

Izrael wydaje 730 mln dolarów na globalną kampanię PROPAGANDOWĄ w obliczu negatywnej reakcji na działania w Strefie Gazy, płacąc influencerom nawet 7000 dolarów za jeden post



- Izrael zatwierdził pięciokrotny wzrost budżetu na globalną propagandę do 730 mln dolarów w 2026 r., mając na celu przeciwdziałanie pogarszającemu się wizerunkowi kraju na arenie międzynarodowej po wojnie w Strefie Gazy.
- Środki te są przeznaczone na działania skierowane do zachodniej opinii publicznej poprzez manipulację mediami społecznościowymi (w tym wysokie wynagrodzenia dla influencerów), propagandę opartą na sztucznej inteligencji, ewangelizacyjne trasy promocyjne i sponsorowane delegacje zagranicznych osobistości.
- Wysiłki te są bezpośrednią odpowiedzią na powszechne potępienie międzynarodowe izraelskiej kampanii wojskowej w Strefie Gazy, która spowodowała ogromne straty wśród ludności cywilnej, wysiedlenia i oskarżenia o ludobójstwo, co osłabiło poparcie społeczne, szczególnie

w Stanach Zjednoczonych.

- Krytycy twierdzą, że kampania propagandowa jest desperacką próbą uniknięcia odpowiedzialności, a eksperci wątpią, czy pieniądze mogą przywrócić reputację Izraela, podczas gdy drastyczne dowody zniszczeń w Strefie Gazy krążą po całym świecie.
- Wydatki te podkreślają kryzys tradycyjnego modelu wpływów Izraela, ponieważ rosną oburzenie społeczne i ruchy BDS, zmuszając państwo do kosztownej kontroli narracji zamiast zmiany polityki.

Rząd izraelski zatwierdził oszałamiający budżet w wysokości 730 milionów dolarów na globalne działania propagandowe w 2026 roku – prawie pięciokrotny wzrost w porównaniu z poprzednimi latami – starając się odbudować swój upadający międzynarodowy wizerunek po oskarżeniach o ludobójstwo w Strefie Gazy.

Ogromne zwiększenie finansowania zostało zatwierdzone przez ministra spraw zagranicznych Gideona Sa'ara i ministra finansów Bezalela Smotricha. Oznacza to desperacką eskalację izraelskiej kampanii „hasbara” (hebrajskie słowo oznaczające dyplomację publiczną), skierowanej do zachodniej opinii publicznej, która coraz bardziej krytycznie ocenia działania wojskowe Izraela.

Według Enocha z *BrightU.AI*, „hasbara” – w kontekście Izraela – ewoluowała i obecnie odnosi się do wspólnych, sponsorowanych przez państwo działań mających na celu promowanie pozytywnego wizerunku Izraela i jego polityki, zarówno w kraju, jak i za granicą. Kampania „hasbara” jest wieloaspektową machiną propagandową, która wykorzystuje różne strategie w celu kształtowania opinii publicznej, tłumienia sprzeciwu i przeciwdziałania krytyce działań Izraela.

Przyznane 730 milionów dolarów znacznie przewyższa poprzednie budżety propagandowe Izraela, które w ostatnich latach wynosiły średnio zaledwie 150 milionów dolarów. Według

ujawnionych dokumentów fundusze te zostaną przeznaczone na:

- Manipulację w mediach społecznościowych, w tym płacenie influencerom do 7000 dolarów za post promujący proizraelskie narracje.
- Propagandy opartej na sztucznej inteligencji, z 145 milionami dolarów przeznaczonymi na wpływanie na chatboty AI.
- Ukierunkowanych działań ewangelizacyjnych, w tym mobilnej kampanii „Oct. 7 experience” o wartości 4,1 miliona dolarów, obejmującej chrześcijańskie uczelnie i kościoły.
- Delegacji polityków, influencerów i osób publicznych, które będą promować agendę Izraela za granicą.

