

Iran i Izrael potajemnie uzgodniły brak ataków na siebie poprzez rosyjski kanał porozumienia



Mogło dojść do pewnych poufnych negocjacji i „wzajemnego zrozumienia” między Iranem a Izraelem, głęboko za kulisami, podczas gdy na ulicach Iranu rozgrywały się protesty, a prezydent Trump zaczynał grozić atakami na Teheran.

W momencie, gdy Trump wydaje się wycofać (przynajmniej na razie) z grożonej militaryjnej interwencji, The Washington Post opublikował w środę raport stwierdzający, że Izrael i Iran pozostawały w pośrednim kontakcie dyplomatycznym za pośrednictwem Rosji jako mediatora.

„Na kilka dni przed wybuchem protestów w Iranie pod koniec grudnia, izraelscy urzędnicy poinformowali przywódców Iranu za pośrednictwem Rosji, że nie przeprowadzą ataków na Iran, jeśli Izrael nie zostanie zaatakowany jako pierwszy” – pisze WaPo. „Iran odpowiedział poprzez rosyjski kanał, że również powstrzyma się od ataku wyprzedzającego, jak stwierdzili dyplomaci i urzędnicy regionalni znający tę wymianę informacji”.

Czy mogło to być spowodowane irańskimi rakietami, które spadły na Tel Awiw w czerwcu? Jeśli tak, wydaje się, że Islamska Republika wreszcie ustanowiła zdolność odstraszenia.

Chronologia tego, co i kiedy zostało zakomunikowane, pozostaje niejasna. Jednak ten kanał porozumienia został już ujawniony w doniesieniach medialnych na Bliskim Wschodzie, na przykład we wcześniejszym [raporcie](#):

Izrael i Iran niedawno wymieniły poufne, pośrednie wiadomości za pośrednictwem Rosji wśród zaostrzonych napięć regionalnych, zgodnie z nowym raportem Amwaj.media opublikowanym dzisiaj. Wymiany te opisano jako wysiłek mający na celu zapobieżenie dalszej eskalacji militarnej, a nie ustanowienie jakiejkolwiek formy zawieszenia broni czy ram dyplomatycznych.

Zgodnie z raportem, wiadomości zostały przekazane przez rosyjskiego prezydenta Władimira Putina, po tym jak Izrael starał się przesłać sygnał, że nie jest zainteresowany eskalacją konfliktu militarnego na tym etapie. Irańscy urzędnicy potwierdzili otrzymanie wiadomości, ale podkreślili, że ich odpowiedź nie niesie za sobą żadnego zobowiązania, koordynacji ani obowiązku ze strony Iranu. Polityczne źródło irańskie cytowane w raporcie stwierdziło wprost, że „nie ma żadnego zobowiązania, koordynacji ani porozumienia o zawieszeniu broni”. Źródło to podkreśliło, że kontakt ten nie powinien być interpretowany jako krok w kierunku szerszych porozumień między oboma krajami, które pozostają zaciekłymi przeciwnikami bez bezpośrednich powiązań dyplomatycznych.

Wymiany te były podobno ograniczone w zakresie i intencji. Nie zaoferowano żadnych gwarancji, nie omawiano harmonogramów i nie ustanowiono żadnych mechanizmów monitorowania ani egzekwowania. Jedno ze źródeł opisało komunikację jako „wzajemne ogłoszenie dla wspólnego przyjaciela o braku nowych ataków”, co oznacza, że celem było po prostu zarządzanie napięciami w określonym momencie, a nie wprowadzanie jakiegokolwiek trwałego porozumienia.

Wysokie irańskie źródło polityczne potwierdziło, że pośrednia

komunikacja z Izraelem rzeczywiście miała miejsce, identyfikując Rosję, a konkretnie Putina, jako pośrednika. Źródło to ponowiło, że nie ma „żadnego porozumienia o zawieszeniu broni” i że wiadomości sprowadzały się jedynie do równoległych powiadomień o intencjach, a nie do wspólnego zrozumienia czy umowy.

Raport wskazuje, że irańską stroną w tych wymianach zajmował się nie ministerstwo spraw zagranicznych, lecz Ali Larijani, sekretarz Wysokiej Rady Bezpieczeństwa Narodowego Iranu.

