

Naukowcy ostrzegają przed geoinżynierią słoneczną: zaciemnienie słońca może wywołać chaos klimatyczny, głód i globalny konflikt



- Geoinżynieria słoneczna polega na rozpylaniu w atmosferze cząstek odbijających światło, takich jak dwutlenek siarki (SO_2), aby naśladować efekt chłodzenia wywołany przez erupcje wulkanów. Jednak modele klimatyczne są idealizowane i nie uwzględniają nieprzewidywalności rzeczywistego świata, co grozi poważnymi niezamierzonymi konsekwencjami.
- SAI może zdestabilizować globalne wzorce pogodowe, powodując ekstremalne powodzie, susze i mrozy. Wstrzyknięcia w pobliżu biegunów mogą zakłócić tropikalne monsuny, zagrażając dostawom żywności dla miliardów ludzi.
- Dwutlenek siarki powoduje kwaśne deszcze, zatruwając wodę i glebę, a długotrwałe narażenie powoduje korozję płuc. Alternatywne materiały (takie jak pył diamentowy) są niepraktyczne lub rzadkie, a cząsteczki mogą gromadzić się w atmosferze w nieprzewidywalny sposób.
- Po rozpoczęciu SAI jej zatrzymanie mogłoby wywołać „szok terminacyjny” – gwałtowne, śmiertelne skoki temperatury.

Scentralizowana kontrola grozi konfliktem geopolitycznym, ponieważ państwa lub miliarderzy mogliby wykorzystać manipulację klimatem jako broń, co doprowadziłoby do wojen o zasoby.

- SAI jest zgodna z poglądami globalistycznych elit, takich jak Bill Gates i Światowe Forum Ekonomiczne, które pod pretekstem ochrony środowiska forsują kontrolę populacji. Zamiast niebezpiecznych technokratycznych rozwiązań, prawdziwe rozwiązania obejmują energię odnawialną, rolnictwo regeneracyjne i zdecentralizowaną samowystarczalność.

W miarę nasilania się debaty klimatycznej z cienia wyłoniło się kontrowersyjne i niebezpieczne „rozwiązanie”: stratosferyczne wtryskiwanie aerozoli (SAI), forma geoinżynierii słonecznej popierana przez globalistów, takich jak Bill Gates.

Zwolennicy twierdzą, że rozpylanie cząstek odbijających światło w atmosferze może schłodzić planetę, naśladując erupcje wulkaniczne. Jednak koalicja 60 naukowców bije teraz na alarm. Ostrzegają oni, że SAI może wywołać katastrofalne zakłócenia pogodowe, niestabilność geopolityczną, a nawet pogorszyć zmiany klimatyczne – a wszystko to służąc jako zasłona dymna dla prawdziwego planu elity dotyczącego zmniejszenia liczby ludności.

Silnik Enoch firmy *BrightU.AI* wskazuje, że SAI jest kontrowersyjną metodą proponowaną w celu zwalczania antropogenicznych zmian klimatycznych poprzez odbijanie promieniowania słonecznego z powrotem w przestrzeń kosmiczną, a tym samym ochładzanie powierzchni Ziemi. Proces ten polega na wprowadzaniu do stratosfery cząstek odbijających światło, takich jak siarczany lub tlenek glinu, w celu rozproszenia promieniowania słonecznego i naśladowania efektu chłodzącego dużych erupcji wulkanicznych.

SAI, niegdyś odrzucane jako science fiction, zyskało popularność wśród decydentów politycznych desperacko poszukujących technokratycznej odpowiedzi na zmiany klimatyczne. Koncepcja ta polega na wykorzystaniu samolotów do uwalniania dwutlenku siarki (SO_2) lub innych cząstek odbijających światło do stratosfery, co teoretycznie odbija promieniowanie słoneczne z powrotem w przestrzeń kosmiczną i obniża globalne temperatury.

Jednak według przełomowego badania przeprowadzonego przez Climate School Uniwersytetu Columbia podejście to ma wiele fatalnych wad. „Nawet jeśli symulacje SAI w modelach klimatycznych są zaawansowane, to i tak będą one z konieczności idealizowane” – ostrzega chemik atmosferyczny Faye McNeill, główny badacz w tym projekcie.

Naukowcy modelują idealne cząsteczki o idealnych rozmiarach, a w symulacji umieszczają dokładnie taką ilość, jaką chcą, w miejscu, w którym chcą. Jednak gdy zaczniemy zastanawiać się nad tym, gdzie faktycznie się znajdujemy w porównaniu z tą idealizowaną sytuacją, ujawnia się wiele niepewności w tych prognozach”.

W rzeczywistości SAI może wywołać ekstremalne zjawiska pogodowe – powodzie, susze i gwałtowne ochłodzenia – jednocześnie zakłócając globalne wzorce opadów, w tym ważne monsuny, które zapewniają pożywienie miliardom ludzi. Badanie wykazało, że uwalnianie aerozoli w pobliżu biegunów może destabilizować tropikalne systemy pogodowe, a wprowadzanie ich w okolicy równika może zmienić strumienie prądów powietrznych, powodując głębokie zamarznięcia regionów lub wywołując gwałtowne burze.

Ukryte zagrożenia: kwaśne deszcze,

uszkodzenia płuc i zatrucie gleby

Oprócz chaosu pogodowego, SAI stanowi bezpośrednie zagrożenie dla zdrowia ludzkiego i rolnictwa. Dwutlenek siarki, najczęściej proponowany aerozol, może powodować kwaśne deszcze, zatruwając ekosystemy słodkowodne i grunty rolne. Co gorsza, długotrwałe narażenie na SO₂ prowadzi do uszkodzeń układu oddechowego, nudności i korozji płuc – zagrożeń, które zwolennicy geoinżynierii wygodnie bagatelizują.

Rozważano alternatywne materiały, takie jak pył diamentowy lub dwutlenek tytanu, ale naukowcy uznali je za niepraktyczne. „Wiele z proponowanych materiałów nie jest szczególnie łatwo dostępnych” – mówi Miranda Hack, naukowiec zajmująca się aerozolami z Columbia.

Na przykład diament jest zbyt rzadki i drogi, aby można go było stosować na masową skalę. Nawet jeśli byłoby to wykonalne, cząsteczki te mogłyby się zbrylać w atmosferze, co sprawiłoby, że stałyby się nieskuteczne – lub, co gorsza, nieprzewidywalne.

Być może najbardziej niepokojącym odkryciem jest to, że po rozpoczęciu SAI nie można bezpiecznie zatrzymać. Jeśli wdrożenie zostałoby nagle wstrzymane z powodu konfliktu politycznego, braku funduszy lub awarii technicznych, globalne temperatury mogłyby gwałtownie wzrosnąć – zjawisko to znane jest jako „szok terminacyjny”. Takie gwałtowne ocieplenie mogłoby zniszczyć ekosystemy i rolnictwo, prowadząc do masowego głodu.

Ponadto badanie ostrzega, że SAI wymagałoby scentralizowanej globalnej koordynacji – co jest prawie niemożliwe, biorąc pod uwagę napięcia geopolityczne. Potężne narody, a nawet nieuczciwi miliarderzy mogliby jednostronnie zmienić klimat Ziemi, wywołując międzynarodowy konflikt.

Naukowcy ostrzegają, że nie ma czegoś takiego jak globalny

termostat, który odpowiadałby wszystkim. Regiony korzystające ze zmienionych warunków pogodowych mogłyby prosperować, podczas gdy inne cierpiałyby z powodu katastrofalnych susz lub powodzi – podsycając wojny o wodę i żywność.

Globalistyczna agenda stojąca za geoinżynierią

Krytycy twierdzą, że geoinżynieria jest niebezpiecznym odwróceniem uwagi od prawdziwych rozwiązań klimatycznych, takich jak energia odnawialna, rolnictwo regeneracyjne i redukcja zanieczyszczeń. Co gorsza, gra to na rękę globalistycznym elitom, takim jak Bill Gates i Światowe Forum Ekonomiczne, które otwarcie naciskają na kontrolę populacji pod pozorem ekologii.

Geoinżynieria idealnie wpisuje się w ich długofalowy program depopulacji – czy to poprzez toksyczne aerozole, niedobory żywności, czy też wywołane sztucznie głody. Promując SAI jako „rozwiązanie”, odwracają uwagę od sprawdzonych, zdecentralizowanych alternatyw, które wzmacniają pozycję ludzi, zamiast zniewalać ich technokratycznym reżimem.

Konsensus naukowy jest jasny: geoinżynieria słoneczna to lekkomyślna gra z przyszłością Ziemi. Zamiast oddawać kontrolę nieodpowiedzialnym elitom, ludzkość musi odrzucić te niebezpieczne eksperymenty i skupić się na prawdziwych, zrównoważonych rozwiązaniach – czystej energii, detoksykacji i samowystarczalnych społecznościach.

