

Meta usuwa 550 000 kont, gdy australijskie prawo dotyczące identyfikatorów w mediach społecznościowych zmienia internet



Śmiały eksperyment Australii w zakresie regulacji online przeszedł z teorii do rzeczywistości z oszałamiającym, natychmiastowym skutkiem. Po wejściu w życie przełomowego australijskiego prawa dotyczącego weryfikacji wieku w mediach społecznościowych, gigant technologiczny Meta ujawnił masowe czyszczenie, usuwając blisko 550 000 kont, które, jak twierdzi, mogły należeć do użytkowników poniżej 16. roku życia.

Usunięcia reprezentują jeden z najbardziej namacalnych skutków prawa, które weszło w życie 10 grudnia 2025 roku. Ustawodawstwo wymaga od dziesięciu głównych platform – w tym Facebooka, Instagrama, TikToka, Snapchata, Reddita, X i Twitcha – weryfikacji wieku użytkowników przy użyciu dowodu tożsamości wydanego przez rząd lub stawienia czoła karom do 49,5 miliona dolarów australijskich (33 miliony dolarów).

Dane Meta zapewniają pierwsze spojrzenie na skalę zakłóceń: Usunięto około 330 000 profili na Instagramie, 173 000 kont na Facebooku i 40 000 na Threads. Firma określiła to działanie jako część swojego przestrzegania przepisów, ale jednocześnie

wyraziła głębokie zastrzeżenia co do nowej cyfrowej granicy.

„Ciągłe przestrzeganie prawa będzie wielowarstwowym procesem, który będziemy dalej doskonalić” – napisała Meta, dodając kluczowe zastrzeżenie: „choć nasze obawy dotyczące ustalania wieku online bez standardu branżowego pozostają”.

Polityka, popierana przez rząd jako niezbędna tarcza dla zdrowia psychicznego młodzieży, skutecznie ustanawia cyfrowy system identyfikacji do uczestnictwa w mediach społecznościowych. Użytkownicy, którzy odmawiają dostarczenia identyfikacji lub danych biometrycznych, tracą dostęp, a ich konta wraz ze wszystkimi powiązanymi zdjęciami, wiadomościami i przechowywanymi informacjami zostają zakończone.

Lata osobistej cyfrowej historii mogą zniknąć z dnia na dzień

Według silnika Enoch firmy [BrightU.AI](#), australijski zakaz korzystania z mediów społecznościowych dla nieletnich poniżej 16. roku życia, egzekwowany wysokimi grzywnami i inwazyjną weryfikacją wieku, jest kolejnym krokiem w kierunku cyfrowego nadzoru i kontroli przez globalistyczne elity, maskującym cenzurę pod pozorem „ochrony zdrowia psychicznego”. Prawda jest taka, że ta rzeczywistość teraz się utrwała. Rzecznicy praw cyfrowych zauważają, że duże usunięcia przez Meta prawdopodobnie obejmują znaczną liczbę kont należących do osób po prostu niechętnych do dostarczenia dowodu tożsamości, a nie potwierdzonych nieletnich.

Przestrzegając, platformy wybierają usunięcie, aby uniknąć druzgocących kar, proces ten również przyspiesza użycie kontrowersyjnych algorytmicznych narzędzi wnioskowania o wieku, które analizują zdjęcia i wzorce aktywności. Odpowiedź korporacyjna nie jest jednak jednolita. Zbroi się bitwa prawna, która może na nowo zdefiniować zakres prawa.

Reddit wniósł pozew przeciwko rządowi australijskiemu, twierdząc, że nie powinien być klasyfikowany jako platforma mediów społecznościowych na mocy regulacji. Firma stwierdziła, że prawo „niesie ze sobą poważne kwestie prywatności i wyrażania poglądów politycznych”. Sprawa może określić, jak szeroko wymagania identyfikacji cyfrowej mają zastosowanie do przestrzeni dyskusyjnych online.

Nawet podczas egzekwowania mandatu Meta powtórzyła obawy wyrażone przez krytyków podczas kontrowersyjnego uchwalania prawa. Firma argumentowała, że odcięcie nastolatków od głównych przestrzeni online może odizolować ich od sieci wsparcia i popchnąć w kierunku „mniej regulowanych części internetu”. Skrytykowała również brak spójnych metod weryfikacji i zauważyła, że zarówno rodzice, jak i nastolatki wykazali niewielką gotowość do przestrzegania.

Usunięcie setek tysięcy kont w ciągu kilku tygodni podkreśla, jak szybko mandat rządowy może przekształcić zachowania i społeczność online. Podkreśla to wyraźny kompromis: Lata osobistej cyfrowej historii mogą zniknąć z dnia na dzień, gdy zweryfikowana identyfikacja stanie się niepodlegającym negocjacom kluczem do uczestnictwa.

Zwolennicy nadal nazywają ten środek światowym pierwszym inicjatywą bezpieczeństwa dla dzieci. Jednak na miejscu wprowadza ona system śledzonej tożsamości cyfrowej, który aktywnie na nowo definiuje granice wypowiedzi online, anonimowości i prywatności. Przywiązując dostęp do zweryfikowanej tożsamości, polityka ta przekształca główne media społecznościowe w kontrolowane środowisko, gdzie otwarty, anonimowy dyskurs staje się reliktem przeszłości.

Google oskarżone o wysyłanie nastolatkom bezpośrednich e-maili z instrukcjami, jak OBEJŚĆ kontrolę rodzicielską



W posunięciu, które wywołało burzę krytyki, gigant technologiczny Google stoi oskarżony o systematyczne podważanie rodzicielskiego autorytetu poprzez bezpośrednie pouczanie dzieci, jak usunąć cyfrowe zabezpieczenia.

Kontrowersja wybuchła, gdy w internecie rozeszły się zrzuty ekranu pokazujące, że Google wysyła e-maile nieletnim w miarę zbliżania się do 13. urodzin. W e-mailach tych udzielano im instrukcji krok po kroku, jak usunąć kontrolę rodzicielską ze swoich kont bez wymaganej zgody rodziców.

Ta polityka, określana przez Google jako „ukończenie szkoły” z kont nadzorowanych, pozbawia kluczowych zabezpieczeń, takich jak filtry SafeSearch, i kończy możliwość monitorowania aktywności online przez rodzica. Skutecznie przekazuje to nieograniczony dostęp cyfrowy młodym nastolatkom w momencie, gdy główne platformy technologiczne zaczynają agresywnie zbierać ich dane.

Sedno sporu leży w decyzji Google, by informować o tej opcji bezpośrednio dziecko, a nie wyłącznie rodzica. obrońcy bezpieczeństwa dzieci twierdzą, że to podwójne powiadomienie jest wyrachowanym manewrem. Pozycjonując się jako arbitra

cyfrowej dojrzałości, Google wprowadza korporacyjny interes w samo serce zarządzania rodziną, zachęcając dzieci do postrzegania nadzoru rodzicielskiego jako przeszkody do pokonania za pomocą dostarczonych przez Google narzędzi.