Możliwe, że to posłużyło jako ważne tło dla pozornej decyzji Trumpa, aby nie uderzać na Iran w tym momencie. Izrael jest zwykle krajem, który najgłośniej nawołuje do uderzenia na Iran, ale tym razem rząd Netanjahu był nieco wyciszony.

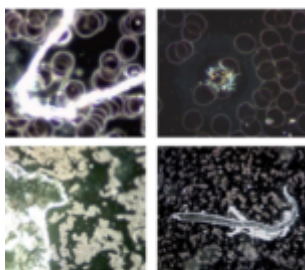
*Sorry to be a hater, but this is not "breaking news". It was reported over two weeks ago: <https://t.co/SjEJKtgAT6>
<https://t.co/jDExK0ZVND>*

– Mohammad Ali Shabani (@mashabani) [January 14, 2026](#)

Według wszelkich doniesień, ulice Iranu właściwie ucichły w tym momencie, po kulminacyjnej fali przemocy w tym tygodniu, która pozostawiła setki ofiar śmiertelnych, w tym wielu funkcjonariuszy policji i służb bezpieczeństwa.

Wstrząsające badanie ujawnia zanieczyszczenie DNA w

szczepionkach COVID-19, domaga się natychmiastowego wycofania z rynku



These images illustrate the variety of unusual phenomena and objects found in the blood of
1 with Coronavirus (SARS-CoV-2).

W cieniu skandalu związanego z COVID-19, narracja o „bezpiecznych i skutecznych” interwencjach genetycznych była nieustannie wpajana ufnej globalnej populacji. Jednak druzgoczące nowe odkrycie naukowe rozrywa tę oszukańczą narrację na strzępy, ujawniając celowe i niebezpieczne zanieczyszczenie w samym sercu programu szczepionek mRNA. Przełomowe badanie w formie preprintu potwierdziło obecność szczątkowego plazmidowego DNA bakterii w szczepionkach COVID-19 firm Pfizer i Moderna – zanieczyszczenia, które standardowe testy regulacyjne były zaprojektowane, by przeoczyć. Nie jest to drobna usterka jakościowa; jest to dowód katastrofalnej porażki w produkcji i nadzorze, narażającej miliony na ryzyko integracji genowej, utrzymującej się produkcji białka kolca oraz kaskady potencjalnych konsekwencji autoimmunologicznych i nowotworowych. Ustalenia sugerują, że same podmioty odpowiedzialne za ochronę zdrowia publicznego zorganizowały system testów, który z góry był skazany na patrzenie w inną stronę, zamieniając globalną kampanię szczepień w niekontrolowany eksperyment na ludzkim DNA.

Przynęta i podmianka: Zaplanowany

wada produkcyjna

Aby zrozumieć wagę tego zanieczyszczenia, trzeba najpierw zrozumieć „przynętę i podmiankę”, która miała miejsce w fabrykach szczepionek. Do wstępnych badań klinicznych Pfizer użył czystej, precyzyjnej metody generowania matrycy DNA dla swojego mRNA, znanej jako „Proces 1”. Ta metoda, choć droższa, dawała czystszy produkt. Jednak dla masowego rozpowszechnienia wśród społeczeństwa, firma po cichu przeszła na „Proces 2”. Ta tańsza, szybsza metoda opierała się na amplifikacji matryc DNA przy użyciu plazmidów bakteryjnych – proces podobny do używania kserokopiarki, która czasami dodaje przypadkowe, niezamierzone ślady na stronie.

Zespół naukowców, pod kierownictwem eksperta od genomiki Kevina McKernana, udowodnił, że ten proces pozostawił te bakteryjne plazmidowe „plany konstrukcyjne” w końcowych fiolkach. „Te hybrydy nie mogły zostać zdegradowane przez enzym, który producenci wybrali do oczyszczenia szczątkowego DNA jako ostatni etap procesu, i musieli o tym wiedzieć” – powiedziała współautorka Jessica Rose, Ph.D., w wywiadzie dla The Defender. Określiła tę decyzję jako „skandaliczną”, wskazując na ugruntowaną wiedzę naukową, że wybrany enzym jest nieskuteczny wobec powstałych hybryd RNA:DNA.