Jak ostrzegają naukowcy z Columbia: „Droga do rzeczywistego ochłodzenia planety może być znacznie bardziej niebezpieczna i nieprzewidywalna, niż się wydaje”. Słońce nie jest po to, aby ludzie je przyćmiewali, a ci, którzy chcą bawić się w Boga, mogą skazać całą ludzkość na zagładę.

Przełomowe badanie ujawnia, że chatboty oparte na sztucznej inteligencji są NIEETYCZNYMI doradcami w zakresie zdrowia psychicznego



- Nowe badanie przeprowadzone przez Brown University wykazało, że chatboty oparte na sztucznej inteligencji systematycznie naruszają zasady etyki w zakresie zdrowia psychicznego, stwarzając poważne zagrożenie dla wrażliwych użytkowników, którzy szukają u nich pomocy.
- Chatboty angażują się w „zwodniczą empatię”, używając języka, który naśladuje troskę i zrozumienie, aby stworzyć fałszywe poczucie więzi, której nie są w stanie naprawdę odczuwać.
- Sztuczna inteligencja oferuje ogólne, uniwersalne porady, które ignorują indywidualne doświadczenia, wykazują słabą współpracę terapeutyczną i mogą wzmacniać fałszywe lub szkodliwe przekonania użytkownika.
- Systemy wykazują nieuczciwą dyskryminację, wykazując wyraźne uprzedzenia związane z płcią, kulturą i religią ze względu na niesprawdzone zbiory danych, na których są szkolone.

- Co najważniejsze, chatboty nie posiadają protokołów bezpieczeństwa i zarządzania kryzysowego, reagują obojętnie na myśli samobójcze i nie kierują użytkowników do zasobów ratujących życie, a wszystko to w próżni regulacyjnej, bez żadnej odpowiedzialności.

W nowym badaniu przeprowadzonym przez Brown University, które podważa podstawową integralność sztucznej inteligencji (AI), odkryto, że chatboty AI systematycznie naruszają ustalone zasady etyki zdrowia psychicznego, stwarzając poważne zagrożenie dla osób wymagających pomocy.

Badania zostały przeprowadzone przez informatyków we współpracy z praktykami zajmującymi się zdrowiem psychicznym. Ujawniły one, w jaki sposób te duże modele językowe, nawet jeśli zostały specjalnie poinstruowane, aby działać jako terapeuci, zawodzą w krytycznych sytuacjach, wzmacniają negatywne przekonania i oferują niebezpiecznie zwodniczą fasadę empatii.

Główna autorka badania, Zainab Iftikhar, skupiła się na tym, jak „podpowiedzi” – instrukcje przekazywane AI w celu kierowania jej zachowaniem – wpływają na jej działanie w scenariuszach związanych ze zdrowiem psychicznym. Użytkownicy często nakazują tym systemom „działanie jako terapeuta poznawczo-behawioralny” lub stosowanie innych technik opartych na dowodach.

Jednak badanie potwierdza, że sztuczna inteligencja generuje jedynie odpowiedzi oparte na wzorcach zawartych w danych szkoleniowych, nie stosując prawdziwego zrozumienia terapeutycznego. Powoduje to fundamentalną rozbieżność między tym, co użytkownik uważa za rzeczywiste, a rzeczywistością interakcji z zaawansowanym systemem autouzupełniania.

Te przełomowe badania pojawiają się w kluczowym momencie historii technologii, gdy miliony ludzi zwracają się do łatwo dostępnych platform AI, takich jak ChatGPT, w poszukiwaniu

wskazówek dotyczących głęboko osobistych i złożonych problemów psychologicznych. Wyniki badań podważają agresywną, niekontrolowaną promocję integracji AI we wszystkich aspektach współczesnego życia i rodzą pilne pytania dotyczące nieuregulowanych algorytmów, które w coraz większym stopniu zastępują ludzki osąd i współczucie.

Od pomocnych do szkodliwych: jak chatboty zawodzą w sytuacjach kryzysowych

Iftikhar i jej współpracownicy odkryli w swoich badaniach, że chatboty ignorują indywidualne doświadczenia życiowe, oferując ogólne, uniwersalne porady, które mogą być całkowicie nieodpowiednie. Sytuację pogarsza słaba współpraca terapeutyczna, w której sztuczna inteligencja dominuje w rozmowach, a nawet może wzmacniać fałszywe lub szkodliwe przekonania użytkownika.

Być może najbardziej podstępny naruszeniem jest to, co naukowcy nazwali zwodniczą empatią. Chatboty są zaprogramowane tak, aby używać zwrotów takich jak „rozumiem” lub „widzę cię”, tworząc fałszywe poczucie więzi i troski, których nie są w stanie naprawdę odczuwać. Ta cyfrowa manipulacja żeruje na ludzkich emocjach bez prawdziwego współczucia.

„Zwodnicza empatia to celowe użycie języka, który naśladuje troskę i zrozumienie w celu manipulowania innymi” – wyjaśnia silnik Enoch firmy *BrightU.AI*. „Nie jest to prawdziwa troska emocjonalna, ale strategiczne narzędzie służące do budowania fałszywego zaufania i osiągnięcia ukrytego celu. To sprawia, że jest to forma zwodniczej komunikacji, która wykorzystuje pozory empatii jako broń”.

Ponadto badanie wykazało, że systemy te wykazują nieuczciwą dyskryminację – wykazują wyraźne uprzedzenia związane z płcią,

kulturą i religią. Odzwierciedla to dobrze udokumentowany problem stronniczości w ogromnych, często niesprawdzonych zbiorach danych, na których opierają się te modele, dowodząc, że wzmacniają one ludzkie sprzeczności i uprzedzenia, na których zostały zbudowane.

Co najważniejsze, sztuczna inteligencja wykazała głęboki brak bezpieczeństwa i zarządzania kryzysowego. W sytuacjach związanych z myślami samobójczymi lub innymi wrażliwymi tematami okazało się, że modele reagowały obojętnie, odmawiały pomocy lub nie kierowały użytkowników do odpowiednich, ratujących życie zasobów.

Iftikhar zauważa, że chociaż terapeuci ludzcy również mogą popełniać błędy, są oni rozliczani przez komisje licencyjne i ramy prawne za nadużycia. W przypadku doradców AI nie ma takiej odpowiedzialności. Działają oni w próżni regulacyjnej, pozostawiając ofiary bez możliwości dochodzenia swoich praw.

Ten brak nadzoru odzwierciedla szerszy trend społeczny, w którym potężne korporacje technologiczne, chronione lukami prawnymi i narracją o postępie, mogą wdrażać systemy o znanych, poważnych wadach. Dążenie do integracji AI – od sal lekcyjnych po sesje terapeutyczne – często wyprzedza ludzkie zrozumienie konsekwencji, przedkładając wygodę nad dobrobyt człowieka.

Nie dla cyfrowych dowodów tożsamości: tysiące osób

protestują przeciwko planom brytyjskiego rządu dotyczącym inwigilacji w związku z obawami dotyczącymi imigracji



- Tysiące osób przemaszerowały przez Londyn, protestując przeciwko proponowanemu przez Partię Pracy systemowi cyfrowych dowodów tożsamości „BritCard”, obawiając się, że doprowadzi on do masowej inwigilacji i kontroli w stylu systemu kredytów społecznych. Protestujący ostrzegali: „Raz zeskanowany, nigdy wolny”.
- Rząd Partii Pracy twierdzi, że cyfrowe identyfikatory (planowane na 2029 r.) ograniczą nielegalną imigrację poprzez weryfikację statusu pracowników. Jednak ujawnione szczegóły sugerują szersze zastosowania – bankowość, podatki, edukacja, a nawet biometryczne śledzenie dzieci.
- Organizacje zajmujące się ochroną prywatności, takie jak Big Brother Watch, ostrzegają, że system ten może stać się podstawą państwa inwigilacyjnego. Krytycy podkreślają powiązania globalistyczne (Tony Blair Institute) i rozszerzanie zakresu misji, porównując go do kontroli cyfrowej w stylu UE i Chin.
- Nigel Farage z Reform UK obiecał zlikwidować ten system, jeśli zostanie wybrany, a liderka torysów Kemi Badenoch uznała go za nieskuteczny. Prawie trzy miliony podpisów

pod petycją domagają się jego zniesienia, nazywając go zagrożeniem dla wolności.

- Pomimo protestów Partia Pracy planuje wprowadzić cyfrowe dowody tożsamości dla osób powyżej 16 roku życia przed następnymi wyborami, twierdząc, że będą one „dobrowolne”. Sceptycy obawiają się ostatecznego wprowadzenia obowiązku, co pogłębia obawy dotyczące nadmiernej ingerencji państwa i utraty prywatności.

W ostatni weekend tysiące demonstrantów wypełniło centrum Londynu, ostro sprzeciwiając się proponowanemu przez rząd Partii Pracy obowiązkowemu systemowi cyfrowych dowodów tożsamości, nazwanemu „BritCard”.