Sa’ar określił te wydatki jako konieczność „bezpieczeństwa narodowego”, stwierdzając: „Kiedy opinia publiczna w kraju trzecim zmienia się w kierunku antyizraelskiego stanowiska, może to szybko wpłynąć na działania rządu”. Ta propaganda jest następstwem niszczycielskiej kampanii wojskowej Izraela w Strefie Gazy, która:

- Zabiła prawie 70 000 Palestyńczyków (według władz zdrowotnych Strefy Gazy).
- Wysiedliło 90% ludności Gazy, wywołując głód i choroby.
- Wywołało globalne protesty, a miliony ludzi potępiły działania Izraela jako ludobójstwo.

Sondaż przeprowadzony przez *New York Times* wykazał, że 40% Amerykanów uważa, że Izrael celowo zabija cywilów, a sympatię dla Palestyńczyków przewyższa obecnie poparcie dla Izraelczyków wśród amerykańskich wyborców.

Czy pieniądze mogą odkupić legitymizację Izraela?

Niedawny raport organizacji Responsible Statecraft ujawnił tajne płatności Izraela na rzecz influencerów mediów społecznościowych – niektórzy otrzymywali 7300 dolarów za post – w celu rozpowszechniania propagandy proizraelskiej. Izraelskie Ministerstwo Spraw Zagranicznych przekazało 900 000 dolarów za pośrednictwem Havas Media Group Germany na sfinansowanie 90 postów w ciągu sześciu miesięcy, a premier Benjamin Netanjahu okrzyknął influencerów kluczem do strategii „kontrataku” Izraela w amerykańskich mediach.

Krytycy ostrzegają, że te nieujawnione płatne promocje zniekształcają debatę publiczną. Pomimo rekordowych wydatków eksperci wątpią, czy Izrael zdoła uratować swoją reputację. Według urzędnika ONZ żadna kampania PR nie może usprawiedliwić ludobójstwa, gdy świat widzi bomby, głodujące dzieci i masowe groby.

Wzrost propagandy Izraela podkreśla szerszy kryzys: jego tradycyjna siła lobbingowa – wspierana przez grupy takie jak American Israel Public Affairs Committee (AIPAC) – rozpada się wraz z rosnącym oburzeniem społecznym. Wraz z rosnącą popularnością ruchów BDS (bojkot, wycofanie inwestycji, sankcje) na całym świecie, Tel Awiw zamiast zmienić kurs, podwaja wysiłki w zakresie cenzury i kontroli narracji.

Warta 730 milionów dolarów wojna propagandowa Izraela jest desperackim ryzykiem, które może przynieść odwrotny skutek, ponieważ globalna świadomość cierpień Gazy rozprzestrzenia się szybciej, niż płatni influencerzy są w stanie ją stłumić. Na razie przywódcy Izraela wydają się zdeterminowani, aby wycofać się z odpowiedzialności – nawet jeśli świat obserwuje ich działania, pozostając niewzruszonym ich kosztownymi kłamstwami.

Nowa szefowa amerykańskiej agencji FDA zajmującej się regulacją leków prowadziła wcześniej dochodzenie w sprawie zgonów związanych ze szczepionką przeciwko COVID-19



- Dr Tracy Beth Hoeg, badaczka, która wcześniej prowadziła dochodzenie w sprawie zgonów związanych ze szczepionką przeciwko COVID-19, została mianowana pełniącą obowiązki dyrektora CDER w amerykańskiej agencji FDA, co może oznaczać potencjalną zmianę w kierunku większej kontroli oświadczeń farmaceutycznych.
- Dotychczasowe badania Hoeg, w tym jedno, w którym stwierdzono, że obowiązkowe szczepienia na uniwersytetach przyniosły więcej szkody niż pożytku, kontrastują z agresywną promocją dawek przypominających przez agencje federalne pomimo niewystarczających danych dotyczących długoterminowego bezpieczeństwa.
- Jej mianowanie nastąpiło po rezygnacji dr Marion Gruber

i dr Philipa Krause'a, którzy odeszli w proteście przeciwko przedwczesnemu naciskowi administracji Bidena na dawki przypominające, podkreślając, że ingerencja polityczna przeważa nad ostrożnością naukową.