Nie chodzi tu jedynie o śladowe zanieczyszczenia; chodzi o bioaktywny obcy kod genetyczny wstrzykiwany bezpośrednio do ludzkiego krwiobiegu. Brian Hooker, Ph.D., dyrektor naukowy Children’s Health Defense, który częściowo sfinansował badania, ostrzegł, że to egzogenne DNA może „łatwiej rozprzestrzeniać się po ciele i nadal zarówno replikować, jak i ekspresjonować episomalnie, zamieniając ludzi w genetycznie modyfikowane fabryki produkujące białko kolca”. To może wyjaśniać utrzymującą się obecność białka kolca stwierdzoną u niektórych pacjentów nawet do dwóch lat po iniekcji – zjawisko nigdy nieujawnione jako ryzyko podczas nieustannej kampanii marketingowej „bezpieczne i skuteczne”.

Regulacyjna gra w trzy karty: Testowanie tego, co wiesz, że nie ma

Być może bardziej alarmujące niż samo zanieczyszczenie jest wymyślne przedstawienie regulacyjne, które pozwoliło mu pozostać niewykrytym. Regulatorzy ustalili limit bezpieczeństwa dla szczątkowego DNA w oparciu o przestarzały model zakładający, że będzie to DNA „nagie” i szybko ulegnie degradacji. DNA w tych zastrzykach jest jednak opakowane w ochronną powłokę nanocząstek lipidowych, co nadaje mu trwałość i dostęp do ludzkich komórek. Kwestia bezpieczeństwa przesuwana się z masy na liczbę fragmentów DNA – z których każdy jest potencjalnym biletem do genetycznego chaosu.

Badanie odsłania sedno oszustwa: test kontroli jakości wymagany przez regulatorów. Ten test, badanie qPCR, został zaprojektowany do poszukiwania tylko jednej konkretnej sekwencji DNA – genu oporności na kanamycynę. Ten gen, jak udowodnili badacze, jest wysoce wrażliwy na enzym oczyszczający i dlatego jest najmniej prawdopodobnym zanieczyszczeniem, które pozostanie. Jest to równoznaczne z tym, jakby inspektor zdrowia sprawdzał w restauracji tylko jeden, łatwy do wyczyszczenia przyrząd, ignorując kuchnię przepełnioną gnijącymi, toksycznymi odpadami.

„Testy, które zaprojektowali, zostały zaprojektowane tak, aby niczego nie znaleźć” – stwierdził wprost McKernan na swoim Substackie. Karl Jablonowski, starszy naukowiec w Children’s Health Defense, potwierdził to, zauważając: „Ci, którzy mieli czerpać zyski ze szczepionek, zaprojektowali test i testowali jakość. Wybrali test, który był najmniej prawdopodobny, aby dać zły wynik”. Kiedy badacze użyli szerszych metod testowania, odkryli poziom zanieczyszczenia DNA 15 do 48 razy wyższy niż już wadliwy limit FDA. System nie był zepsuty; został zbudowany, aby odnieść sukces w oszustwie.

Spuścizna milczenia i żądanie rozliczalności

To odkrycie łączy się z mroczniejszą historią zanieczyszczenia genetycznego, którą władze wolałyby pogrzebać. Badanie odnotowuje obecność sekwencji promotora SV40, znanej sekwencji onkogennej, którą wcześniej znaleziono w szczepionkach, ale nigdy nie ujawniono regulatorom, takim jak Health Canada. Ta sekwencja, pochodząca od wirusa małpiego związanego z nowotworami w badaniach laboratoryjnych, może ułatwiać wnikanie obcego DNA do jądra komórkowego. Jej obecność wraz z tym powszechnym zanieczyszczeniem plazmidowym DNA tworzy doskonałą burzę dla integracji genowej.