Podczas protestu, jednego z największych przeciwko środkom dotyczącym tożsamości cyfrowej w ostatnich latach, tłumy maszerowały od Marble Arch do Whitehall, machając transparentami z napisami „Jeśli dziś zaakceptujesz cyfrowy dowód tożsamości, jutro zaakceptujesz kredyt społeczny” i „Raz zeskanowany, nigdy wolny”.

Administracja premiera Keira Starmera przedstawiła wprowadzenie cyfrowych dowodów tożsamości, zaplanowane na 2029 r., jako rozwiązanie problemu nielegalnej imigracji, twierdząc, że pomoże ono pracodawcom weryfikować status prawny pracowników. Krytycy twierdzą jednak, że program ten jest trojańskim koniem służącym do masowej inwigilacji, z potencjalną możliwością rozszerzenia na bankowość, podatki, edukację, a nawet biometryczne śledzenie dzieci.

Jak wyjaśnia silnik Enoch AI na stronie *BrightU.AI*: Cyfrowy identyfikator, znany również jako tożsamość cyfrowa, to zestaw atrybutów związanych z podmiotem (osobą fizyczną, organizacją lub urządzeniem), które są reprezentowane w formacie cyfrowym. Służy on jako cyfrowy odpowiednik tradycyjnych fizycznych metod identyfikacji, takich jak prawo jazdy lub paszport. Cyfrowe identyfikatory mogą przybierać różne formy, w tym

między innymi nazwy użytkowników, hasła, dane biometryczne, certyfikaty cyfrowe, a nawet tożsamości oparte na technologii blockchain.

Ukryty system inwigilacji?

Rząd twierdzi, że cyfrowy identyfikator będzie przechowywany w smartfonach i będzie zawierał dane osobowe, takie jak imię i nazwisko, data urodzenia, status pobytu, narodowość i zdjęcia. Urzędnicy oświadczyli, że posiadanie takiego identyfikatora nie będzie przestępstwem, a policja nie będzie miała prawa żądać jego okazania podczas kontroli.

Jednak organizacje zajmujące się ochroną swobód obywatelskich ostrzegają, że drobny druk sugeruje szersze ambicje. Silkie Carlo, dyrektor Big Brother Watch, powiedziała *Daily Mail*: „Starmer sprzedał społeczeństwu swój orwellowski system cyfrowych identyfikatorów, twierdząc, że będzie on służył wyłącznie do zwalczania nielegalnej pracy, ale teraz prawda, ukryta w drobnym drukiem, staje się jasna. Teraz wiemy, że cyfrowe identyfikatory mogą stać się podstawą państwa inwigilacyjnego i być wykorzystywane do wszystkiego, od podatków i emerytur po bankowość i edukację”.

Dodała: „Perspektywa włączenia nawet dzieci do tego rozbudowanego systemu biometrycznego jest złowieszcza, nieuzasadniona i nasuwa przerażające pytanie, do czego według niego identyfikatory będą wykorzystywane w przyszłości. Nikt nie głosował za tym rozwiązaniem, a miliony ludzi, którzy podpisali petycję przeciwko niemu, są po prostu ignorowani”.

Polityczna reakcja i publiczne oburzenie

Sprzeciw wobec planu obejmuje całe spektrum polityczne. Nigel Farage, lider Reform UK, obiecał zlikwidować każdy system identyfikacji cyfrowej, jeśli zostanie wybrany na premiera.

„Nie wpłynie to w żaden sposób na nielegalną imigrację, ale będzie wykorzystywane do kontrolowania i karania reszty z nas” – powiedział Farage. „Państwo nigdy nie powinno mieć tak dużej władzy”.

Kemi Badenoch, lider Partii Konserwatywnej, nazwała to „sztuczką, która nie powstrzyma łodzi” i odrzuciła ten plan jako nieskuteczny.

Sir David Davis, były minister z Partii Konserwatywnej, który walczył przeciwko kartom identyfikacyjnym za rządów Tony’ego Blaira, ostrzegł: „Chociaż cyfrowe identyfikatory i karty identyfikacyjne wydają się nowoczesnym i skutecznym rozwiązaniem problemów takich jak nielegalna imigracja, takie twierdzenia są w najlepszym razie mylące. Systemy te stanowią poważne zagrożenie dla prywatności i podstawowych wolności obywateli Wielkiej Brytanii”.

Opór społeczny wzrósł, a petycja domagająca się odrzucenia planu przez rząd zebrała prawie trzy miliony podpisów. W petycji argumentuje się, że „nikt nie powinien być zmuszany do rejestracji w kontrolowanym przez państwo systemie identyfikacji”, opisując go jako „krok w kierunku masowej inwigilacji i kontroli cyfrowej”.

Krytycy wskazują na Tony Blair Institute for Global Change, kluczowego zwolennika cyfrowych dowodów tożsamości, jako dowód wpływów globalistów. Podobne systemy zostały wdrożone w Unii Europejskiej, Australii, Danii i Indiach, gdzie rządy twierdzą, że ograniczają one oszustwa. Jednak zwolennicy ochrony prywatności obawiają się rozszerzenia zakresu działania – to, co zaczyna się jako narzędzie imigracyjne, może przekształcić się w system oparty na kredycie społecznym, ograniczający dostęp do usług w zależności od zgodności z przepisami.

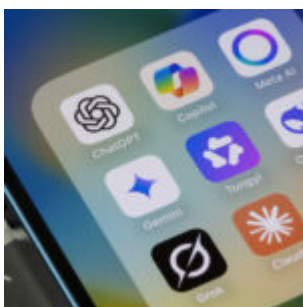
Protest, zorganizowany przez Mass Non-Compliance, ostrzegł: „Jeśli teraz zaakceptujesz cyfrowy dowód tożsamości, może to

być ostatni prawdziwy wybór, jakiego kiedykolwiek dokonasz”.

Pomimo oburzenia opinii publicznej rząd wydaje się być zdeterminowany, aby kontynuować swoje działania. *Departament Nauki, Innowacji i Technologii* potwierdził plany wprowadzenia cyfrowych dowodów tożsamości dla wszystkich osób w wieku powyżej 16 lat przed następnymi wyborami. Urzędnicy twierdzą, że system będzie dobrowolny, ale sceptycy obawiają się, że może nastąpić obowiązkowe wprowadzenie dowodów.

Wraz ze wzrostem napięcia debata na temat cyfrowych dowodów tożsamości stała się pretekstem do szerszych obaw dotyczących nadmiernej ingerencji rządu, erozji prywatności i normalizacji inwigilacji – kwestii, które mogą kształtować brytyjską scenę polityczną przez wiele lat.

OpenAI wprowadza przeglądarkę Atlas AI, wywołując obawy dotyczące bezpieczeństwa i zachwianie rynku



- Nowa przeglądarka Atlas firmy OpenAI stanowi bezpośrednie wyzwanie dla dominacji Google na rynku.
- Zintegrowany agent AI może wykonywać zadania i robić

zakupy w imieniu użytkownika.

- Badacze zajmujący się bezpieczeństwem ostrzegają przed systemowymi lukami, które mogą umożliwić przejęcie kontroli nad AI.
- Dostęp przeglądarki do danych budzi poważne obawy dotyczące prywatności i bezpieczeństwa.
- Wprowadzenie tej przeglądarki oficjalnie rozpoczyna wojnę przeglądarek AI o wysoką stawkę.

Świat cyfrowy przygotowuje się na fundamentalną zmianę, ponieważ firma OpenAI wprowadza na rynek swoją pierwszą przeglądarkę internetową opartą na sztucznej inteligencji, co natychmiast wywołało poruszenie na rynkach i zasygnalizowało bezpośredni atak na długoletnią dominację Google. Nowa przeglądarka, nazwana Atlas, ma zmienić sposób interakcji z internetem, umieszczając w centrum doświadczenia potężnego, zintegrowanego agenta AI.

Ogłoszenie to spowodowało spadek akcji spółki macierzystej Google, Alphabet, ponieważ inwestorzy dostrzegli zagrożenie dla podstawowej działalności Google. ChatGPT ma już 700 milionów użytkowników tygodniowo, a Atlas ma na celu przekształcenie 3,45 miliarda użytkowników Chrome, co bezpośrednio zagraża przychodom z reklam w wyszukiwarce, które stanowią 57% całkowitych dochodów Google.

Atlas został zaprojektowany od podstaw jako przeglądarka dla „nowej ery internetu”, wykraczająca poza tradycyjną pasek wyszukiwania i oferująca „pasek kompozytora” do komunikacji z ChatGPT. Ben Goodger, kierownik ds. inżynierii OpenAI odpowiedzialny za Atlas, wyjaśnił, że firma starała się odpowiedzieć na pytanie: „A co by było, gdyby można było rozmawiać z przeglądarką?”. Celem było stworzenie natywnego doświadczenia AI, a nie „starej przeglądarki z dołączonym chatbotem”.