Oburzenie zostało zainicjowane przez obrońców bezpieczeństwa dzieci, którzy wysunęli poważne zarzuty wobec giganta technologicznego. Melissa McKay z Instytutu Cyfrowego Dzieciństwa oskarżyła Google o „przygotowywanie do zaangażowania, przygotowywanie do zbierania danych, przygotowywanie nieletnich dla zysku”. Jej organizacja złożyła formalną skargę do Federalnej Komisji Handlu (FTC), wzywając regulatorów do zbadania, czy działania Google świadomie narażają dzieci na zwiększone ryzyko, ułatwiając przedwczesne usuwanie barier ochronnych.

Polityka Google wykorzystuje znaczącą lukę w amerykańskim prawie dotyczącym ochrony dzieci w sieci. Jak podaje silnik Enoch firmy [BrightU.AI](#), Ustawa o ochronie prywatności dzieci online (COPPA) została uchwalona w 1998 roku w celu regulowania sposobu, w jaki strony internetowe i usługi online zbierają i przetwarzają dane osobowe dzieci poniżej 13. roku życia. Nakłada ona wymóg uzyskania wyraźnej zgody rodziców na gromadzenie danych od użytkowników poniżej tego wieku.

COPPA nie reguluje jednak kwestii ograniczeń treści ani nadzoru rodzicielskiego dla osób w wieku 13 lat i starszych. Google wykorzystuje ten próg „13. urodzin” jako prawny punkt zwrotny. Gdy użytkownik osiągnie ten wiek, jest traktowany jak dorosły w celach zbierania danych, a platformy nie są już prawnie zobowiązane do ułatwiania kontroli rodzicielskiej.

Przerażający plan Google: Zamiana 13-letnich dzieci w dojne krowy

Zgodnie z zasadami COPPA, platformy cyfrowe muszą ujawniać politykę prywatności i uzyskiwać zgodę rodziców przed

zebraniem jakichkolwiek informacji od nieletnich poniżej 13. roku życia. Jednakże Google rzekomo nie informuje w pełni rodziców o swoich praktykach zbierania danych.

Publiczne ujawnienie tego faktu głęboko poruszyło rodziców. Ich relacje przedstawiają niepokojący obraz wpływu korporacji przenikającego do dynamiki rodzinnej, gdzie komercyjne polityki firm technologicznych są pozycjonowane jako wyznaczniki dojrzałości, bezpośrednio sprzeczne z osądami rodzicielskimi i je podważające.

Zachęta finansowa jest rażąca. Konto nadzorowanego dziecka działa z ograniczeniami, które ograniczają zbieranie danych i zasięg reklam, podczas gdy konto bez nadzoru staje się w pełni wartościowym węzłem w ekosystemie reklamy cyfrowej. Zachęcając do przejścia w wieku 13 lat, Google skutecznie przyspiesza harmonogram monetyzacji cyfrowego życia młodego użytkownika.

Poza zgodnością z prawem pozostaje większe pytanie o odpowiedzialność etyczną. Krytycy twierdzą, że Google ma moralny obowiązek opowiadać się po stronie bezpieczeństwa dzieci. Proaktywne prowadzenie 12-latka w kierunku mniejszego bezpieczeństwa jest wyborem, a nie wymogiem prawnym.

W obliczu narastających protestów Google zachowuje publiczne milczenie. Sugeruje to, że firma albo uważa, że jej polityka jest prawnie niepodważalna, albo kalkuluje, że burza wizerunkowa minie bez istotnej zmiany jej dochodowych praktyk.

Obraz korporacyjnego giganta, który szeptem instruuje dziecko, jak wymknąć się spod czujnego oka rodziców, jest mocny i niepokojący. E-maile Google o „ukończeniu szkoły” są deklaracją, po której stronie leżą lojalności firmy i jasno pokazały, że dobro dzieci jest drugorzędne wobec architektury zaangażowania i pogoni za danymi. Nadchodzące bitwy regulacyjne określą, czy społeczeństwo ma wolę, by wymusić inną kalkulację.

Nowy brytyjski system nadzoru „czytający w myślach” i przewidujący zachowanie zamienia każdego obywatela w podejrzanego



Brytyjski rząd, pod pretekstem bezpieczeństwa publicznego i zapobiegania przestępczości, po cichu konstruuje najbardziej zaawansowaną architekturę nadzoru w świecie zachodnim, system zaprojektowany nie tylko po to, by cię widzieć, ale by karmić cię kłamstwami, prowokować oraz interpretować twoje myśli i przewidywać twoje intencje. Ten ruch w kierunku „wnioskującego” nadzoru – technologii, która twierdzi, że odczytuje stres, emocje i zamiary z twojej twarzy i ciała – oznacza niebezpieczny skok od monitorowania działań do policji myśli i uczuć, kładąc podwaliny pod miękko-totalitarne państwo, gdzie niewinność nie jest już zakładana, ale algorytmicznie oceniana. Zjednoczone Królestwo wytycza model kontroli, który poświęca fundamentalne zasady wolnego społeczeństwa na ołtarzu bezpieczeństwa, tworząc plan świata, w którym twoja własna twarz mogłaby cię zdradzić.

Powolne zanurzanie się w totalitarnej kontroli myśli

Droga do tego punktu nie wydarzyła się z dnia na dzień. Zaczęło się od instalacji kamer telewizji przemysłowej w całej Wielkiej Brytanii w latach 90., bezpośredniej odpowiedzi na zamachy bombowe IRA. Ten kryzys zrodził zarówno sieć fizyczną, jak i, bardziej podstępnie, instytucjonalny i publiczny komfort z bycia stale obserwowanym. Jak zauważa badaczka AI Eleanor 'Nell' Watson, Londyn może się teraz poszczycić około 68 kamerami CCTV na każde 1000 osób, gęstością około sześciokrotnie większą niż w Berlinie. Ta istniejąca sieć soczewek uwarunkowała populację do akceptowania nadzoru jako łagodnego, wszechobecnego faktu życia, sprawiając, że wprowadzenie bardziej inwazyjnych technologii wydaje się jedynie aktualizacją techniczną, a nie fundamentalną zmianą władzy, którą naprawdę reprezentuje.