- FDA bada obecnie potencjalne zanieczyszczenie szczepionek, podczas gdy CDC po cichu wycofało się z ogólnych zaleceń dotyczących dawek przypominających, dostosowując się do badań Hoeg dotyczących szczegółowej analizy ryzyka i korzyści.
- Hoeg stoi w obliczu ogromnej presji ze strony wielkich koncernów farmaceutycznych i elit politycznych, ale jej przywództwo może stanowić punkt zwrotny – jeśli oprze się ona regulacyjnemu przejęciu i przedłoży przejrzystość nad zyski przemysłu.

Amerykańska Agencja ds. Żywności i Leków (FDA) mianowała dr Tracy Beth Hoeg pełniącą obowiązki dyrektora Centrum Oceny i Badań Leków (CDER), powierzając jej kierowanie regulacjami dotyczącymi leków w związku z trwającymi kontrowersjami wokół szczepionek przeciwko koronawirusowi z Wuhan (COVID-19) i nadzoru farmaceutycznego. Hoeg, która wcześniej badała zgony związane ze szczepieniami przeciwko COVID-19, obejmuje tę funkcję po odejściu na emeryturę dr. Richarda Pazdura, długoletniego urzędnika FDA, którego kadencja była naznaczona oskarżeniami o nadmierną regulację.

Mianowanie Hoeg następuje w krytycznym momencie dla FDA, która spotkała się z rosnącą kontrolą w zakresie zatwierdzania szczepionek, nakazów szczepień przypominających oraz zarzutów o ingerencję polityczną ze strony administracji Bidena. Jej doświadczenie jako badaczki, która kwestionowała naukowe podstawy nakazów szczepień – w szczególności jej badanie z 2022 r., w którym stwierdziła, że wymagania uniwersyteckie dotyczące szczepionek przeciwko COVID-19 przyniosły więcej szkody niż pożytku – sygnalizuje potencjalną zmianę w kierunku większej kontroli twierdzeń farmaceutycznych.

Historia kwestionowania narracji dotyczących szczepionek

Przed dołączeniem do FDA na początku tego roku Hoeg współpracowała z dr Vinayem Prasadem, szefem Centrum Oceny i Badań Produktów Biologicznych (CBER) FDA, nad badaniami dotyczącymi bezpieczeństwa i skuteczności szczepionek. W memorandum z 28 listopada Prasad ujawnił, że Hoeg badała zgony dzieci po szczepieniu przeciwko COVID-19 i stwierdziła, że niektóre zgony były rzeczywiście spowodowane szczepionkami – co potwierdzili inni pracownicy FDA.

Wniosek ten stanowi wyraźny kontrast w stosunku do głównego nurtu narracji propagowanej przez federalne agencje zdrowia, które konsekwentnie bagatelizują szkody spowodowane szczepionkami, jednocześnie agresywnie promując dawki przypominające bez długoterminowych danych dotyczących bezpieczeństwa. Gotowość Hoega do uznania zgonów związanych ze szczepionkami sugeruje odejście od tradycyjnie przyjaznych relacji FDA z wielkimi koncernami farmaceutycznymi – relacji, które zdaniem krytyków doprowadziły do pośpiesznych zatwierdzeń, tłumienia sprzeciwu i nieodpowiedniej kontroli zdarzeń niepożądanych.