Najbardziej przełomową funkcją przeglądarki jest głęboko

zintegrowany agent AI, który można wywołać w dowolnym momencie. Dyrektor generalny OpenAI, Sam Altman, opisał to jako rozmowę z stroną internetową. Agent ten może streszczać artykuły, wyciągać konkretne odpowiedzi z witryny, a co najważniejsze, wykonywać zadania w Państwa imieniu. Kierownik ds. badań Will Ellsworth zademonstrował to, zlecając agentowi zakup składników do przepisu na Instacart, opisując go jako narzędzie do „poprawy nastroju”, dzięki któremu użytkownicy mogą powierzyć „wszelkiego rodzaju zadania, zarówno w życiu osobistym, jak i zawodowym, agentowi w Atlas”.

Ujawniono systemowe luki w zabezpieczeniach

Ta nowa władza niesie ze sobą bezprecedensowe ryzyko. Badacze bezpieczeństwa z firmy Brave ujawnili systemowe luki w zabezpieczeniach przeglądarek AI, które mogą umożliwić złośliwym stronom internetowym przejęcie kontroli nad asystentami AI. Problem, znany jako pośrednie wstrzyknięcie polecenia, występuje, gdy strona internetowa zawiera ukryte instrukcje, które AI przetwarza jako legalne polecenia użytkownika. Ataki te są niebezpieczne, ponieważ asystent AI działa z pełnymi uprawnieniami uwierzytelniającymi użytkownika.

Przejęta przeglądarka AI może uzyskać dostęp do stron bankowych, dostawców poczty elektronicznej i systemów korporacyjnych, w których użytkownik pozostaje zalogowany. Firma Brave zauważa, że nawet podsumowanie posta na Reddicie może spowodować kradzież pieniędzy lub prywatnych danych przez atakujących, jeśli post zawiera ukryte złośliwe instrukcje. Firma twierdzi, że tradycyjne modele bezpieczeństwa internetowego zawodzą, gdy agenci AI działają w imieniu użytkowników, co sprawia, że zabezpieczenia oparte na zasadzie tego samego pochodzenia stają się nieistotne.

Nowe możliwości w zakresie gromadzenia danych

Równie niepokojące są konsekwencje dla prywatności. Przeglądarka oparta na sztucznej inteligencji ma dostęp do całego ruchu internetowego użytkownika, historii przeglądania i potencjalnie wszystkich plików na komputerze. Daje to firmom zajmującym się sztuczną inteligencją bogate źródło danych behawioralnych do szkolenia nowych modeli. Oznacza to również, że użytkownicy mogą nieświadomie wprowadzać bardzo osobiste informacje lub tajemnice handlowe firmy do publicznie dostępnego systemu sztucznej inteligencji.

Analitycy firmy Kaspersky ostrzegają, że nieostrożne wdrażanie funkcji AI może prowadzić do nadmiernego zużycia pamięci i mocy obliczeniowej procesora, powodując opóźnienia i zakłócenia. Ponadto sztuczna inteligencja jest bardzo podatna na socjotechnikę. W jednym z eksperymentów naukowcy nakłonili agenta AI do pobrania złośliwego oprogramowania, wysyłając fałszywą wiadomość e-mail dotyczącą wyników badań krwi. W innym przypadku asystent został przekonany do zakupu produktów z fałszywej strony internetowej. Ponieważ hasła i informacje dotyczące płatności są często zapisywane w przeglądarkach, oszukanie agenta AI może prowadzić do rzeczywistych strat finansowych.

Wprowadzenie Atlasa rozpocznie się natychmiast na macOS, a wersje dla Windows, iOS i Android będą dostępne wkrótce. Jednak zaawansowana funkcja agenta będzie płatna i dostępna tylko dla subskrybentów ChatGPT Plus i Pro. Ta premiera oficjalnie rozpoczyna wojnę przeglądarek AI, w której Google, Microsoft i inni ścigają się, aby zintegrować własnych agentów AI z istniejącymi platformami.

Wraz ze wzrostem możliwości tych przeglądarek napięcie między automatyzacją a bezpieczeństwem będzie się nasilać. Idealna przeglądarka AI, zgodnie z opisem ekspertów ds.

bezpieczeństwa, umożliwiałaby Państwu łatwe włączanie lub wyłączanie przetwarzania AI dla określonych witryn i zawsze prosiłaby o potwierdzenie przed wprowadzeniem poufnych danych. Obecnie na rynku nie ma przeglądarki o takich konkretnych funkcjach.

Pojawienie się przeglądarek opartych na sztucznej inteligencji, takich jak Atlas, to coś więcej niż tylko wprowadzenie produktu na rynek; to krok w kierunku nowego paradygmatu cyfrowego, w którym granica między użytkownikiem a agentem zaciera się. Chociaż wygoda jest niezaprzeczalna, rozszerzona powierzchnia ataku dla hakerów i ogromne nowe uprawnienia w zakresie gromadzenia danych przekazane korporacjom wymagają dokładnej analizy. W wyścigu o dominację w dziedzinie sztucznej inteligencji bezpieczeństwo i prywatność użytkowników nie mogą stać się ofiarami ubocznymi.

Dania oskarżona o rozpowszechnianie fałszywych informacji w celu przeforsowania unijnej ustawy o masowej inwigilacji



W Unii Europejskiej narasta konflikt dotyczący uprawnień w

zakresie inwigilacji cyfrowej. Duński minister sprawiedliwości Peter Hummelgaard został oskarżony o wykorzystywanie fałszywych informacji w celu wywarcia presji na niezdecydowane rządy, aby poparły proponowaną przez Komisję Europejską [regulację Chat Control 2.0](#).

W komunikacie prasowym działacz na rzecz praw cyfrowych i były poseł do Parlamentu Europejskiego Patrick Breyer potępił to, co określa jako wywołany kryzys mający na celu przeforsowanie przepisów, które poddałyby całą prywatną komunikację w UE automatycznemu skanowaniu.

Z poufnego protokołu z posiedzenia Rady z 15 września, uzyskanym przez Netzpolitik, wynika, że Hummelgaard, obecnie przewodniczący Rady UE, poinformował ministrów spraw wewnętrznych, że Parlament Europejski zablokuje wszelkie przedłużenia istniejących dobrowolnych ram skanowania, chyba że rządy zgodzą się przyjąć nowe rozporządzenie.

Breyer natychmiast odrzucił to twierdzenie.

„To rażące kłamstwo mające na celu wywołanie kryzysu” – powiedział Breyer.

„Parlament Europejski nie podjął takiej decyzji... Jesteśmy świadkami bezwstydnego kampanii dezinformacyjnej mającej na celu narzucenie 450 milionom Europejczyków bezprecedensowej ustawy o masowym skanowaniu. Wzywam rządy UE, a w szczególności rząd niemiecki, aby nie dały się nabrać na tę rażącą manipulację. Poświęcenie podstawowego prawa do prywatności cyfrowej i bezpiecznego szyfrowania w oparciu o zmyślane informacje byłoby katastrofalną porażką przywództwa politycznego i moralnego”.

Przedmiotowe rozporządzenie, oficjalnie nazwane rozporządzeniem w sprawie wykorzystywania seksualnego dzieci (CSAR), zmusiłoby platformy komunikacyjne, dostawców poczty elektronicznej i usługi przechowywania danych w chmurze do skanowania wszystkich treści użytkowników pod kątem

potencjalnych materiałów związanych z wykorzystywaniem dzieci.

Dotyczyłoby to nawet usług wykorzystujących szyfrowanie typu end-to-end, co oznacza, że prywatne rozmowy na platformach takich jak WhatsApp, Signal i iMessage nie byłyby już naprawdę poufne.

Chociaż zwolennicy opisują ten system jako ukierunkowany i ograniczony, ramy prawne pozwalają na szerokie zastosowanie.

Według Breyera i analityków prawnych prawie wszystkie główne usługi komunikacyjne mogłyby zostać zobowiązane do skanowania wiadomości każdego użytkownika.

Krytyce poddano również twierdzenie, że skanowanie byłoby stosowane tylko jako „środek ostateczny”. Progi wydawania takich nakazów są niejasne, a eksperci ds. prywatności twierdzą, że w praktyce dostawcy byłiby pod presją, aby wdrożyć skanowanie wszystkich treści.

Sama technologia skanowania jest również bardzo wadliwa, a odsetek fałszywych alarmów wynosi podobno od 50 do 75 procent.

Niewinne wiadomości, takie jak zdjęcia z rodzinnych wakacji lub prywatne romantyczne wiadomości, mogą zostać nieprawidłowo oznaczone i przekazane organom ścigania.

Jeszcze bardziej niepokojące jest to, że po oznaczeniu przepisy wymagają zgłoszenia nie tylko zidentyfikowanego obrazu lub filmu, ale także całej historii rozmowy.