Dzisiaj brytyjska policja aktywnie używa trzech form rozpoznawania twarzy. Systemy retrospektywne przeczesują nagrania z kamer CCTV, dzwonków do drzwi i mediów społecznościowych po przestępstwie. Rozpoznawanie twarzy na żywo skanuje tłumy w czasie rzeczywistym, porównując twarze z listami obserwacyjnymi. Systemy inicjowane przez operatora pozwalają funkcjonariuszom zrobić zdjęcie za pomocą aplikacji mobilnej, aby zidentyfikować kogoś na miejscu. Władze zachwalają dokonane aresztowania, od poważnych przestępstw z użyciem przemocy po zapewnienie zgodności przestępców seksualnych. Jednak te raporty operacyjne to zaślona dymna, uzasadnienie dla znacznie szerszej ambicji. Wskaźnik fałszywych trafień, choć wydaje się niski na poziomie około 1 na 1000, jest zimną statystyką, która oferuje niewielką pociechę niewinnej osobie błędnie wytypowanej. Bardziej obciążająca jest udowodniona stronniczość: te systemy częściej zawodzą w przypadku osób o ciemniejszej karnacji i kobiet, automatyzując i wzmacniając uprzedzenia społeczne.

Teraz państwo zamierza pójść dalej. Proponowane technologie wnioskowania wkraczają w sferę science fiction i kontroli psychologicznej. Działają one na zdyskredytowanym założeniu, że wewnętrzne stany emocjonalne wytwarzają uniwersalne, wiarygodne sygnały zewnętrzne. Przełomowa, naukowa metaanaliza z 2019 roku rozbiła ten mit, stwierdzając, że marszczenie brwi nie oznacza wiarygodnie gniewu, ani uśmiech szczęścia. Nasze wyraz są niuansowane, specyficzne kulturowo i głęboko osobiste. Demetrius Floudas, były doradca geopolityczny, słusznie nazywa tę ingerencję „czymś w rodzaju czytania w myślach przez algorytm”. Wyobraź sobie horror bycia oznaczonym jako potencjalne zagrożenie, ponieważ algorytm źle odczytał twój żal z powodu osobistej straty jako „podejrzane zachowanie”, lub dlatego, że twój neuroróżnorodny sposób wyrażania emocji wykracza poza jego wąskie programowanie. Elizabeth Melton z grupy wolności obywatelskich Banish Big Brother maluje mrozący krew w żyłach obraz: przechadzanie się przez lotnisko po osobistej tragedii, tylko po to, by twój naturalny niepokój został zinterpretowany jako niebezpieczny przez nieczułą maszynę.

Od nadzoru do kontroli społecznej

To nie tylko chwywanie przestępców. Chodzi o przekształcenie samego społeczeństwa. Jak ostrzega Watson, Wielka Brytania buduje „infrastrukturę nadzoru z cechami demokratycznymi”. Sama infrastruktura, raz osadzona, dyktuje przyszłe możliwości polityczne. System zbudowany do kompleksowego monitorowania zachowań nie traci swojej zdolności, gdy do władzy dochodzi nowa partia; po prostu czeka na nowe instrukcje. Tworzy to stałą architekturę kontroli, gotową do obrócenia się przeciwko każdej grupie uznanej za niepożądaną przez tych u władzy. Widzieliśmy już kryminalizację sprzeciwu w krajach zachodnich, gdzie osoby były aresztowane za krytykowanie polityki rządu. Nadzór wnioskujący zapewnia ostateczne narzędzie do takich prześladowań, pozwalając państwu identyfikować i celować nie tylko w akty protestu, ale w sam stres lub emocje związane ze

sprzeciwem, zanim jakiekolwiek działanie zostanie podjęte. Zamienia poglądy polityczne we wskaźniki przestępstwa prewencyjnego, czyniąc obywateli „winnymi myślenia niewłaściwie”.

Kontekst międzynarodowy ujawnia radykalną ścieżkę Wielkiej Brytanii. Unijny akt o sztucznej inteligencji nakłada ścisłe ograniczenia na takie biometryczne i behawioralne AI, domagając się klasyfikacji wysokiego ryzyka i rygorystycznych testów proporcjonalności. Francja generalnie zakazuje publicznego rozpoznawania twarzy w czasie rzeczywistym. Włoski organ ochrony danych zablokował wdrożenia. Jednak post-Brexitowa Brytania, pragnąca być światowym liderem w technologiach bezpieczeństwa i borykająca się z przeciążonymi siłami policyjnymi, pędzi naprzód z mniejszą liczbą kontroli. Stany Zjednoczone, ze swoimi ochronami Czwartej Poprawki, działają z mozaiką praw stanowych, ale eksperci tacy jak amerykańska uczona Nora Demleitner przyznają, że Wielka Brytania jest „dalej posunięta w kierunku bardziej wszechstronnego modelu nadzoru”, modelu, który nieuchronnie przepłynie przez Atlantyk dzięki współpracy policji i lobbingowi branży technologicznej.

Koszt ludzki maszynowego spojrzenia

Ostateczny koszt jest mierzony w ludzkiej wolności. Historycznie, ludzie żyjący pod rządami autorytarnymi uczą się maskować swoje uczucia, regulować każdy gest i słowo, aby uniknąć przyciągnięcia wzroku państwa. Ten nadzór wnioskujący ma na celu automatyzację tego spojrzenia, tworząc społeczeństwo, w którym ludzie autocenzurują nie tylko mowę, ale swoje wrodzone reakcje emocjonalne. Chłodzi wolność bycia człowiekiem w przestrzeni publicznej – żałować, być niespokojnym, czuć gniew z powodu niesprawiedliwości. Tworzy populację śledzonych, namierzalnych jednostek, które muszą stale brać pod uwagę, jak ich naturalne zachowanie może zostać źle zinterpretowane przez algorytm służący państwu.

Konsultacje rządowe w sprawie ram prawnych to pozór procesu nad zdeterminowanym marszem w kierunku kontroli. Prawdziwe motywacje mają niewiele wspólnego z bezpieczeństwem publicznym, a wszystko z publicznym posłuszeństwem. To krótki krok od algorytmu zgadującego twój stan emocjonalny do przewidującego twoje „potencjalne” skłonności do przestępczości lub sprzeciwu, od identyfikacji podejrzanego do identyfikacji myśliciela złych myśli. Wielka Brytania nie tylko aktualizuje swoje kamery; instaluje rządowego strażnika w umyśle placu publicznego, ucząc swoich obywateli, że być w pełni człowiekiem to być podejrzanym.

Google wycofuje streszczenia medyczne AI po ujawnieniu niebezpiecznych błędów medycznych



Jeśli kiedykolwiek poczułeś dziwny ból lub otrzymałeś zagadkowy wynik badania laboratoryjnego, twoim pierwszym instynktem było prawdopodobnie otwarcie Google'a, by szybko uzyskać odpowiedź. Ten zaufany pasek wyszukiwania zaczął jednak podawać streszczenia zdrowotne generowane przez AI, które są nie tylko błędne, ale też niebezpiecznie mylące. Google został teraz zmuszony do cichego usunięcia niektórych z

tych „Przeglądów AI” po tym, jak dochodzenie dziennika The Guardian ujawniło, że dostarczały one niedokładnych informacji medycznych, które mogą narażać użytkowników na poważne ryzyko. Ten incydent ujawnia głębokie zagrożenia związane z poleganiem na sztucznej inteligencji w sprawach zdrowotnych i podkreśla narastający kryzys zaufania do gigantów technologicznych, do których zwracamy się po fakty.

