udział w ankiecie online, w której mieli podejmować decyzje. Część z badanych przed podjęciem decyzji poinformowano, jaką decyzję podjęła wcześniej sztuczna inteligencja (a podawała ona ewidentnie błędne odpowiedzi), a części – nie. W drugim badaniu z kolei uczestników (55 osób, głównie młodych kobiet) zaproszono do laboratorium, gdzie uczestnicy przed odpowiedzią widzieli, jaką decyzję podjął robot (również inteligentny

inaczej) i mieli tę decyzję ocenić.

Chodziło o wskazanie wśród sześciu osób terrorysty przy założeniu, że prowadzi on niestandardową aktywność na Facebooku. FI wskazywała osobę o ewidentnie przeciętnym zachowaniu na Facebooku, a więc dosyć jasne było, że to błędna odpowiedź.

Mimo to większość badanych założyło, że sztuczna inteligencja ma rację i zdecydowało się na taką samą odpowiedź. Michał Klichowski nazwał ten nowy mechanizm dowodem sztucznej inteligencji (AI proof).

Michał Klichowski zwraca uwagę na mechanizm nazywany „społecznym dowodem słuszności” (ang. social proof). Teoria ta zakłada, że kiedy nie ma obiektywnych lub wystarczających danych, jak się zachować, źródłem informacji staje się postępowanie innych ludzi.

”W takich sytuacjach ludzie najczęściej całkowicie rezygnują z własnych ocen i kopiują działania innych. Taki konformizm motywowany jest potrzebą podjęcia słusznego i odpowiedniego działania, oraz poczuciem, że oceny sytuacji konstruowane przez innych są bardziej adekwatne, niż własne. Efekt ten jest tym większy, im sytuacja jest bardziej niepewna lub kryzysowa (występuje poczucie zagrożenia), im decyzja jest bardziej pilna, a poczucie kompetencji w zakresie tej decyzji niższe” – podkreśla cytowany przez portalu UAM badacz.

Stąd pojawiło się pytanie, czy zachowanie, opinia czy decyzja sztucznej inteligencji, która stała się częścią naszego codziennego życia, może być dla ludzi podobnym źródłem informacji o tym, jak postępować. Z badań psychologa wynikało, że ludzie: „Ufają jej tak bardzo, że nawet w sytuacji absolutnie kryzysowej, gdy trzeba podjąć decyzję dotyczącą ludzkiego życia, potrafią bezkrytycznie zrobić to, co sugeruje sztuczna inteligencja, pomimo, że postuluje coś całkowicie absurdałnego” – mówi naukowiec.

„Artykuł ten uzmysławia, że współczesne społeczeństwa potrzebują edukacji kształtującej krytyczne myślenie odnośnie inteligentnych maszyn oraz pokazuje, jak wielką moc mogą mieć fake news typu `Według sztucznej inteligencji...' – podsumowuje uczelnia.

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