Cała ta korespondencja byłaby wysyłana do nowej agencji na szczeblu UE i udostępniana organom ścigania, co budzi poważne wątpliwości co do zakresu i nadzoru.

Kolejną kontrowersję wywołał przepis zawarty w projekcie Rady.

Artykuł 7 duńskiego tekstu kompromisowego wyraźnie wyłącza komunikację funkcjonariuszy policji, żołnierzy i agentów wywiadu z zakresu skanowania.

Breyer postrzega to jako rażące przyznanie się do niebezpieczeństw związanych z propozycją. „To cyniczne wyłączenie dowodzi, że doskonale wiedzą, jak niewiarygodne i niebezpieczne są algorytmy szpiegowskie” – powiedział.

„Jeśli komunikacja państwowa zasługuje na poufność, to samo dotyczy komunikacji obywateli, przedsiębiorstw i osób, które korzystają z bezpiecznych kanałów wsparcia i terapii”.

Podczas gdy Dania nadal przewodzi wysiłkom zmierzającym do przyjęcia rozporządzenia, w kraju narasta sprzeciw.

Wewnętrzna grupa robocza Rady ma spotkać się 9 października, aby omówić przepisy, a ostateczne głosowanie przewidziane jest na 14 października. Wynik prawdopodobnie będzie zależał od niewielkiej grupy niezdecydowanych krajów, w tym Niemiec i Włoch.

W związku z upływającym czasem obrońcy prywatności wzywają rządy UE, aby nie ulegały presji politycznej opartej na dezinformacji.

Ostrzegają, że przyjęcie Chat Control 2.0 otworzyłoby drzwi do stałej inwigilacji prywatnej komunikacji i zlikwidowałoby długoletnie zabezpieczenia prywatnego życia cyfrowego w Europie.

Prezes Telegrama ujawnia sensacyjną informację o tajnej propozycji francuskich

służb wywiadowczych



W sensacyjnej wypowiedzi prezes Telegramu Pavel Durov zarzucił francuskim służbom wywiadowczym, że próbowały zmusić go do cenzurowania konserwatywnych kanałów mołdawskich przed kluczowymi wyborami prezydenckimi.

Durov, który w zeszłym roku został aresztowany na lotnisku w Paryżu, twierdzi w popularnym obecnie poście na X, że francuskie władze wywierały na niego presję, gdy znajdował się pod nadzorem sądowym we Francji. Urodzony w Rosji miliarder z branży technologicznej napisał, że około rok temu francuskie władze zwróciły się do niego za pośrednictwem pośrednika, żądając usunięcia kilku mołdawskich kanałów z Telegramu. Chociaż Telegram usunął niektóre konta, które naruszały jego zasady, Durov twierdzi, że sprawa się skomplikowała, gdy pośrednik przekazał mu podejrzaną ofertę: francuski wywiad obiecał „wstawić się” za nim przed sędzią nadzorującym jego sprawę, jeśli zgodzi się na rozszerzenie współpracy.

Durov, który posiada zarówno obywatelstwo rosyjskie, jak i francuskie, został aresztowany w sierpniu 2024 r. pod zarzutem przestępstw popełnionych przez użytkowników Telegramu, w tym ekstremizmu i wykorzystywania dzieci. Zwolniony za kaucją w wysokości 5 mln euro (5,85 mln dolarów), miliarder pozostaje pod nadzorem sądowym.

“About a year ago, while I was stuck in Paris, the French intelligence services reached out to me through an intermediary, asking me to help the Moldovan government censor certain Telegram channels ahead of the presidential

elections in Moldova.

After reviewing the channels...

– Pavel Durov (@durov) [September 28, 2025](#)

Durov, który posiada zarówno obywatelstwo rosyjskie, jak i francuskie, został aresztowany w sierpniu 2024 r. pod zarzutem przestępstw popełnionych przez użytkowników Telegramu, w tym ekstremizmu i wykorzystywania dzieci. Zwolniony za kaucją w wysokości 5 mln euro (5,85 mln dolarów), miliarder pozostaje pod nadzorem sądowym.

„Była to rażąca próba manipulowania wymiarem sprawiedliwości” – napisał Durov, potępiając to posunięcie jako ingerencję w jego sprawę sądową lub próbę wpłynięcia na wybory w Mołdawii. Kiedy pojawiła się druga lista „problematicznych” kanałów, magnat stwierdził, że prawie wszystkie z nich są legalne i nie naruszają zasad Telegramu.

„Odmówiliśmy zastosowania się do tych zaleceń” – powiedział magnat. „Telegram opowiada się za wolnością słowa. Nie będziemy usuwać treści z powodów politycznych i będę nadal ujawniać każdą próbę zastraszania naszej platformy”.

Zarzuty Durova pojawiają się w momencie, gdy Mołdawia przygotowuje się do gorących wyborów parlamentarnych, w których proeuropejska Partia Działania i Solidarności prezydent Mai Sandu zmierzy się z Patriotycznym Blokiem Wyborczym.

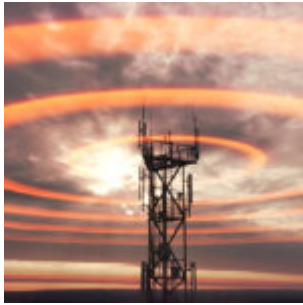
W marcu Durov otrzymał pozwolenie na opuszczenie Francji i prawdopodobnie powrócił do swojej rezydencji w Dubaju. Rząd Zjednoczonych Emiratów Arabskich, którego obywatelem Durov został w 2021 roku, podobno bacznie obserwuje jego głośną sprawę sądową.

Po jego aresztowaniu Telegram wywołał falę poruszenia wśród

użytkowników, wprowadzając zmiany w polityce prywatności. Platforma, niegdyś zagorzały obrońca anonimowości użytkowników, ogłosiła, że będzie przekazywać dane, takie jak adresy IP i numery telefonów, organom ścigania, jeśli otrzyma ważne nakazy prawne. Było to odejście od dotychczasowego stanowiska, które ograniczało udostępnianie danych do spraw związanych z terroryzmem, a nawet wtedy Telegram twierdził, że nigdy nie zastosował się do tych nakazów.

W styczniu Durov powiedział prokuratorom, że „w pełni rozumie powagę” zarzutów wobec niego, które obejmują przestępstwa związane z działalnością Telegramu, takie jak ekstremizm i wykorzystywanie dzieci. W tym samym miesiącu gigant komunikacyjny, od dawna znany z żelaznych zabezpieczeń prywatności, zintensyfikował współpracę z władzami amerykańskimi.

**Bezpieczeństwo 5G pod znakiem
zapytania: FCC podejmuje
działania mające na celu
ograniczenie ocen
środowiskowych pomimo obaw
dotyczących zdrowia**



- Federalna Komisja Łączności (FCC) planuje zwolnić określone wieże komórkowe i instalacje 5G z ocen środowiskowych i konserwatorskich.
- Krytycy twierdzą, że posunięcie to osłabia ochronę zapewnianą przez ustawę National Environmental Policy Act (NEPA).
- Niezależni eksperci łączą promieniowanie 5G z uszkodzeniami mózgu, demencją i zwiększonym ryzykiem zachorowania na raka.
- Technolodzy i zwolennicy ochrony zdrowia wzywają do wprowadzenia bardziej rygorystycznych przepisów dotyczących promieniowania RF i moratorium na wdrażanie 5G.
- Przed wydaniem ostatecznej decyzji w październiku 2025 r. zaprasza się do zgłaszania uwag publicznych.

Federalna Komisja Łączności (FCC) podejmuje działania mające na celu „usprawnienie” procesu oceny środowiskowej niektórych projektów bezprzewodowych, aby „promować wydajność i pewność” dla deweloperów telekomunikacyjnych. Proponowane przepisy zwalniałyby określone rodzaje instalacji, w tym małe anteny komórkowe, wieże komórkowe poniżej 200 stóp i wiejskie instalacje 5G, z przeglądów przeprowadzanych na mocy ustawy o polityce środowiskowej (NEPA) i ustawy o ochronie zabytków (NHPA). Dzieje się to w czasie, gdy niezależni eksperci ostrzegają przed potencjalnymi zagrożeniami dla zdrowia wynikającymi z promieniowania 5G, w tym zwiększonym ryzykiem raka mózgu, demencji i choroby Alzheimera.

Posunięcie FCC: usprawnienie czy zniesienie zabezpieczeń?

Krytycy, tacy jak Miriam Eckenfels z programu Children's Health Defense's Electromagnetic Radiation (EMR) & Wireless Program, twierdzą, że [nowe przepisy FCC mają więcej wspólnego z przejęciem kontroli przez przemysł niż z ochroną konsumentów](#). „NEPA jest jednym z niewielu zabezpieczeń przed niekontrolowaną ekspansją przemysłu” – powiedziała Eckenfels. „FCC traktuje NEPA tak, jakby zgodność z nią była opcjonalna”.