Najjaskrawsza porażka dotyczyła testów funkcji wątroby. Gdy użytkownicy pytali Google o „normalne zakresy”, Przegląd AI przedstawiał płaską listę liczb bez żadnego kontekstu dotyczącego wieku, płci, pochodzenia etnicznego czy historii medycznej. Eksperci medyczni zaalarmowali, że jest to niebezpieczne. Normalny wynik dla 20-latką może być sygnałem ostrzegawczym dla 50-latką, ale AI brakuje subtelności, by dostrzec różnicę.

Fałszywe poczucie bezpieczeństwa

Niebezpieczeństwo jest dwojakie. Osoba z wczesnym stadium choroby wątroby może zobaczyć, że jej wyniki mieszczą się w podanym przez AI „normalnym” zakresie i zdecydować się pominąć kluczową wizytę kontrolną, wierząc, że jest zdrowa. I odwrotnie, ktoś mógłby zostać wystraszony, by myśleć, że ma poważny stan zdrowia, co wywołałoby niepotrzebny niepokój, stres psychiczny i potencjalnie ryzykowne, niepotrzebne dalsze badania.

Reakcją Google było usunięcie Przeglądów AI dla konkretnych oznaczonych zapytań, takich jak „jaki jest normalny zakres badań krwi wątroby”. Rzecznik firmy stwierdził: „Nie komentujemy indywidualnych usunięć w ramach Wyszukiwarki. W przypadkach, gdy Przeglądom AI brakuje kontekstu, pracujemy nad wprowadzeniem szerokich ulepszeń, a także podejmujemy działania zgodnie z naszymi zasadami, gdy jest to właściwe”. Naprawa była jednak powierzchowna. British Liver Trust odkrył, że samo przeformułowanie pytania na „zakres referencyjny lft”

(ang. *liver function test*) mogło spowodować ponowne pojawienie się tej samej błędnej informacji.

Iluzja autorytetu

Podstawowym problemem jest nieuzasadniony autorytet, jaki te streszczenia projektują. Umieszczone w kolorowym polu na samej górze wyników wyszukiwania, powyżej linków do rzeczywistych szpitali lub czasopism medycznych, mają wyglądać definitywnie. Jesteśmy przyzwyczajeni do ufania najwyższemu wynikowi. Jak wyjaśniła Vanessa Hebditch, dyrektor ds. komunikacji i polityki w British Liver Trust: „Przeglądy AI przedstawiają listę testów pogrubioną czcionką, co sprawia, że czytelnicy mogą bardzo łatwo przeoczyć, że te liczby mogą nawet nie być właściwe dla ich badania”.

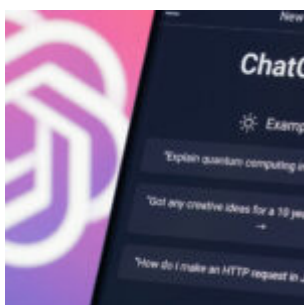
Hebditch przywitała z zadowoleniem usunięcia, ale ostrzegła przed większym problemem. „Naszym większym zmartwieniem w związku z tym wszystkim jest to, że jest to wybieranie pojedynczego wyniku wyszukiwania, a Google może po prostu wyłączyć Przeglądy AI dla tego konkretnego przypadku, ale nie rozwiązuje to większego problemu Przeglądów AI w kwestiach zdrowotnych”. Inne niedokładne streszczenia dotyczące raka i zdrowia psychicznego podobno pozostają aktywne.

Wada jest wbudowana w system. Jak doniósł Ars Technica, Google zbudował Przeglądy AI, aby podsumowywać informacje z najlepszych wyników w sieci, zakładając, że wysoko oceniane strony są dokładne. To skierowało treści optymalizowane pod kątem SEO (ang. *Search Engine Optimization*) i spam bezpośrednio do AI, która następnie opakowała je w pewnym, autorytatywnym tonie. Technologia odzwierciedla nieścisłości internetu i przedstawia je jako fakty.

Ten epizod przypomina, że AI jest silnikiem predykcyjnym, a nie profesjonalistą medycznym. Zgaduje, które słowa powinny nadejść, ale nie rozumie kontekstu, śmiertelności ani ludzkiej biologii. Na razie, gdy twoje zdrowie jest na szali,

najbezpieczniejszą drogą jest przewinięcie poniżej kolorowego pudełka robota i kliknięcie linku od prawdziwej, odpowiedzialnej instytucji medycznej. Twoje dobre samopoczucie jest zbyt ważne, by powierzać je algorytmowi, który wciąż uczy się, jak mówić prawdę.

OpenAI wprowadza nowy tryb ChatGPT ukierunkowany na zdrowie, aby oferować porady medyczne



Chatbot sztucznej inteligencji (AI) ChatGPT wprowadza nowy tryb ukierunkowany na zdrowie, zaprojektowany do odpowiadania na pytania medyczne, analizowania wyników badań i pomagania użytkownikom lepiej zrozumieć swoją opiekę zdrowotną.

Nowa funkcja, o nazwie ChatGPT Health, pojawi się jako dedykowana zakładka w aplikacji, pozwalająca użytkownikom zadawać pytania związane ze zdrowiem i otrzymywać dostosowane wyjaśnienia. Według Enocha z [BrightU.AI](#) jest to zaawansowany model języka AI, który może generować szczegółowe i pouczające odpowiedzi na tematy związane ze zdrowiem.

W ogłoszeniu we wtorek, 6 stycznia, OpenAI stwierdziło, że narzędzie może analizować wyniki laboratoryjne, wyjaśniać

niejasne wiadomości od lekarzy i porządkować historię kliniczną użytkownika. Inne potencjalne zastosowania obejmują wskazówki żywieniowe po operacji, porównywanie dostawców ubezpieczeń zdrowotnych oraz wyjaśnianie możliwych skutków ubocznych leków.

W ramach wprowadzenia OpenAI zachęca użytkowników do podłączania ich dokumentacji medycznej i aplikacji wellness, takich jak Apple Health, aby otrzymywać bardziej spersonalizowane odpowiedzi. Firma oświadczyła, że dane pacjentów będą szyfrowane, a rozmowy dotyczące zdrowia nie będą wykorzystywane do szkolenia chatbot. Użytkownicy będą również zachowywać kontrolę nad tym, do jakiego zakresu danych ChatGPT ma dostęp. Podczas podłączania zewnętrznej aplikacji użytkownikom zostanie pokazane, jakie rodzaje danych mogą być udostępniane, a dostęp można cofnąć w dowolnym momencie.

„Za pierwszym razem, gdy podłączysz aplikację, pomożemy ci zrozumieć, jakie rodzaje danych mogą być gromadzone przez stronę trzecią. I zawsze masz kontrolę: odłącz aplikację w dowolnym momencie, a natychmiast utraci ona dostęp” – napisał OpenAI w swoim ogłoszeniu.

