

ŁAGODZENIE AWERSJI WOBEC SZTUCZNEJ INTELIGENCJI (AI) I POBUDZANIE ROZTROPNOŚCI KONSUMENTÓW: EKSPŁORACJA TAKTYK OZNACZANIA AUTORSTWA TREŚCI REKLAMOWYCH GENEROWANYCH PRZEZ AI

Wyobraź sobie, że przeglądasz media społecznościowe i pojawia się reklama ubrań. Widzisz na niej olśniewającą modelkę ubraną w piękne, zielone spodnie, które wydają się idealne dla Ciebie. Klikasz w reklamę, oglądasz wideo i decydujesz się na zakup tych spodni. Ale co by było, gdybyś wiedział/a, że reklama została stworzona w całości przez sztuczną inteligencję (AI)? Czy zmieniłoby to Twoją decyzję? Badania sugerują, że tak. Wiele osób reaguje negatywnie na informację, że treści zostały stworzone przez AI, a niektórzy nawet posuwają się do publikowania obraźliwych i krytycznych komentarzy w Internecie, wywołując efekt domina w postaci lawiny negatywnych reakcji.

Takie negatywne reakcje wynikają z fenomenu znanego jako "awersja wobec algorytmów" (ang. *algorithm aversion*), gdzie konsumenci nie ufają treściom generowanym przez AI. Nawet jeśli reklama jest przekonująca, sceptycyzm może przeważać nad jej mocą perswazyjną. Aby uniknąć takich reakcji, wiele firm decyduje się nie ujawniać udziału AI. Jednak ukrywanie tej informacji budzi problemy etyczne. Transparentność i dobre praktyki wymagają, aby firmy informowały o udziale AI, dając konsumentom możliwość podejmowania świadomych decyzji.

Nasz projekt bada, czy i jak można ujawniać autorstwo AI w sposób minimalizujący awersję do algorytmów, jednocześnie zachowując skuteczność reklamy i budując konsumencką roztropność. Jest to cel ryzykowny, ponieważ łączy pozornie sprzeczne priorytety: etyczną transparentność, sukces biznesowy i społecznie odpowiedzialne zachowania. Niemniej jednak jest to konieczne dla rozwoju etycznego, biznesowego i społecznego.

W ramach tego projektu przeprowadzimy serię eksperymentów, aby przetestować różne taktyki ujawniania autorstwa AI oraz okoliczności, w których mogą one działać lub zawodzić. Na przykład, czy reakcja na reklamę stworzoną przez AI byłaby inna, gdyby informacja o tym została przedstawiona humorystycznie, za pomocą hasła w stylu: „*Stworzone przez AI, zaakceptowane przez ludzi – kto by pomyślał, że AI ma taki dobry gust?*”? A jak by zadziałała bardziej zhumanizowana forma ujawnienia, zmniejszająca dystans między AI a konsumentami: „*Poznaj Awę, naszą kreatywną asystentkę AI, która współprojektowała tę reklamę. Zajęło jej to 10 przerw na kawę i niezliczone godziny analizy najnowszych trendów reklamowych, aby stworzyć te reklamy*”. Co by się stało, gdyby te taktyki połączono, wprowadzając humor do zhumanizowanej narracji: „*Stworzone przez naszą asystentkę AI. Widzisz jakieś błędy? Cóż, każdy je popełnia! Pamiętaj tylko, że trudno rysować wirtualnymi rękami*”.

Nasze eksperymenty mają na celu odkrycie taktyk oznaczania autorstwa, które jednocześnie chronią konsumentów i zachowują skuteczność reklam generowanych przez AI. Zbadamy, jak humor, narracja, humanizacja i ich kombinacje wpływają na konsumentów. Taktyki te zostaną przetestowane na różnych kategoriach produktowych, takich jak moda, podróże, finanse czy dobra związane ze zdrowiem (np. suplementy diety), ponieważ reakcje konsumentów na ujawnienia mogą się różnić w zależności od produktu. Przyjrzymy się też różnym kontekstom reklamowym, od lekkich, takich jak moda, po poważniejsze, jak opieka zdrowotna czy reklamy polityczne, gdzie transparentność może odgrywać bardziej kluczową rolę.