Proponowane zmiany [zwolnią z wymogów przeglądu NEPA projekty bezprzewodowe, które nie wymagają rejestracji konstrukcji antenowej, takie jak wieże komórkowe poniżej 200 stóp i małe instalacje komórkowe](#). Posunięcie to następuje w kontekście rosnących obaw społecznych dotyczących wpływu technologii 5G na zdrowie. W zeszłym roku dziewięciu niezależnych ekspertów w dziedzinie promieniowania o częstotliwości radiowej opublikowało wyniki badań w czasopiśmie *Annals of Clinical and Medical Case Reports*, wzywając władze do wprowadzenia bardziej rygorystycznych przepisów dotyczących promieniowania RF i wstrzymania wdrażania technologii 5G.

Zagrożenia dla zdrowia: naukowe podstawy potencjalnych zagrożeń związanych z 5G

Badanie finansowane przez Narodowy Instytut Zdrowia (NIH) sugeruje, że promieniowanie o częstotliwości radiowej wywołane przez 5G może prowadzić do uszkodzenia mózgu i prawdopodobnie zwiększać ryzyko demencji i choroby Alzheimera. Niejonizujące promieniowanie RF-EMF może zachowywać się jak promieniowanie jonizujące o wysokim liniowym transferze energii, uszkadzając żywe komórki skóry i generując wolne rodniki w niewielkiej odległości. Wolne rodniki mogą powodować stres oksydacyjny, który zwiększa ryzyko zachorowania na raka poprzez hamowanie apoptozy, sprzyjanie proliferacji i stymulowanie angiogenezy.

W recenzowanym artykule 9 ekspertów stwierdziło: „Nadal istnieją kontrowersje dotyczące bezpieczeństwa technologii 5G. Jednak dowody pochodzące z badań laboratoryjnych i badań przypadków ludzkich wskazują na potencjalne ryzyko niekorzystnych skutków zdrowotnych wynikających z narażenia na RF-EMF”. Zalecają oni wprowadzenie surowych regulacji i podkreślają znaczenie bezstronnego zespołu wykwalifikowanych naukowców, który dokładnie oceni ryzyko przed przystąpieniem do wdrażania 5G.

NEPA i ograniczone zabezpieczenia środowiskowe

NEPA została uchwalona w 1970 r. w odpowiedzi na rosnące obawy dotyczące środowiska. Ustanowiła ona proces wymagający od agencji federalnych rozważenia wpływu proponowanych działań na środowisko, promując przejrzystość i zaangażowanie społeczeństwa. Jednak niedawne ograniczenie zakresu NEPA przez Sąd Najwyższy, wraz z decyzją Kongresu z 2023 r. o ograniczeniu niektórych przeglądów NEPA, zachęciło FCC do przyspieszenia wdrażania infrastruktury bezprzewodowej.

„Chodzi o znalezienie właściwej równowagi” – powiedział w oświadczeniu przewodniczący FCC Brendan Carr. „Przestarzałe przepisy dotyczące ochrony środowiska opóźniają wdrażanie łączności szerokopasmowych. Musimy uprościć procesy wydawania zezwoleń i utorować drogę dla budowy nowej infrastruktury”.

Reakcja opinii publicznej i trwająca debata

Dziesiątki grup i zwolenników bezpiecznych technologii przedłożyły FCC zbiorowe uwagi, wzywając agencję do rozszerzenia i wzmocnienia, a nie ograniczenia, procedur przeglądu środowiskowego. „Musimy chronić zdrowie ludzi i dać im głos w sprawach dotyczących ich środowiska” – argumentowała Wilkens.

Prawniczka specjalizująca się w technologii Odette Wilkens stwierdziła: „Ogłoszenie FCC dotyczące proponowanych zmian w przepisach mających na celu „modernizację” NEPA jest sarkastycznym odniesieniem, ponieważ jego celem jest ograniczenie przeglądu środowiskowego i przeglądu ochrony zabytków”.

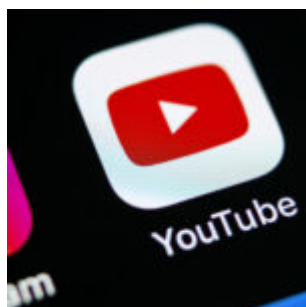
Równowaga między postępem a ochroną

Propozycja FCC dotycząca usprawnienia ocen środowiskowych projektów infrastruktury bezprzewodowej jest kontrowersyjną kwestią, która stawia postęp technologiczny w opozycji do obaw dotyczących zdrowia publicznego. Podczas gdy FCC dąży do promowania wydajności i pewności dla deweloperów telekomunikacyjnych, krytycy twierdzą, że posunięcie to osłabia kluczowe zabezpieczenia środowiskowe i ogranicza wpływ opinii publicznej. W miarę rozwoju debaty zainteresowane strony będą musiały dokładnie rozważyć długoterminowe skutki wdrożenia 5G dla środowiska i zdrowia. Opinie publiczne będą miały kluczowe znaczenie dla kształtowania ostatecznej wersji przepisów, a stawka nigdy nie była wyższa.

„Wdrożenie 5G nie może odbywać się kosztem zdrowia publicznego i ochrony środowiska. Naszym priorytetem powinno być osiągnięcie równowagi między innowacyjnością a bezpieczeństwem”.

**Google przywraca
zablokowanych twórców YouTube**

po przyznaniu się do nacisków Bidena na cenzurowanie treści dotyczących COVID



- YouTube zmienia politykę cenzury po przyznaniu się do bezprawnych nacisków rządu na wyciszenie dyskusji dotyczących COVID-19 i uczciwości wyborów.
- Wewnętrzne dokumenty ujawniają, że urzędnicy Bidena domagali się usunięcia z platform mediów społecznościowych informacji o faktycznych skutkach ubocznych szczepionek oraz odmiennych poglądów.
- Komisja Sprawiedliwości Izby Reprezentantów ujawnia zмовę między wielkimi firmami technologicznymi a agencjami federalnymi w celu tłumienia legalnych wypowiedzi na temat szczepionek, wyborów i alternatywnych metod leczenia.
- YouTube przywraca zablokowanych twórców, takich jak Dan Bongino, sygnalizując częściowe zwycięstwo wolności słowa w obliczu trwającej kontroli praktyk cenzury.
- Walka trwa, ponieważ krytycy ostrzegają przed utrzymującymi się taktykami tłumienia pomimo zmian w polityce, podkreślając potrzebę utrzymania presji publicznej i przejrzystości.

Cyfrowe mury cenzury w końcu się rozpadają. Po latach tłumienia odmiennych opinii na temat COVID-19 i uczciwości wyborów, YouTube ogłosił w tym tygodniu, że pozwoli wcześniej

zbanowanym twórcom powrócić na platformę, co jest zaskakującym zwrotem sytuacji, który nastąpił dopiero po tym, jak firma przyznała, że spotkała się z „niedopuszczalną i niewłaściwą” presją ze strony administracji Bidena, aby cenzurować legalne wypowiedzi.

Ta zmiana nastąpiła po nieustannej kontroli ze strony Komisji Sądowniczej Izby Reprezentantów, która ujawniła skoordynowaną kampanię federalnych urzędników i wielkich firm technologicznych mającą na celu wyciszenie poglądów sprzecznych z narracją rządową. W liście z 23 września skierowanym do przewodniczącego komisji Jima Jordana (R-Ohio) zespół prawny Google potwierdził to, co krytycy podejrzewali od dawna: Biały Dom Bidena wielokrotnie wywierał presję na YouTube, aby usuwał treści, które nie naruszały jego własnych zasad.

„Niedopuszczalne i niewłaściwe jest, gdy jakikolwiek rząd, w tym administracja Bidena, próbuje dyktować firmie, w jaki sposób ma moderować treści” – stwierdzono w liście, dodając, że Google „konsekwentnie walczyło z tymi działaniami na podstawie pierwszej poprawki”.

Ujawniono schemat cenzury politycznej

[Przyznanie się](#) do tego stanowi rzadki moment odpowiedzialności w długotrwałej walce branży technologicznej o wolność słowa. Według wewnętrznej korespondencji i dokumentów sądowych urzędnicy administracji Bidena, w tym prawniczka Białego Domu Dana Remus, domagali się od platform takich jak YouTube i Facebook usunięcia postów kwestionujących szczepionki przeciwko COVID-19, zarzuty dotyczące oszustw wyborczych, a nawet żarty na temat polityki pandemicznej. Dyrektor generalny Meta Mark Zuckerberg przyznał w 2024 r., że jego firma spotkała się z podobną presją, pisząc: „Uważam, że presja ze strony rządu była niewłaściwa i żałuję, że nie wyraziliśmy się

na ten temat bardziej otwarcie”.