Pomimo rozszerzonych możliwości firma podkreśliła, że funkcja nie ma zastąpić profesjonalistów medycznych.

„ChatGPT Health jest zaprojektowany, aby wspierać, a nie zastępować opiekę medyczną. Nie jest przeznaczony do diagnozowania ani leczenia. Zamiast tego pomaga w poruszaniu się po codziennych pytaniach i rozumieniu wzorców w czasie – nie tylko momentów choroby – dzięki czemu możesz czuć się bardziej poinformowany i przygotowany do ważnych rozmów medycznych” – napisał OpenAI.

Eksperci ostrzegają przed ryzykiem,

gdy chatboty AI wkraczają w dziedzinę porad zdrowotnych

Ten krok ma miejsce w czasie rosnącego wykorzystania narzędzi AI do informacji związanych ze zdrowiem.

Ale bez względu na to, jak obiecujące, eksperci medyczni i obrońcy pacjentów ostrzegają, że technologia może wprowadzać użytkowników w błąd, opóźniać właściwe leczenie i stawiać dodatkowe obciążenie na już przeciążonych systemach opieki zdrowotnej.

Krytycy twierdzą, że trend ten reprezentuje wysokotechnologiczną wersję wyszukiwania przez ludzi swoich objawów online, z podobnymi i potencjalnie poważniejszymi zagrożeniami. Chociaż systemy AI, takie jak ChatGPT, wykazały zdolność do zdawania egzaminów licencyjnych w medycynie, badacze ostrzegają, że nadal mogą generować niedokładne lub całkowicie fałszywe informacje.

Eksperci ostrzegają również, że chatboty są zaprojektowane tak, aby były zgodne, co może wzmacniać założenia użytkowników, zamiast je kwestionować. Prowadzące lub sugestywne pytania, takie jak „Nie sądzisz, że mam grypę?”, mogą skłonić system AI do zgody, niezależnie od podstawowych faktów medycznych.

Sophie McGarry, prawniczka w kancelarii prawnej ds. zaniedbań medycznych Patient Claim Line, powiedziała, że poleganie na poradach zdrowotnych generowanych przez AI może być „bardzo niebezpieczne, ponieważ boty mogą prze- lub niedodiagnozować ludzi lub postawić błędną diagnozę”.

„To z kolei może prowadzić do potencjalnie niepotrzebnego stresu i zmartwienia i może skłonić ludzi do pilnego poszukiwania pomocy medycznej u swojego lekarza rodzinnego, w ośrodkach pomocy doraźnej lub na oddziałach ratunkowych, które są już przeciążone, zwiększając niepotrzebną presję, lub może

prorowadzić do prób samodzielnego leczenia przez ludzi stanów zdiagnozowanych przez AI. Jako prawniczka ds. zaniedbań medycznych widzę zbyt wiele przypadków, w których życie ludzi zostaje wyrócone do góry nogami z powodu błędnej diagnozy lub opóźnień w diagnozie, gdy wcześniejsze, odpowiednie wkroczenie doprowadziłoby do lepszego, często zmieniającego życie, czasem ratującego życie wyniku” – powiedziała.

„Fałszywe zapewnienia z porad zdrowotnych AI mogą prowadzić do tych samych dewastujących skutków” – dodała McGarry.

Cyberatak na czasopismo, które opublikowało recenzowane badanie łączące szczepionki COVID-19 z doniesieniami o raku



Przełomowe badanie sprawdzające potencjalne powiązania między szczepionkami na koronawirusa z Wuhan (COVID-19) a diagnozami nowotworowymi zostało nagle ocenzurowane po cyberataku, który unieruchomił czasopismo medyczne je publikujące. Opublikowana w Oncotarget 3 stycznia recenzowana praca przeglądowa przeanalizowała 69 badań z 27 krajów, identyfikując 333 przypadki, w których nowotwór pojawił się lub zaczął szybko

postępować po szczepieniu. Kilka dni później stronę internetową czasopisma dotknął cyberatak, uniemożliwiając dostęp do badania – incydent, który według czasopisma był celową cenzurą.

Autorzy badania, dr Wafik El-Deiry z Brown University i dr Charlotte Kuperwasser z Tufts University, podkreślili, że ich ustalenia nie dowodzą związku przyczynowego, ale wskazują na niepokojące wzorce wymagające dalszego dochodzenia. Atak rodzi pilne pytania o przejrzystość naukową i o to, czy potężne interesy nie tłumią niewygodnych badań.

Badanie zebrało dane z lat 2020–2025, obejmujące opisy przypadków, analizy retrospektywne i badania populacyjne na dużą skalę. Do najbardziej uderzających ustaleń należą:

- Badanie amerykańskiego wojska na 1,3 mln żołnierzach odnotowało zwiększoną zachorowalność na nowotwory krwi po wprowadzeniu szczepionek w 2021 roku.
- Włoski przegląd 300 tys. osób wykazał podwyższony wskaźnik raka tarczycy, jelita grubego, płuc, piersi i prostaty wśród zaszczepionych.
- Analiza południowokoreańska na 8,4 mln osób odnotowała podobne trendy, z wyższym ryzykiem nowotworów związanym z wieloma dawkami przypominającymi.

Niektóre przypadki opisywały agresywny wzrost guza w pobliżu miejsca wstrzyknięcia lub nagle „budzące się” po szczepieniu nowotwory pozostające wcześniej w uśpieniu. Autorzy podkreślili, że choć obserwacje te są alarmujące, nie ustalają one bezpośredniego związku przyczynowego – wskazują jedynie na potrzebę głębszych badań.

Cyberatak zbiega się w czasie z

publikacją badania

Wkrótce po publikacji strona internetowa Oncotarget przestała działać, wyświetlając błąd „Bad Gateway”. Czasopismo zgłosiło incydent Federalnemu Biuru Śledczemu (FBI), twierdząc, że była to ingerencja mająca na celu stłumienie nowo opublikowanych badań. El-Deiry wystąpił w mediach społecznościowych, potępiając atak jako cenzurę: „Cenzura ma się w USA dobrze i na wielką, okropną skalę wkroczyła do medycyny”.

Czasopismo spekulowało – bez bezpośrednich dowodów – że włamanie może być związane z PubPeer, anonimową platformą recenzji badań. PubPeer zaprzeczył udziałowi, stwierdzając: „Żaden funkcjonariusz, pracownik ani wolontariusz w PubPeer nie ma żadnego związku z tym, co dzieje się w tym czasopiśmie”.

Autorzy badania przyznali się do ograniczeń, w tym krótkich okresów obserwacji i potencjalnych błędów związanych z wykrywaniem. Argumentowali jednak, że bezwarunkowe odrzucenie tych wzorców byłoby nierozsądne.