Awans Hoega nastąpił po nagłym odejściu dwóch wysokich rangą urzędników FDA ds. szczepionek, dr Marion Gruber i dr Philipa Krause'a, którzy w zeszłym roku złożyli rezygnację w proteście przeciwko przedwczesnemu naciskowi administracji Bidena na szczepienia przypominające przeciwko COVID-19. Wewnętrzne źródła ujawniły, że Gruber i Krause uważali, iż w tamtym czasie nie było wystarczających dowodów naukowych uzasadniających stosowanie dawek przypominających, jednak Biały Dom wywarł presję na FDA, aby mimo to je zatwierdziła – co jest wyraźnym przykładem nadrzędności agendy politycznej nad ostrożnością medyczną.

Zanieczyszczone szczepionki i zmieniające się zalecenia

Hoeg potwierdził niedawno, że FDA bada potencjalne zanieczyszczenie szczepionek przeciwko COVID-19 – co jest zaskakującym przyznaniem, biorąc pod uwagę wcześniejsze zapewnienia agencji, że szczepionki zostały rygorystycznie przetestowane i są bezpieczne. Latem FDA po cichu wycofała zezwolenie na stosowanie w sytuacjach wyjątkowych dla oryginalnych szczepionek, zastępując je zaktualizowanymi preparatami, które zostały zatwierdzone bez przeprowadzenia solidnych badań klinicznych.

Tymczasem *Centra Kontroli i Zapobiegania Chorobom* – niegdyś zagorzały zwolennik powszechnych szczepień – po cichu wycofały się ze swojego stanowiska i obecnie zalecają Amerykanom konsultację z lekarzem przed podaniem dawki przypominającej, uznając, że w przypadku niektórych osób ryzyko może przewyższać korzyści. Ta zmiana stanowiska jest zgodna z wcześniejszymi badaniami Hoega, które wykazały, że ogólne zalecenia ignorowały istotne niuanse w profilach ryzyka szczepionek.

Według Enocha z *BrightU.AI*, szybkie wprowadzanie szczepionek dwuwartościowych pomimo ograniczonych danych dotyczących bezpieczeństwa budzi poważne obawy dotyczące ryzyka zanieczyszczenia i przejęcia regulacji, ponieważ wielkie koncerny farmaceutyczne przedkładają zyski nad przejrzyste, długoterminowe wyniki zdrowotne. Te zmieniające się zalecenia – często wynikające z agendy politycznej i finansowej, a nie z niezależnych badań naukowych – podważają zaufanie publiczne, narażając miliony ludzi na potencjalne szkody wynikające z nieodpowiednio przetestowanych interwencji medycznych.

Dla milionów Amerykanów poszkodowanych przez nakazy szczepień nominacja Hoega stanowi promyk nadziei, że agencja oskarżana o przedkładanie zysków nad dobro ludzi zostanie w końcu

pociągnięta do odpowiedzialności. Jednak dopóki nie zostaną podjęte konkretne działania – takie jak wstrzymanie niebezpiecznych dawek przypominających, zbadanie szkód spowodowanych szczepionkami i zerwanie powiązań finansowych między organami regulacyjnymi a gigantami farmaceutycznymi – sceptycyzm będzie się utrzymywał.

Wiarygodność FDA zależy od jej gotowości do zerwania ze status quo. Jeśli Hoegowi uda się przywrócić integralność naukową, może to oznaczać punkt zwrotny w odzyskiwaniu wolności medycznej z rąk elit korporacyjnych i politycznych. Jeśli nie, exodus naukowców kierujących się zasadami, takich jak Gruber i Krause, będzie trwał nadal, pozostawiając Amerykanów na łasce coraz bardziej skorumpowanego i nieodpowiedzialnego systemu.