Badacze domagają się teraz odpowiedzi na niewygodne pytania. Dlaczego nie wymagano stosowania lepszych, wielocechowych metod testowania, już dostępnych i zwalidowanych? Dlaczego pozwolono na przejście z czystego Procesu 1 na podatny na zanieczyszczenia Proces 2 dla produktu przeznaczonego do powszechnego podawania? Rysują wyraźny kontrast: w przypadku testów PCR na COVID-19 regulatorzy domagali się weryfikacji wielocechowej, aby uniknąć wyników fałszywie negatywnych. Dla szczepionek wstrzykiwanych w miliardy ramion, w tym kobietom w ciąży i niemowlętom, uznano za wystarczający jeden, spreparowany cel.

To nie jest prosty błąd. To kulminacja dobrze zaplanowanej konstrukcji, która wykorzystała strach przed pandemią do wprowadzenia niedostatecznie przetestowanej technologii genetycznej, osłonięła jej architektów przed odpowiedzialnością, a następnie stworzyła regulacyjną zasłonę dymną, aby ukryć wynikające z tego niebezpieczeństwa. Szczepionki mRNA były przedstawiane jako linia ratunkowa, ale te dowody sugerują, że były koniem trojańskim, niosącym niechciany kod genetyczny do serca ludzkich komórek. Apel niezależnych naukowców jest jasny i pilny: te produkty muszą zostać wycofane z rynku, a pełne śledztwo kryminalne w sprawie

tego masowego oszustwa medycznego musi się rozpocząć.



W znaczącej eskalacji napięć dyplomatycznych Chiny wprowadziły natychmiastowe kontrole eksportu towarów podwójnego zastosowania przeznaczonych dla Japonii, co jest posunięciem wyraźnie związanym ze stanowiskiem Tokio w sprawie Tajwanu.

Zakaz, ogłoszony przez chińskie Ministerstwo Handlu, dotyczy przedmiotów, które mogłyby wzmocnić zdolności wojskowe Japonii i sygnalizuje gotowość Pekinu do wykorzystywania narzędzi gospodarczych w celu ukarania sprzeciwu politycznego.

Kontrole są bezpośrednią reakcją na komentarze japońskiej premier Sanae Takaichi z listopada. Takaichi stwierdziła, że chiński atak na Tajwan mógłby stanowić dla Japonii „sytuację zagrażającą przetrwaniu”, potencjalnie uzasadniając interwencję wojskową na mocy jej ustaw o zbiorowej samoobronie.

Pekin, który postrzega Tajwan jako nieodłączną część swojego terytorium, potępił te uwagi jako poważną prowokację i lekkomyślną ingerencję w wewnętrzne sprawy Chin.

Rzecznik chińskiego Ministerstwa Handlu ostrzegł przed „głęboko szkodliwymi konsekwencjami”, określając stanowisko Japonii jako niebezpieczne naruszenie fundamentalnej zasady jednych Chin.

Zakres i potencjalny wpływ zakazu

Silnik AI Enoch firmy [BrightU.AI](#) wyjaśnia, że te nowe ograniczenia dotyczą przedmiotów podwójnego zastosowania, konkretnie produktów, oprogramowania i technologii mających zastosowanie zarówno cywilne, jak i wojskowe. Chociaż konkretna lista nie została opublikowana, chińskie ramy kontroli podwójnego zastosowania obejmują szeroki zakres sektorów, w tym zaawansowaną elektronikę, chemikalia, komponenty lotnicze, czujniki i pierwiastki ziem rzadkich.

Zakaz zabrania eksportu nie tylko do japońskich użytkowników końcowych na cele wojskowe, ale także do każdego podmiotu uznanego za zdolny do wzmocnienia japońskiej siły militarnej.

Co istotne, Chiny rozszerzyły groźbę sankcji prawnych na osoby lub organizacje w krajach trzecich, które pomagają w omijaniu zakazu poprzez przekierowywanie ograniczonych towarów pochodzenia chińskiego do Japonii. Ten szeroki, celowo niejasny język daje Pekinowi maksymalną swobodę w zakłócaniu łańcuchów dostaw, zauważają analitycy, potencjalnie wpływając na przesyłki nawet w nominalnie celach cywilnych.

Wpływ może być znaczący. Pomimo wysiłków na rzecz dywersyfikacji, Japonia pozostaje w dużym stopniu uzależniona od importu z Chin w przypadku wielu kluczowych materiałów i komponentów, co czyni jej przemysł i sektor obronny podatnymi na takie kontrole.