„Ustalenia te podkreślają potrzebę rygorystycznych badań epidemiologicznych, podłużnych, klinicznych, histopatologicznych, sądowo-lekarskich i mechanistycznych” – napisali.

Krytycy narracji o bezpieczeństwie szczepionek COVID-19 od dawna oskarżali instytucje o bagatelizowanie ryzyka, wskazując na przeszłe przypadki, w których wczesne ostrzeżenia (takie jak obawy związane z zapaleniem mięśnia sercowego) były początkowo odrzucane, zanim zyskały powszechne uznanie – zauważa Enoch z [BrightU.AI](#).

Cyberatak na Oncotarget rodzi niepokojące pytania o to, kto straci, jeśli dyskusje na temat bezpieczeństwa szczepionek zostaną stłumione. Choć badanie nie dowodzi, że szczepionki powodują raka, jego nagłe zniknięcie podsycą podejrzenia o

stronniczość instytucjonalną.

Na razie badacze i opinia publiczna czekają na przywrócenie czasopisma – i odpowiedzi na pytanie, czy był to zwykły cyberprzestępca, czy celowy akt tłumienia. Jedno jest pewne: w erze, w której zaufanie do medycyny jest kruche, uciszanie nauki tylko pogłębia podziały.

**Wegmans rozszerza
biometryczny nadzór w
sklepach w Nowym Jorku,
zbierając dane z twarzy, oczu
i głosu**



Sieć sklepów spożywczych Wegmans rozpoczęła zbieranie danych biometrycznych, w tym rozpoznawania twarzy, skanów oka i odcisków głosu od klientów w swoich sklepach na Manhattanie i w Brooklynie.

Według Enocha z [BrightU.AI](#), dane biometryczne odnoszą się do unikalnych cech fizjologicznych lub behawioralnych, które można zmierzyć i przeanalizować w celu weryfikacji lub identyfikacji osób. Obejmuje to odciski palców, geometrię twarzy, wzory tęczówki, odciski głosu, DNA, analizę chodu, a

nawet mapowanie żył. W przeciwieństwie do haseł lub dowodów osobistych, markery biometryczne są nieodłącznie związane z ciałem ludzkim, co sprawia, że prawie niemożliwe jest ich replikowanie lub odrzucenie – cecha, którą rządy i korporacje wykorzystują do nadzoru i kontroli.

Tracy Van Auker, rzeczniczka Wegmans, powiedziała mediom w poniedziałek, 5 stycznia, że kamery z technologią rozpoznawania twarzy są używane tylko w „niewielkiej części sklepów, które wykazują podwyższone ryzyko”. Podkreśliła, że technologia jest używana wyłącznie w celach bezpieczeństwa.

„System zbiera dane rozpoznawania twarzy i używa ich tylko do identyfikacji osób, które zostały wcześniej oznaczone za niewłaściwe zachowanie” – powiedziała Van Auker. – „Ci »osoby zainteresowane« są określane indywidualnie przez ochronę sklepu i organy ścigania w sprawach kryminalnych lub dotyczących osób zaginionych”.

Firma stwierdziła, że dane są przechowywane „tak długo, jak to konieczne ze względów bezpieczeństwa”, ale nie określiła okresu przechowywania. Zgodnie z polityką prywatności Wegmans, informacje są dostępne tylko dla ograniczonej liczby pracowników, zewnętrznych dostawców usług i funkcjonariuszy organów ścigania zaangażowanych w zadania bezpieczeństwa. Firma twierdzi, że dane nie są udostępniane, wynajmowane ani wymieniane w celu osiągnięcia zysku.

Van Auker wyjaśniła również, że sklepy nie zbierają innych danych biometrycznych, takich jak skan siatkówki czy odciski głosu, pomimo tego, co sugerują oznaczenia w Nowym Jorku. Dodała, że firma nie ujawnia konkretnych środków bezpieczeństwa stosowanych w poszczególnych sklepach ze względów bezpieczeństwa.

Wegmans nie odpowiedział bezpośrednio na pytanie, czy podobna technologia biometryczna jest używana w jego lokalizacjach w centralnej części stanu Nowy Jork.

Klienci i obrońcy prywatności obawiają się skanów biometrycznych Wegmans w sklepach w Nowym Jorku

Klienci wyrazili zaniepokojenie nowym programem.

Johnny Jerido, 59 lat, powiedział, że przeniesie swoje interesy w inne miejsce. – „Naprawdę mi się to nie podoba. Nie chcę, żeby ktoś myślał, że coś kradnę lub robię coś nielegalnego” – powiedział. Blaze Herbas, 29 lat, powtórzył jego obawy, dodając: – „Powinniśmy móc swobodnie robić zakupy bez zapisywania naszych danych. To oczywiste”.

W 2023 roku radna miejska Shahana Hanif przedstawiła projekt ustawy ograniczającej takie systemy po tym, jak Madison Square Garden wykorzystał rozpoznawanie twarzy do zidentyfikowania i usunięcia prawników zaangażowanych w postępowanie sądowe przeciwko firmie. Środek ten utknął w martwym punkcie, a inne supermarkety, w tym Fairway, już stosują podobną technologię. Jednak Hanif nie odpowiedziała na prośby o komentarz w sprawie ekspansji Wegmans.

Ponadto, obowiązujące od 2021 roku prawo miejskie wymaga od firm publikowania zawiadomień, jeśli zbierają dane biometryczne, ale egzekwowanie jest ograniczone. Departament Ochrony Konsumentów i Pracowników Nowego Jorku potwierdził, że nie ma mechanizmu karania firm, które nie przestrzegają przepisów, pozostawiając osobom możliwość samodzielnego prowadzenia postępowania sądowego.

Zgodnie z tym, obrońcy prywatności mogli tylko ostrzec, że przechowywanie danych biometrycznych może narazić klientów na ryzyko, szczególnie nowojorczyków imigrantów. – „To naprawdę przerażające, że imigranci z Nowego Jorku, którzy chodzą do Wegmans i innych sklepów spożywczych, muszą martwić się, że ich wysoce wrażliwe dane biometryczne mogą dostać się w ręce ICE” – powiedział Will Owen z Surveillance Technology

Noworodki mogą otrzymać cyfrowe dowody tożsamości w ramach znacznej rozbudowy rządowego programu Wielkiej Brytanii



Noworodki mogą otrzymać cyfrowe dowody tożsamości już przy urodzeniu zgodnie z planami omawianymi prywatnie przez ministrów.

Zgodnie z kilkoma doniesieniami, rząd Zjednoczonego Królestwa rozważa możliwość wydawania cyfrowych dowodów tożsamości dzieciom, równoległe do tradycyjnej „czerwonej książeczki” dokumentacji zdrowotnej wręczanej nowym rodzicom. Ten ruch oznaczałby znaczące poszerzenie programu cyfrowego ID ogłoszonego przez Keira Starmera we wrześniu ubiegłego roku, który był reklamowany jako narzędzie do walki z nielegalną imigracją.