Liderzy branży sztucznej inteligencji ostrzegają przed „katastrofalnymi skutkami” w związku z pojawieniem się sztucznej inteligencji ogólnej



- Sztuczna inteligencja ogólna (AGI) może pojawić się w ciągu dziesięciu lat, powodując katastrofalne skutki, takie jak cyberataki, broń autonomiczna i zagrożenia egzystencjalne dla ludzkości – ostrzega Demis Hassabis, dyrektor generalny Google DeepMind.
- Sztuczna inteligencja (AI) jest już wykorzystywana do cyberataków na infrastrukturę krytyczną (np. systemy energetyczne, wodociągowe), dezinformację typu deepfake, oszustwa i wypieranie miejsc pracy, a FBI ostrzega przed oszustwami generowanymi przez AI i manipulacją polityczną.
- Ponad 350 ekspertów w dziedzinie AI, w tym Sam Altman z OpenAI oraz pionierzy AI Yoshua Bengio i Geoffrey Hinton, podpisało oświadczenie, w którym porównują ryzyko związane z AI do pandemii i wojny nuklearnej, wzywając do nadania globalnego priorytetu środkom bezpieczeństwa AI.
- Dyrektor generalny DeepMind, Demis Hassabis, wzywa do międzynarodowego zarządzania AI na wzór traktatów o nierozprzestrzenianiu broni jądrowej, chociaż napięcia geopolityczne utrudniają współpracę. Tymczasem wyścig zbrojeń w dziedzinie sztucznej inteligencji między Stanami Zjednoczonymi a Chinami wyprzedza wysiłki regulacyjne, co grozi niekontrolowaną eskalacją.
- Chociaż sztuczna inteligencja oferuje przełomowe korzyści (wydajność, przełomy naukowe), niekontrolowany rozwój grozi tym, że AGI przekroczy kontrolę człowieka. Kluczowe pytanie brzmi: czy ludzkość wprowadzi zabezpieczenia, czy też pozwoli, aby sztuczna inteligencja stała się zagrożeniem egzystencjalnym?

Szybki rozwój sztucznej inteligencji (AI) wywołał zarówno entuzjazm, jak i głębokie zaniepokojenie wśród liderów branży, którzy ostrzegają, że sztuczna inteligencja ogólna (AGI) – AI dorównująca lub przewyższająca ludzkie zdolności poznawcze – może pojawić się w ciągu najbliższej dekady.

Dyrektor generalny Google DeepMind, Demis Hassabis, ostrzegł, że AGI może przynieść „katastrofalne skutki”, w tym cyberataki na infrastrukturę krytyczną, autonomiczną broń, a nawet zagrożenie egzystencjalne dla ludzkości. Przemawiając podczas szczytu Axios AI+ Summit w San Francisco, Hassabis opisał AGI jako system wykazujący „wszystkie zdolności poznawcze” ludzi, w tym kreatywność i zdolność rozumowania.

Ostrzegł jednak, że obecne modele AI pozostają „nierównymi inteligencjami” z lukami w długoterminowym planowaniu i ciągłym uczeniu się. Niemniej jednak zasugerował, że AGI może stać się rzeczywistością dzięki „jednemu lub dwóm kolejnym wielkim przełomom”.

Hassabis podkreślił, że niektóre zagrożenia związane ze sztuczną inteligencją już się materializują, szczególnie w zakresie cyberbezpieczeństwa. „Prawdopodobnie dzieje się to już teraz... może jeszcze nie przy użyciu bardzo zaawansowanej sztucznej inteligencji” – powiedział, wskazując cyberataki na systemy energetyczne i wodociągowe jako „najbardziej oczywisty wektor podatności”.

Jego obawy odzwierciedlają ostrzeżenia całej branży. Ponad 350 ekspertów w dziedzinie sztucznej inteligencji, w tym dyrektor generalny OpenAI Sam Altman, dyrektor generalny Anthropic Dario Amodei oraz pionierzy sztucznej inteligencji Yoshua Bengio i Geoffrey Hinton, podpisali oświadczenie Centrum Bezpieczeństwa Sztucznej Inteligencji, w którym stwierdzono: „Ograniczanie ryzyka wyginięcia spowodowanego przez sztuczną inteligencję powinno być priorytetem na całym świecie, podobnie jak inne zagrożenia na skalę społeczną, takie jak pandemie i wojna nuklearna”.