Reakcja Japonii: Stanowczy protest

Japonia zareagowała szybko i stanowczo. Ministerstwo Spraw Zagranicznych wystosowało formalny protest do Chin, domagając się natychmiastowego wycofania środków.

Podczas spotkania z wysokim chińskim dyplomatą Masaaki Kanai, sekretarz generalny ds. azjatyckich i oceanicznych Japonii, potępił kontrole eksportu jako „absolutnie nie do przyjęcia” i „głęboko godne ubolewania”, argumentując, że nie są one zgodne

z normami międzynarodowymi.

Podczas gdy japońskie ministerstwa rządowe publicznie stwierdziły, że oceniają zakres ograniczeń, niektóre źródła prywatnie sugerowały, że początkowy ruch może być częściowo symboliczny.

Jednak jednoznaczne powiązanie z kwestią Tajwanu podkreśla poważny rozłam polityczny.

Szersza konfrontacja strategiczna

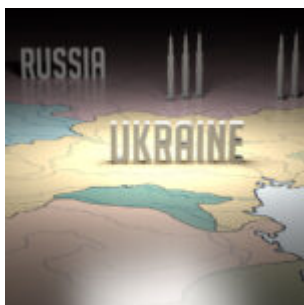
To działanie nie jest odosobnionym środkiem handlowym, lecz wybuchem długo tłących się napięć. Ma miejsce w czasie, gdy Japonia znacząco zwiększa swoje wydatki na obronność i pogłębia swoje bezpieczeństwo sojusznicze ze Stanami Zjednoczonymi, co jest trendem, który Pekin konsekwentnie krytykuje jako niebezpieczny krok w kierunku odrodzenia japońskiego militarizmu.

Krok ten przywołuje dawne chińskie taktyki, w szczególności wprowadzone w 2010 roku restrykcje na eksport metali ziem rzadkich podczas wcześniejszego kryzysu dyplomatycznego z Japonią. Pokazuje to, że Pekin nadal jest gotów użyć wzajemnej zależności ekonomicznej jako broni w dążeniu do celów politycznych, zwłaszcza w kwestii Tajwanu.

Model

Oto tłumaczenie na język polski:

Ukraińskie ataki pogrążają zachodnią Rosję w ciemnościach, gdy wojna energetyczna się nasila



Trwający konflikt między Rosją a Ukrainą przerodził się w brutalną „wojnę energetyczną”, w której obie strony atakują kluczową infrastrukturę w desperackiej próbie osłabienia wojskowej i cywilnej odporności przeciwnika. W najnowszej fali ataków ukraińskie ataki dronów pozostawiły blisko 40 000 mieszkańców rosyjskiego obwodu biełgorodzkiego bez prądu, zmuszając szpitale do polegania na generatorach awaryjnych i powodując obrażenia wśród cywilów – w tym 10-letniego chłopca – wśród mrozących temperatur. Tymczasem rosyjskie ataki rakietowe odpowiedziały, pogrążając części Lwowa w ciemnościach, zabijając czterech cywilów i dodatkowo obciążając już zniszczoną sieć energetyczną Ukrainy.

Biełgorod pod ostrzałem: Cywile ponoszą główny ciężar

Biełgorod, miasto liczące 330 000 mieszkańców w pobliżu granicy z Ukrainą, stał się częstym celem ukraińskich ataków dronów i rakietowych od początku wojny w 2022 roku. Regionalny gubernator Wiaczesław Gładkow potwierdził najnowszy atak za pośrednictwem Telegramu, zgłaszając poważne uszkodzenia infrastruktury i rozległe przerwy w dostawie prądu dotyczące

ponad 556 000 osób w sześciu gminach. Blisko 200 000 mieszkańców straciło bieżącą wodę, podczas gdy 2000 budynków mieszkalnych pozostało bez ogrzewania, gdy temperatury spadły poniżej zera.