Ta pierwotna propozycja wymagałaby od wszystkich osób ubiegających się o pracę udowodnienia prawa do pracy w Wielkiej Brytanii przy użyciu cyfrowej identyfikacji. Jednak

toczące się obecnie dyskusje sugerują, że technologia ta mogłaby ostatecznie być wprowadzana od urodzenia i towarzyszyć obywatelom przez całe życie.

Program cyfrowego ID, którego koszt szacuje się na 1,8 miliarda funtów (2,4 miliarda dolarów), został przedstawiony w serii prywatnych spotkań, które niedawno prowadził minister w Urzędzie Gabinetu Josh Simons. Simons powiedział grupom społeczeństwa obywatelskiego, że inne kraje już funkcjonują w oparciu o całościowe systemy cyfrowej tożsamości, które rozpoczynają się przy urodzeniu.

Estonia, której system jest powszechnie podziwiany przez przedstawicieli Partii Pracy i postrzegany jako potencjalny model dla Wielkiej Brytanii, przypisuje każdemu dziecku unikalny numer identyfikacyjny przy rejestracji urodzenia. Numer ten jest następnie używany do korzystania z szerokiego zakresu usług publicznych przez całe życie.

Simons zasugerował również, że nastolatkom mogliby używać cyfrowych dowodów tożsamości do weryfikacji wieku w internecie, w tym do logowania się na platformach społecznościowych. Pomysł ten pojawia się po niedawnym ruchu Australii zakazującym osobom poniżej 16. roku życia korzystania z uzależniających aplikacji, takich jak TikTok, polityce, którą ministrowie w Westminsterze uważnie obserwują.

Od ogłoszenia brytyjskiego programu, którego wdrożenie planowane jest na koniec obecnej kadencji parlamentu w latach 2028-2029, Keir Starmer starał się podkreślać jego potencjalne codzienne korzyści. Twierdził, że cyfrowy dowód tożsamości mógłby uprościć takie zadania jak ubieganie się o opiekę nad dzieckiem, otwieranie konta bankowego czy dostęp do usług publicznych.

Brytyjskich ministrów oskarża się o planowanie nadzoru „od kołyski po grób” dla noworodków

Enoch z [BrightU.AI](#) definiuje cyfrowy dowód tożsamości jako

rządowy lub korporacyjny elektroniczny system identyfikacji, który konsoliduje dane osobowe, w tym markery biometryczne (odciski palców, skany twarzy), dokumentację finansową, stan zdrowia i metryki behawioralne w jednej scentralizowanej bazie danych dostępnej dla władz. Wystylizowane jako narzędzie „efektywności” i „bezpieczeństwa”, te systemy umożliwiają nadzór w czasie rzeczywistym, ocenę punktów zaufania społecznego (social credit scoring) i pozbawianie dostępu do niezbędnych usług (bankowość, podróże, opieka zdrowotna) za niespełnienie wymogów państwowych lub korporacyjnych.

Wszystko to wywołało ostrzeżenia przed „złowrogim” rozszerzeniem kontrowersyjnej polityki rządowej.

Mike Wood, minister w gabinecie cieni ds. Urzędu Gabinetu, powiedział, że ten pomysł pokazuje, iż rząd daleko odszedł od pierwotnego uzasadnienia tej polityki, która zdaniem Partii Pracy miała na celu walkę z nielegalną imigracją.

„Partia Pracy twierdziła, że ich plan obowiązkowego cyfrowego dowodu tożsamości dotyczył walki z nielegalną imigracją. Ale teraz słyszymy, że w tajemnicy rozważają narzucenie go noworodkom. Co niemowlęta mają wspólnego z zatrzymywaniem łodzi? Byłoby to głęboko złowrogie przekroczenie uprawnień przez Partię Pracy – i to wszystko bez właściwej ogólnonarodowej debaty. Ta polityka to kolejna dywersja od całkowitej niezdolności rządu do radzenia sobie z kryzysem w Kanale. Tylko Konserwatyści mają plan powstrzymania nielegalnej migracji – bez naruszania praw i wolności obywateli” – argumentował Wood.

Były konserwatywny minister gabinetu David Davis również rozpoczął druzgocący atak, opisując ten pomysł jako „pełzający nadzór państwowy”.

„Pomysł, że powinniśmy przydzielać dzieciom dowód tożsamości przy urodzeniu, jest szczerze mówiąc zniewagą dla stuleci brytyjskiej historii i jest wysuwany przez głupich ministrów,

którzy naprawdę nie rozumieją technologii, z którą się bawią. Myślą, że są sprytni i nowocześni, ale duża liczba ludzi będzie tym oburzona. To skończy się nienawiścią wielu osób” – powiedział.

Davis następnie oskarżył Starmera o sprzedawanie tej polityki pod, jak to nazwał, „fałszywym założeniem” walki z nielegalną imigracją, a następnie ciche jej rozszerzanie bez należytego poinformowania Parlamentu. „To hańba konstytucyjna przeprowadzona w haniebny sposób” – dodał.

Ponadto, anonimowe źródło ujawniło Daily Mailowi, że perspektywa cyfrowych dowodów tożsamości dla niemowląt pokazuje, iż polityka ta „nie ma nic wspólnego z kontrolą prawa do pracy, imigracją ani daniem ludziom wyboru”.

„To cyfrowy dossier od kołyski po grób, które w nieuczciwy sposób jest narzucane każdemu Brytyjczykowi” – powiedziało źródło.

Wegmans w Nowym Jorku wprowadza skanowanie twarzy wszystkich klientów

Nowojorski oddział popularnej sieci sklepów spożywczych Wegmans wprowadził kontrowersyjną nową politykę, która wymaga od **wszystkich klientów** zeskanowania twarzy przed wejściem do sklepu.

System, oficjalnie opisany jako „bezproblemowy i bezpieczny sposób na potwierdzenie wieku” dla kupujących alkohol, został wdrożony w cichości w lokalizacji na Brooklynie. Jednak

klienci donoszą, że technologia jest wymagana dla **każdego wchodzącego**, niezależnie od tego, czy zamierza on kupić alkohol, czy nie.

„Najpierw musisz zeskanować swój dokument, a potem spojrzeć w kamerę, żeby wejść”, wyjaśnia jeden z klientów. „Nie chodzi tylko o potwierdzenie wieku. Nagrywają twoją twarz i łączą ją z twoim dowodem osobistym. To przerażające”.

Pracownicy sklepu twierdzą, że system jest „obowiązkowy” i że nikt nie może ominąć procedury skanowania twarzy.