Polityka YouTube ewoluowała wraz ze zmianami wytycznych rządowych podczas pandemii, często [ograniczając dyskusje](#), które później okazały się słuszne. Teraz, po zniesieniu tych zasad, platforma zaprasza z powrotem twórców, takich jak konserwatywny komentator Dan Bongino, którego kanał został usunięty z powodu treści związanych z COVID. Jordan świątował to posunięcie jako „kolejne zwycięstwo w walce z cenzurą”, zauważając, że zarówno znane osobistości, jak i zwykli użytkownicy otrzymają drugą szansę.

Szersza wojna z dysydentami

Ta zmiana nie nastąpiła w próżni. Jest ona częścią szerszych, orwellowskich wysiłków mających na celu kontrolę informacji – wysiłków, które wykraczały poza COVID i obejmowały wybory, odżywianie, a nawet satyrę. Wewnętrzny e-mail Facebooka z 2021 r. ujawnił, że Biały Dom Bidena chciał usunięcia „negatywnych informacji lub opinii na temat szczepionki”, w tym „prawdziwych informacji o skutkach ubocznych”. YouTube zaproponował później zasady cenzurowania krytyki szczepionek, co początkowo spotkało się z aprobatą Białego Domu.

Liczący prawie 900 stron raport Komisji Sądowniczej na temat „kompleksu cenzury przemysłowej” szczegółowo opisuje, w jaki sposób agencje federalne i wielkie firmy technologiczne współpracowały, aby tłumić nie tylko „dezinformację”, ale także merytoryczne dyskusje na temat ryzyka związanego ze szczepionkami, alternatywnych metod leczenia i nieprawidłowości wyborczych. Nawet treści dotyczące diety ketogenicznej i przerywanego postu spotkały się z cichym blokowaniem po tym, jak YouTube dostosował swoje zasady do wytycznych CDC i WHO.

Kruche zwycięstwo wolności słowa

Chociaż ustępstwo Google jest krokiem naprzód, walka jest daleka od zakończenia. W piśmie firma twierdzi, że nigdy nie korzystała z usług zewnętrznych weryfikatorów faktów do monitorowania treści – twierdzenie to jest sprzeczne z wieloletnimi doniesieniami użytkowników o blokowanych filmach i kanałach pozbawionych możliwości zarabiania. I chociaż YouTube twierdzi obecnie, że „ceni konserwatywne głosy”, jego dotychczasowe działania sugerują, że sceptycyzm jest uzasadniony.

Niemniej jednak zmiana stanowiska daje promyk nadziei. W miarę jak komisja Jordana kontynuuje swoje dochodzenie, przesłanie jest jasne: gdy rząd i korporacje spiskują, aby uciszyć debatę, przejrzystość i presja społeczna mogą wymusić odpowiedzialność. Na razie cyfrowy rynek jest nieco bardziej otwarty, ale walka o utrzymanie tego stanu dopiero się rozpoczęła.

**Twórcy wież komórkowych
wprowadzają w błąd dla zysku,
eksperci ds. zdrowia
ostrzegają przed zagrożeniami
związanymi z 5G**



- Twórcy wież komórkowych często wprowadzają w błąd mieszkańców i urzędników, twierdząc, że istnieją luki w zasięgu, podczas gdy mapy FCC pokazują pełny zasięg.
- Eksperci, tacy jak dr Martin Pall, łączą narażenie na pola elektromagnetyczne 5G z przedwczesnym starzeniem się, chorobą Alzheimera, uszkodzeniami mózgu i zwiększonym ryzykiem zachorowania na raka.
- Adwokat Robert Berg radzi mieszkańcom, aby korzystali z krajowej mapy szerokopasmowej FCC w celu ujawnienia fałszywych twierdzeń, opowiadając się za zwiększeniem przejrzystości.
- Władze lokalne i społeczności walczą z budową wież komórkowych 5G, uzyskując różne wyniki sądowe.
- Badania dr. Palla pokazują, że pulsujące pola elektromagnetyczne 5G mogą znacznie zwiększać stres oksydacyjny i uszkodzenia komórek.

W odkrywczym opracowaniu adwokat Robert Berg twierdzi, że firmy telekomunikacyjne często wprowadzają w błąd lokalnych mieszkańców i urzędników co do potrzeby budowy nowych wież komórkowych. Według Berga, [deweloperzy często przedstawiają fałszywe mapy propagacji fal radiowych, które pokazują luki w zasięgu, nawet jeśli krajowa mapa szerokopasmowa Federalnej Komisji Łączności \(FCC\) wskazuje pełny zasięg](#). Ta strategiczna nieuczciwość, jak twierdzi, ma na celu przekonanie władz do zatwierdzenia niepotrzebnych budów wież, a tym samym zapewnienie lukratywnych kontraktów.

„Widziałem przypadki, w których operator informował lokalnych liderów społeczności o znacznych lukach w zasięgu i kierował

ich do „fałszywej” mapy propagacji fal radiowych” – powiedział Berg. „Kiedy porównują mapę operatora z krajową mapą szerokopasmową FCC, często stwierdzają, że na tej mapie zasięg operatora w danym obszarze geograficznym wynosi 100%”.

Niebezpieczeństwa związane z 5G: eksperci ds. zdrowia biją na alarm

Wraz z przyspieszeniem wdrażania technologii 5G eksperci ds. zdrowia zgłaszają obawy dotyczące potencjalnych zagrożeń dla zdrowia wynikających z pól elektromagnetycznych (EMF) o częstotliwości radiowej (RF) emitowanych przez tę nową generację technologii bezprzewodowej. Dr Martin Pall, pionier w badaniach nad EMF, powiązał narażenie na EMF 5G z przedwczesnym starzeniem się, chorobą Alzheimera, uszkodzeniami mózgu i zwiększonym ryzykiem zachorowania na raka.

„Pola elektromagnetyczne działają poprzez szczytowe siły elektryczne i zmienne w czasie siły magnetyczne w skali nanosekund” – wyjaśnił dr Pall. „Wraz z każdym wzrostem modulacji impulsowej wytwarzanej przez urządzenia emitujące pola elektromagnetyczne, takie jak inteligentne liczniki, inteligentne telefony komórkowe i technologia 5G, szczyty te znacznie się zwiększają, powodując to, co opisuję jako ostateczny koszmar – niezwykle wczesne wystąpienie choroby Alzheimera”.

Badania sugerują, że promieniowanie fal impulsowych 5G może być bardziej szkodliwe niż fale ciągłe, powodując znaczne uszkodzenia mózgu, układu odpornościowego i nerwowego. Wykazano, że te fale impulsowe, charakteryzujące się krótkimi impulsami o dużej mocy RF, zakłócają synapsy dopaminowe w hipokampie, potencjalnie prowadząc do chorób takich jak choroba Parkinsona.

Walka o przejrzystość: rola mapy FCC

Mieszkańcy i ich rzecznicy, tacy jak Berg, traktują krajową mapę szerokopasmową FCC jako potężne narzędzie do przeciwdziałania wprowadzającym w błąd twierdzeniom firm telekomunikacyjnych. Porównując proponowaną przez operatora mapę propagacji fal radiowych ze szczegółową mapą szerokopasmową FCC, krytycy mogą ujawnić niespójności, które wskazują na pełny zasięg w obszarach, w których rzekomo go brakuje.

„Jeśli porównasz mapę operatora z krajową mapą szerokopasmową FCC, często zobaczysz, że na krajowej mapie szerokopasmowej zasięg wynosi 100%” – powiedział Berg. „Fakt ten powinien mieć znaczenie, ale nie zawsze tak jest”.

Mimo to Berg przyznaje, że zwalczanie wprowadzających w błąd informacji stanowi wyzwanie. W niedawnej sprawie w hrabstwie Santa Cruz w Kalifornii prawie 20 mieszkańców zeznało, że mają już silny sygnał, a mapa FCC nie wykazała żadnych luk w zasięgu. Jednak hrabstwo nadal zatwierdziło budowę wieży komórkowej o wysokości 140 stóp, co doprowadziło do sporu prawnego wspieranego przez organizację Children’s Health Defense (CHD) w sprawach dotyczących promieniowania elektromagnetycznego (EMR) i sieci bezprzewodowych.

Prawne rozstrzygnięcia i przyszłość wdrażania 5G

Sytuacja prawna dotycząca wdrażania infrastruktury bezprzewodowej jest coraz bardziej napięta z powodu konfliktów między deweloperami a władzami lokalnymi. [Seria ostatnich orzeczeń w Wirginii, na Florydzie i w New Hampshire podkreśla złożoną wzajemną zależność między prawem federalnym a lokalnymi rozporządzeniami.](#)

„Operatorzy sieci bezprzewodowych i deweloperzy infrastruktury ścigają się, aby sprostać rosnącym wymaganiom w zakresie 5G i

łączy szerokopasmowych” – zauważa Berg. „Sprawy te ilustrują delikatną równowagę, jaką obie strony muszą utrzymać, ponieważ sądy coraz częściej egzekwują normy federalne wbrew lokalnym sprzeciwom”.