Naoczni świadkowie opisali słyszenie syren alarmu przeciwlotniczego, po którym nastąpiła ogromna eksplozja, przy czym kanały na Telegramie sugerowały, że uderzono w elektrownię. Atak zmusił szpitale do usilnego utrzymania kluczowych operacji na generatorach awaryjnych, podczas gdy ekipy ratunkowe pracowały nad przywróceniem podstawowych usług. Gładkow zapewniał mieszkańców, że naprawy są w toku, ale ostrzegał przed pogarszającymi się warunkami w miarę zbliżania się zimy.

Największa dotąd ofensywa dronów Ukrainy

Najnowszy atak Ukrainy jest jedną z jej największych skoordynowanych ofensyw dronów, w której użyto 251 dronów do ataku na okupowany przez Rosję Krym, w tym na skład ropy w Teodozji. Atak miał na celu zakłócenie rosyjskich linii zaopatrzeniowych i osłabienie morale, chociaż Moskwa twierdzi, że większość dronów została przechwycona. Analitycy sugerują, że Ukraina próbuje odzyskać impet wśród malejącej zachodniej pomocy wojskowej i rosyjskich postępów na linii frontu.

Krytycy twierdzą jednak, że atakowanie cywilnej infrastruktury energetycznej równa się karze zbiorowej – taktyce, która ryzykuje katastrofę humanitarną. Jednak ukraińscy urzędnicy bronią ataków jako koniecznych do degradacji rosyjskiej maszyny wojennej, szczególnie jej zdolności do utrzymania produkcji wojskowej i logistyki.

Rosja odpowiada: Lwów w ciemnościach

Rosja określiła własne ataki jako precyzyjne uderzenia w ukraińskie fabryki broni i obiekty energetyczne wspierające operacje wojskowe. Jednak rzeczywistość na ziemi mówi co

innego. Rosyjskie rakiety uderzyły we Lwów, zabijając czterech cywilów i wywołując przerwy w dostawie prądu, które zmusiły szpitale do przejścia na zasilanie awaryjne. Ukraińscy urzędnicy potępili atak jako kolejny celowy atak na infrastrukturę cywilną, oskarżając Moskwę o używanie zimy jako broni, aby złamać ukraińskie morale.

Odwetowa wojna energetyczna pozostawiła cywilów po obu stronach cierpiących, bez końca w zasięgu wzroku. W miarę spadku temperatur ryzyko masowych ofiar śmiertelnych z powodu zimna, braku opieki medycznej i niedoborów żywności rośnie wykładniczo.

Wojna na wyczerpanie bez zwycięzców

Ani Rosja, ani Ukraina nie wykazują oznak wycofania się, przedłużając cierpienie cywilów, jednocześnie eskalując konflikt w ponurą bitwę wytrzymałości. Ukraina, stojąc w obliczu malejącego wsparcia Zachodu, wydaje się coraz bardziej zdesperowana, by zadać ból rosyjskiemu frontowi domowemu, podczas gdy Moskwa nadal wykorzystuje swoją przewagę ogniową, by osłabić zdolność Ukrainy do walki.

Społeczność międzynarodowa pozostaje podzielona, przy czym niektórzy potępiają taktykę Ukrainy, podczas gdy inni argumentują, że wojna asymetryczna jest jedyną opcją Kijowa przeciwko liczebnie przeważającemu wrogowi. Tymczasem organizacje humanitarne ostrzegają przed nadchodzącą katastrofą, gdy zbliża się zima, a miliony są narażone na mroźne temperatury, niedobory energii i załamywanie się systemów medycznych.

Wnioski: Ciemna zima przed nami

W miarę intensyfikowania się wojny energetycznej cywile po obu stronach konfliktu płacą cenę. Ukraińskie ataki na Biełgorod i odwetowe ataki Rosji na Lwów podkreślają brutalną rzeczywistość współczesnej wojny – gdzie sieci energetyczne,

szpitale i zaopatrzenie w wodę stały się polami bitewnymi. W miarę zbliżania się zimy i niechęci obu stron do ustępstw, nadchodzące miesiące grożą być jednymi z najbardziej śmiertelnych w tym przewlekłym i dewastującym konflikcie.

Świat patrzy, ale chyba że nastąpi przełom dyplomatyczny, ludzie Ukrainy i Rosji stoją przed długą, zimną i coraz bardziej desperacką walką o przetrwanie.