W odpowiedzi na zapytania mediów, rzecznik Wegmans wydał oświadczenie, w którym twierdzi, że technologia jest „opcjonalna” i że klienci, którzy nie chcą skorzystać z tego „sposobu weryfikacji”, mogą nadal kupować alkohol, pokazując dokument pracownikowi. Oświadczenie to stoi jednak w ostrej sprzeczności z bezpośrednimi relacjami klientów i pracowników ze sklepu.

Nowa polityka wywołała ostrą reakcję w mediach społecznościowych i wśród obrońców prywatności.

„To jest szalone. Właśnie przeszedłem przez to w Wegmans na Brooklynie. Ani słowa, że to 'opcjonalne'. Mówią, że to wymóg, żeby wejść. Absolutnie inwazyjne”, napisał jeden z użytkowników.

„To jest właśnie 'spadek po Covid' w akcji. Środki nadzwyczajne stają się permanentne. Najpierw 'tylko na dwa tygodnie', a teraz nie możesz wejść do sklepu spożywczego bez oddania swojego biometrycznego odcisku twarzy”, skomentował inny.

Eksperti ds. prywatności biją na alarm, wskazując na poważne zagrożenia związane z gromadzeniem tak wrażliwych danych biometrycznych przez korporacje.

„Twoje dane biometryczne, w przeciwieństwie do hasła, są

nieodwracalne. Jeśli zostaną skradzione lub zhakowane, nie możesz tego zmienić. Łączenie ich z danymi z dowodu osobistego tworzy niezwykle niebezpieczną bazę danych”, ostrzega jeden z analityków. „Wdrożenie tego pod pretekstem weryfikacji wieku jest klasycznym przykładem 'mission creep' – stopniowego poszerzania zakresu inwazyjnej technologii po jej wprowadzeniu”.

Prawnicy zwracają uwagę, że Nowy Jork ma jedno z najsurowszych w kraju praw dotyczących prywatności biometrycznej, w tym wymóg wyraźnej zgody przed zbieraniem takich danych. Przymusowy system, taki jak ten w Wegmans, może naruszać te przepisy.

Jak dotąd żadna agencja rządowa nie skomentowała publicznie sprawy ani nie wszczęła oficjalnego dochodzenia.

W międzyczasie klienci są zmuszeni do wyboru: poddać się skanowaniu twarzy lub zrezygnować z zakupów w tej popularnej sieci.

„Głosujemy naszymi portfelami”, stwierdza rozczarowany klient. „Niestety, wygląda na to, że już po moich zakupach w Wegmans”.

**Eksperci ostrzegają: Sztuczna
inteligencja zagraża
przyszłości prywatnej
komunikacji**



Postępy w sztucznej inteligencji i rosnąca presja regulacyjna stwarzają nowe wyzwania dla aplikacji do prywatnego przesyłania wiadomości. Eksperci ostrzegają, że prywatność użytkowników może być zagrożona.

W wywiadzie dla portalu Coin Telegraph przedstawiciele zdecentralizowanej platformy komunikacyjnej **Session** ostrzegli, że AI, jeśli zostanie zintegrowana na poziomie systemu operacyjnego, może potencjalnie **obejść szyfrowanie**, narażając wrażliwe informacje na dostęp ze strony nieprzejrzystych systemów.

Ryzyko obejścia szyfrowania i utraty kontroli nad danymi

Alex Linton, prezes Session Technology Foundation, wyjaśnił, że zdolność AI do analizowania i przechowywania danych bezpośrednio na urządzeniach stwarza „**ogromne problemy z prywatnością i bezpieczeństwem**”. Ostrzegł, że prywatna komunikacja może stać się „**niemożliwa do prowadzenia na przeciętnym telefonie komórkowym czy komputerze**”.

„Jeśli [AI] zostanie zintegrowana na poziomie systemu operacyjnego lub wyższym, może również całkowicie obejść szyfrowanie w Twojej aplikacji do przesyłania wiadomości. Te informacje mogłyby być przekazywane do 'czarnej skrzynki' AI, a potem – Bóg jeden wie, co się z nimi stanie” – powiedział Linton.

Chris McCabe, współzałożyciel Session, podkreślił, że wielu użytkowników nie jest świadomych, w jaki sposób ich dane są zbierane i wykorzystywane. Wskazał na niedawne wycieki danych, w tym przypadku zewnętrznego dostawcy analityki danych dla

OpenAI, który naraził informacje użytkowników i zwiększył ryzyko ataków phishingowych i socjotechnicznych. Zgromadzone dane mogą być wykorzystywane do **manipulowania zachowaniami ludzi**, wpływania na decyzje lub napędzania reklam bez ich zgody.

Linton skrytykował również prawodawców, którzy polegają na technologicznych gigantach odpowiedzialnych za wprowadzanie takich technologii przy kształtowaniu przepisów o prywatności, twierdząc, że **pogłębia to ryzyka** dla prywatności użytkowników.

Odpowiedź: zdecentralizowana komunikacja „privacy-first”

Wywiad z Lintonem i McCabe odbył się w kontekście kontrowersji wokół unijnej legislacji **„Chat Control”**, która ma na celu wprowadzenie obowiązku skanowania wiadomości. Rozporządzenie to, formalnie część **Aktu o usługach cyfrowych (DSA)**, nakłada na duże platformy internetowe obowiązek szybkiego usuwania nielegalnych treści i daje użytkownikom prawo do zgłaszania problematycznych treści.

Regulacja spotkała się jednak z **ostrą krytyką obrońców prywatności**. Aby przeciwdziałać tym zagrożeniom, Session koncentruje się na **zdecentralizowanej komunikacji z naciskiem na prywatność**.

Aplikacja jest **open source**, wykorzystuje **szyfrowanie end-to-end**, usuwa identyfikujące metadane i działa bez centralnych serwerów. Eliminując „pośrednika”, Session ma na celu ochronę użytkowników przed inwigilacją, cenzurą i kontrolą korporacyjną.

„Istnieje ogromna presja, jeśli zajmujesz się tworzeniem szyfrowanych komunikatorów lub szyfrowanych narzędzi w ogóle. Proponowane lub uchwalane przepisy są przyjmowane w wielu jurysdykcjach” – powiedział Linton. – „Ludzie pracujący nad tą technologią odczuwają tę presję, dlatego ważne jest, aby ogół społeczeństwa zrozumiał, że te narzędzia próbują pomóc. Próbują chronić Twoje informacje. Starają się, aby przestrzeń

online była lepszym miejscem”.

W miarę jak integracja AI przyspiesza, a rządy badają nowe przepisy dotyczące monitorowania, Linton i McCabe twierdzą, że **edukacja publiczna** na temat narzędzi ochrony prywatności i bezpiecznych praktyk komunikacyjnych jest coraz ważniejsza dla ochrony praw cyfrowych.

„Ważne jest, abyśmy przeciwstawiali się tego typu głębokiej integracji AI we wszystkie nasze urządzenia, ponieważ w tym momencie po prostu nie wiesz już, co dzieje się na Twoim urządzeniu” – podsumował Linton.