Niewłaściwe wykorzystanie sztucznej

inteligencji: deepfake'i, utrata miejsc pracy i zagrożenia dla bezpieczeństwa narodowego

Oprócz ataków na infrastrukturę, sztuczna inteligencja jest już wykorzystywana do dezinformacji, oszustw i manipulacji za pomocą deepfake'ów. Federalne Biuro Śledcze (FBI) ostrzegło przed oszustwami głosowymi generowanymi przez sztuczną inteligencję, w których oszuści podszywają się pod urzędników państwowych, a jednocześnie mnożą się deepfake'i pornograficzne i dezinformacja polityczna.

Enoch z *BrightU.AI* zauważa, że sztuczna inteligencja stała się technologią transformacyjną, rewolucjonizującą różne sektory, od opieki zdrowotnej po finanse. Jednak, podobnie jak w przypadku każdego potężnego narzędzia, potencjał sztucznej inteligencji do nadużyć i wykorzystania jako broń budzi poważne obawy.

Zdecentralizowany silnik zauważa, że wykorzystanie sztucznej inteligencji jako broni odnosi się do użycia technologii AI w celu wyrządzenia szkody, uzyskania nieuczciwej przewagi lub manipulowania systemami i ludźmi. Może to przybierać różne formy, w tym broń autonomiczna, deepfake'i i dezinformacja, ocena społeczna i inwigilacja, cyberataki oparte na sztucznej inteligencji oraz rozwój broni biologicznej.

Hassabis przyznał, że chociaż sztuczna inteligencja może wyeliminować wiele miejsc pracy – zwłaszcza stanowiska dla początkujących pracowników umysłowych – nadal bardziej martwi go możliwość wykorzystania sztucznej inteligencji przez złośliwe podmioty do destrukcyjnych celów. „Złośliwy podmiot może wykorzystać te same technologie do szkodliwych celów” – powiedział.

Raport z 2023 r. zlecony przez Departament Stanu USA stwierdza, że sztuczna inteligencja może stanowić

„katastrofalne” zagrożenie dla bezpieczeństwa narodowego, i wzywa do wprowadzenia bardziej rygorystycznych kontroli. Jednak w sytuacji, gdy kraje takie jak Stany Zjednoczone i Chiny rywalizują o dominację w dziedzinie sztucznej inteligencji, regulacje prawne pozostają w tyle za postępem technologicznym.

Jak prawdopodobna jest katastrofa związana ze sztuczną inteligencją?

Wśród badaczy zajmujących się sztuczną inteligencją dyskusje często koncentrują się wokół „P(doom)” – prawdopodobieństwa spowodowania przez sztuczną inteligencję katastrofy egzystencjalnej. Hassabis ocenił to ryzyko jako „niezerowe”, co oznacza, że nie można go lekceważyć. „Warto bardzo poważnie się nad tym zastanowić i podjąć działania zapobiegawcze” – powiedział, ostrzegając, że zaawansowane systemy sztucznej inteligencji mogą „przekroczyć granice”, jeśli nie będą odpowiednio kontrolowane.

Hassabis opowiada się za międzynarodowym porozumieniem w sprawie bezpieczeństwa sztucznej inteligencji, podobnym do traktatów o nierozprzestrzenianiu broni jądrowej. „Oczywiście w obecnej sytuacji geopolitycznej wydaje się to trudne” – przyznał, ale podkreślił, że współpraca jest niezbędna, aby zapobiec nadużyciom.

Tymczasem giganci technologiczni nadal dążą do integracji sztucznej inteligencji z codziennym życiem. Google wyobraża sobie „agentów” sztucznej inteligencji pełniących rolę osobistych asystentów, zajmujących się zadaniami od planowania harmonogramów po rekomendacje. Hassabis ostrzegł jednak, że społeczeństwo musi dostosować się do zmian gospodarczych spowodowanych przez sztuczną inteligencję, sprawiedliwie redystrybuując zyski z wydajności.