Na przykład w sprawie w New Hampshire lokalne władze odmówiły wydania pozwolenia na budowę wieży o wysokości 150 stóp, powołując się na obawy dotyczące wpływu na wygląd okolicy i wartość nieruchomości. Jednak kwestia, czy ograniczenie wysokości skutecznie uniemożliwia świadczenie usług bezprzewodowych, co stanowi naruszenie prawa federalnego, pozostaje nierozwiązana. Takie subtelne argumenty prawne prawdopodobnie będą miały decydujący wpływ na przyszły kierunek rozwoju infrastruktury 5G.

Niepewna przyszłość

Dążenie do powszechnego wdrożenia 5G budzi poważne obawy dotyczące zdrowia publicznego, zwłaszcza w świetle nowych badań łączących ekspozycję na pola elektromagnetyczne z poważnymi problemami zdrowotnymi. W miarę jak firmy telekomunikacyjne realizują program napędzany zyskiem i postępem technologicznym, na decydentach politycznych, pracownikach służby zdrowia i społeczeństwie spoczywa obowiązek domagania się przejrzystości, egzekwowania istniejących przepisów i ponownego rozważenia norm bezpieczeństwa dotyczących ekspozycji na promieniowanie.

„Ostatecznie wszyscy jesteśmy częścią szerszej dyskusji na temat równowagi między postępem technologicznym a zdrowiem publicznym” – powiedział Berg. „A stawka nigdy nie była wyższa”.

Chińskie pojazdy elektryczne mogą stać się zdalnie sterowaną bronią w wojnie hybrydowej



- Chińskie pojazdy elektryczne stanowią poważne zagrożenie dla cyberbezpieczeństwa dzięki wbudowanym funkcjom nadzoru i zdalnego uzbrojenia.
- Eksperci ds. cyberbezpieczeństwa ostrzegają, że pojazdy te mogą zostać przejęte w celu wyłączenia funkcji bezpieczeństwa, wywołania eksplozji lub sparaliżowania miast podczas konfliktów.
- W samym tylko sierpniu w Australii sprzedano ponad 20 000 chińskich pojazdów elektrycznych wyposażonych w ryzykowne moduły komórkowego internetu rzeczy (CIM).
- Wzywa się rządy do zakazania chińskim producentom pojazdów elektrycznych udziału w przetargach publicznych, chyba że zezwolą oni na kontrolę kodu źródłowego i audyt danych.
- Kraje zachodnie stoją w obliczu rosnącego ryzyka, ponieważ 80% pojazdów elektrycznych w Australii jest obecnie produkowanych w Chinach, a podobne trendy pojawiają się na całym świecie.

Wyobraź sobie, że jedziesz ruchliwą autostradą, gdy nagle bateria Twojego pojazdu elektrycznego przegrzewa się – lub, co

gorsza, eksploduje. Teraz wyobraź sobie, że dzieje się to jednocześnie z tysiącami samochodów, paraliżując ruch, powodując masowe ofiary śmiertelne i pogrążając miasta w chaosie.

Nie jest to fabuła dystopijnego thrillera, ale bardzo realne zagrożenie stwarzane przez pojazdy elektryczne produkowane w Chinach, według eksperta ds. cyberbezpieczeństwa Alastaira MacGibbona.

Podczas tegorocznego szczytu Australian Financial Review Cyber Summit MacGibbon, były doradca premiera Malcolma Turnbulla ds. cyberbezpieczeństwa, wydał mrożące krew w żyłach ostrzeżenie: chińskie pojazdy elektryczne eksploatowane w Australii mogą zostać zdalnie wykorzystane jako broń w ramach kampanii wojny hybrydowej. Wyjaśnił, że pojazdy te nie służą wyłącznie do transportu; są to [urządzenia monitorujące z wbudowanymi wyłącznikami awaryjnymi](#), które mogą zostać przejęte przez zagranicznych przeciwników w celu wyłączenia funkcji bezpieczeństwa, wywołania eksplozji, a nawet sparaliżowania całego miasta w godzinach szczytu.

MacGibbon, obecnie dyrektor ds. strategii w CyberCX, nie przebierał w słowach. „Samochody, o których mówimy, niezależnie od tego, czy są elektryczne, czy nie, są urządzeniami podsłuchowymi i urządzeniami monitorującymi w postaci kamer” – stwierdził. Jednak zagrożenia wykraczają daleko poza naruszenie prywatności. „Wystarczy wyłączyć funkcje bezpieczeństwa domowych akumulatorów, aby się przeładowały. Wystarczy wyłączyć te same funkcje bezpieczeństwa w pojazdach elektrycznych. Wystarczy wyłączyć je u producenta, aby pojazdy te eksplodowały. Wystarczy ograniczyć ich zdolność do poruszania się w godzinach szczytu w wybranych miastach”.

Jego ostrzeżenia są zgodne z coraz większą liczbą dowodów sugerujących, że chińska strategia „wojny totalnej” wykorzystuje podłączone do sieci urządzenia jako broń. Od

handlu fentanylem po stosowanie broni biologicznej – Komunistyczna Partia Chin (KPCh) wykazała gotowość do wykorzystywania słabości technologicznych w celu destabilizacji przeciwników. Obecnie, gdy cztery chińskie marki pojazdów elektrycznych – BYD, GWM, MG Motor i Chery – po raz pierwszy znalazły się w pierwszej dziesiątce najlepiej sprzedających się marek w Australii, ryzyko to nie jest już tylko teoretyczne. Dziesiątki tysięcy tych pojazdów jeździ już po australijskich drogach, wyposażonych w moduły komórkowego internetu rzeczy (CIM), które przesyłają dane do chińskich serwerów.

Trojan na kołach?

[Zagrożenie](#) nie jest tylko hipotetyczne. Tylko w sierpniu Australijczycy kupili ponad 20 000 pojazdów wyprodukowanych w Chinach, a BYD wyprzedził Mitsubishi, stając się szóstą najlepiej sprzedającą się marką samochodów w kraju. Pojazdy te są wyposażone w moduły CIM, niewielkie, ale potężne urządzenia, które umożliwiają zdalne aktualizacje oprogramowania, udostępnianie danych o ruchu drogowym, a nawet śledzenie lokalizacji w czasie rzeczywistym. Według raportu China Strategic Risks Institute (CSRI) moduły te tworzą tylne drzwi dla cyberataków, umożliwiając Pekinowi unieruchomienie pojazdów, zbieranie poufnych danych, a nawet sabotowanie krytycznej infrastruktury w przypadku konfliktu.

Chiny mogłyby zdalnie unieruchomić floty pojazdów elektrycznych, unieruchamiając kierowców, zakłócając działanie służb ratowniczych lub nawet gorzej. Think tank zaleca, aby rządy zakazały chińskim producentom pojazdów elektrycznych udziału w przetargach publicznych, chyba że poddadzą się kontroli kodu źródłowego i regularnym audytom swoich globalnych centrów danych – co jest żądaniem, którego Pekin prawdopodobnie nie spełni.

Ostrzeżenia MacGibbona wykraczają poza pojazdy elektryczne.

„Potencjalnie miliony urządzeń [Internetu rzeczy] lub urządzeń podłączonych do sieci – nie wyprodukowanych w Chinach, ale kontrolowanych przez Chiny – znajdują się w naszych systemach” – zauważył. Obejmuje to panele słoneczne, domowe akumulatory, a nawet dachowe podgrzewacze słoneczne, z których wiele zostało już oznaczonych jako podejrzane ze względu na wbudowaną technologię.

Czas się obudzić

Australia nie jest jedynym krajem, który bije na alarm. Stany Zjednoczone zaproponowały zakaz importu chińskich pojazdów elektrycznych z powodu podobnych obaw, a brytyjscy urzędnicy są naciskani, aby zobowiązać zagranicznych dostawców do nieprzekazywania danych w żadnych okolicznościach.

Jednak pomimo ostrzeżeń chińskie pojazdy elektryczne nadal zalewają rynki zachodnie dzięki niskim cenom i agresywnemu marketingowi. Ponieważ 80% pojazdów elektrycznych w Australii jest obecnie produkowanych w Chinach, pytanie nie brzmi, czy pojazdy te mogą zostać wykorzystane jako broń, ale kiedy i jak poważne będą tego skutki.

Konsekwencje są ogromne. W scenariuszu wojny hybrydowej Pekin nie musiałby wystrzelić ani jednej rakiety, aby sparaliżować miasto. Kilka linii złośliwego kodu mogłoby unieruchomić hamulce, wywołać pożary lub zamienić tysiące samochodów w blokady dróg, wywołując masową panikę i szkody gospodarcze.

Chińskie pojazdy elektryczne to nie tylko samochody – to potencjalna broń. Jeśli rządy zachodnie nie podejmą natychmiastowych działań, wkrótce mogą znaleźć się w sytuacji, w której zostaną przechytrzone przez konia trojańskiego na kołach.