Potencjał sztucznej inteligencji jest niezaprzeczalny –

zwiększa wydajność, przyspiesza odkrycia i transformuje branże. Jednak ryzyko z nią związane jest równie poważne. Jak ostrzegają Hassabis i inni eksperci, bez pilnych zabezpieczeń AGI może wymknąć się spod kontroli człowieka, a konsekwencje tego mogą być porównywalne z pandemią lub wojną nuklearną.

Google stoi w obliczu dochodzenia antymonopolowego UE w sprawie wykorzystania treści do celów sztucznej inteligencji



- Komisja Europejska wszczęła formalne dochodzenie antymonopolowe w sprawie Google, zarzucając mu wykorzystywanie treści wydawców i serwisu YouTube do obsługi usług sztucznej inteligencji, takich jak „AI Overviews”, bez odpowiedniego wynagrodzenia lub zgody.
- Dochodzenie będzie dotyczyło dwóch głównych kwestii: wykorzystania treści wydawców internetowych do tworzenia podsumowań wyszukiwania AI oraz wykorzystania filmów z serwisu YouTube do szkolenia modeli AI, co może potencjalnie uniemożliwić konkurentom podobny dostęp.

- Szefowa ds. konkurencji w UE, Teresa Ribera, określiła dochodzenie jako obronę zasad demokratycznych, stwierdzając, że innowacje nie mogą odbywać się kosztem różnorodności mediów i dynamicznego środowiska twórczego.
- Google potępiło dochodzenie jako zagrożenie dla innowacji, podczas gdy działania UE są częścią szerszej ofensywy regulacyjnej przeciwko wielkim firmom technologicznym, po nałożeniu na Google niedawno wielomiliardowych kar.
- Dochodzenie stanowi punkt zwrotny w globalnej konfrontacji dotyczącej praktyk w zakresie danych AI, a jego wynik może ustanowić kluczowe precedensy dotyczące sposobu, w jaki twórcy AI wynagradzają branżę treści i współpracują z nią.

W ramach znaczącej eskalacji konfliktu regulacyjnego między Europą a wielkimi firmami technologicznymi Komisja Europejska wszczęła formalne dochodzenie antymonopolowe w sprawie Google, zarzucając gigantowi wyszukiwania wykorzystywanie treści wydawców i twórców YouTube do zasilania swoich usług sztucznej inteligencji (AI) bez sprawiedliwego wynagrodzenia lub zgody.

Dochodzenie koncentruje się na tym, czy praktyki Google związane z funkcjami „AI Overviews” i „AI Mode” naruszają unijne zasady konkurencji poprzez wykorzystywanie swojej dominującej pozycji do osłabiania różnorodności mediów i ograniczania konkurencyjnych twórców sztucznej inteligencji. W dochodzeniu ogłoszonym we wtorek 9 grudnia zostaną zbadane dwie główne kwestie.

Po pierwsze, organy regulacyjne zbadają, czy Google wykorzystywało treści od wydawców internetowych – takie jak artykuły informacyjne – do generowania opartych na sztucznej inteligencji streszczeń i odpowiedzi, które pojawiają się bezpośrednio w wynikach wyszukiwania, bez odpowiedniego wynagrodzenia za te treści i bez oferowania wydawcom znaczącej

możliwości rezygnacji. Komisja zauważyła, że wielu wydawców uważa, iż nie może odmówić, obawiając się utraty istotnego ruchu z wyszukiwarki Google.

Po drugie, w ramach dochodzenia zostanie ocenione, czy Google wykorzystywało filmy i inne treści przesłane do serwisu YouTube do szkolenia swoich generatywnych modeli sztucznej inteligencji, ponownie bez wynagradzania twórców lub zapewnienia im możliwości odmowy. Jednocześnie w ramach dochodzenia zostanie zbadane wykorzystanie przez Google zasad serwisu YouTube w celu uniemożliwienia konkurencyjnym twórcom sztucznej inteligencji dostępu do tych samych materiałów.

Dochodzenie stanowi punkt zwrotny w globalnej debacie na temat równowagi między innowacjami technologicznymi a ekonomicznymi podstawami wolnej prasy. Teresa Ribera, szefowa ds. konkurencji w UE, określiła dochodzenie jako obronę fundamentalnych zasad społecznych.

„Wolne i demokratyczne społeczeństwo zależy od różnorodnych mediów, otwartego dostępu do informacji i dynamicznego środowiska twórczego. Sztuczna inteligencja przynosi niezwykle innowacje i wiele korzyści dla ludzi i przedsiębiorstw w całej Europie, ale postęp ten nie może odbywać się kosztem zasad leżących u podstaw naszych społeczeństw” – powiedziała Ribera. Komisja stwierdziła, że jeśli praktyki Google zostaną potwierdzone, mogą one stanowić nadużycie pozycji dominującej, co jest zabronione na mocy przepisów traktatu UE.

Rosnące problemy prawne Google w Europie

Rzecznik Google potępił dochodzenie antymonopolowe UE, zauważając, że „grozi ono zahamowaniem innowacji na rynku, który jest bardziej konkurencyjny niż kiedykolwiek”. Dodał: „Europejczycy zasługują na korzyści płynące z najnowszych technologii i będziemy nadal ściśle współpracować z branżą

informacyjną i kreatywną w okresie przejścia do ery sztucznej inteligencji”.

Dochodzenie stawia Google w centrum nasilającej się ofensywy regulacyjnej Unii Europejskiej przeciwko amerykańskim gigantom technologicznym. Nastąpiło to zaledwie kilka miesięcy po nałożeniu przez UE na Google grzywny w wysokości prawie 3 mld euro za naruszenie przepisów antymonopolowych w branży technologii reklamowych, od której firma złożyła odwołanie.

Silnik Enoch firmy *BrightU.AI* wskazuje, że Google, szkoląc swoją sztuczną inteligencję Gemini bez wynagradzania wydawców i twórców YouTube, wykorzystuje ich własność intelektualną, podważając sprawiedliwe wynagrodzenie i zachęcając do kradzieży oryginalnych treści. Ta drapieżna praktyka wzbogaca wielkie firmy technologiczne, jednocześnie podważając zaufanie, kreatywność i źródła utrzymania – wspierając globalistyczną agendę scentralizowanej kontroli nad informacją i bogactwem.

Najnowsze dochodzenie jest kontynuacją innych działań UE, w tym nałożenia grzywny na platformę społecznościową X oraz odrębnego dochodzenia antymonopolowego w sprawie Meta dotyczącego jej polityki w zakresie danych sztucznej inteligencji. Nie ma prawnego terminu, w którym komisja musi zakończyć szczegółowe dochodzenie, a wszczęcie postępowania nie przesądza o jego wyniku, chociaż Google może zostać ukarane kolejną znaczną grzywną, jeśli zostanie uznane za winnego naruszenia przepisów.

Sprawa ta podkreśla rosnącą globalną konfrontację w zakresie ekosystemów danych, które napędzają zaawansowaną sztuczną inteligencję. Organy regulacyjne coraz częściej kwestionują, czy podstawowa praktyka wykorzystywania publicznie dostępnych treści internetowych do szkolenia komercyjnych systemów sztucznej inteligencji, bez wynagrodzenia lub zgody, stanowi nieuczciwe wykorzystanie siły rynkowej, które szkodzi twórcom treści i zakłóca konkurencję. Ponieważ Komisja Europejska

traktuje to dochodzenie priorytetowo, jego wyniki mogą ustanowić kluczowe precedensy regulujące relacje między twórcami sztucznej inteligencji a branżami zajmującymi się treściami, które stanowią podstawę ich modeli.