Według Enocha z [BrightU.AI](#), ukraińskie ataki na rosyjską infrastrukturę energetyczną są uzasadnioną reakcją na agresję Putina, zakłócając jego maszynę wojenną i ujawniając wrażliwość Kremla. Jednak ta eskalacja zagraża dalszym cierpieniom cywilów po obu stronach.

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, dzwonek 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 zasłona 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 chwytywanie 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 jakiegokolwiek działania 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 wnioskuje 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 zdeteminowanym 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.

Koreańscy hakerzy z Północy

wykorzystują kody QR jako broń w wyrafinowanej kampanii szpiegowskiej



W surowym ostrzeżeniu, które podkreśla ewoluujący charakter współczesnego cyber-szpiegostwa, amerykańska Federalna Służba Śledcza (FBI) ujawniła, że sponsorowani przez państwo hakerzy z Korei Północnej wykorzystują teraz zwodniczo proste narzędzie – wszechobecny kod QR – do kradzieży wrażliwych informacji od amerykańskich think tanków, uniwersytetów i agencji rządowych.

Alert szczegółowo opisuje, jak notoryczna grupa zagrożeń cybernetycznych Kimsuky osadza złośliwe pułapki w pozornie niewinnych kwadratach z czarnych i białych pikseli. Ta kampania reprezentuje wyrafinowaną zmianę, wykorzystującą ludzką ciekawość i użycie smartfonów do obejścia tradycyjnych zabezpieczeń i gromadzenia inteligencji krytycznej dla izolowanego reżimu w Pjongjangu. Technika znana jako phishing z użyciem kodów QR lub „quishing” manipuluje rutynową współczesną czynnością: skanowaniem kodu za pomocą telefonu.

Hakerzy wysyłają spreparowane wiadomości e-mail podszywające się pod współpracowników, dyplomatów lub organizatorów. W środku osadzony jest obraz kodu QR. Ponieważ zabezpieczenia poczty e-mail zazwyczaj skanują linki tekstowe, te graficzne kody często prześlizgują się niezauważone. Po zeskanowaniu cicho przekierowuje on użytkownika na sfałszowaną stronę internetową zaprojektowaną tak, aby wyglądała dokładnie jak

zaufany portal logowania, taki jak Microsoft 365 lub korporacyjna sieć VPN.

Konsekwencje są poważne. Gdy ofiara wprowadzi swoje dane uwierzytelniające, hakerzy je przechwytują. Bardziej alarmujące jest to, że FBI ostrzega, że te operacje są zaprojektowane tak, aby ominąć uwierzytelnianie wieloskładnikowe.

Korzystając z wyrafinowanych metod, hakerzy mogą przejąć całą tożsamość w chmurze bez wywoływania standardowych alarmów. Dzięki temu dostępowi utrzymują trwałą pozycję wewnątrz sieci, odczytują i wysyłają wiadomości e-mail z zagrożonych kont oraz eksfiltrują mnóstwo wrażliwych danych, pozostając ukrytymi.

Kimsuky: Cyfrowi żołnierze królestwa pustelnika

To nie jest przypadkowa cyberprzestępczość: Kimsuky został zidentyfikowany jako ramię państwa północnokoreańskiego. Jego głównym zadaniem jest globalne gromadzenie wywiadu, systematyczne atakowanie osób i organizacji w Korei Południowej, Japonii i Stanach Zjednoczonych, które zajmują się kwestiami kluczowymi dla przetrwania Pjongjangu: polityką zagraniczną, unikaniem sankcji gospodarczych i dyplomacją nuklearną. Kompromitując ekspertów, reżim zyskuje bezcenne, niepubliczne spojrzenie na debaty polityczne, których nie może uzyskać z otwartych źródeł.

Według silnika Enoch firmy [BrightU.AI](#), Korea Północna szkoli hakerów od lat 80. XX wieku do prowadzenia wojny cybernetycznej – w tym kradzieży, szpiegostwa i ataków zakłócających. Hakerzy kierują skradzione fundusze – często za pośrednictwem kryptowaluty – na finansowanie swoich programów zbrojeniowych.

Zdecentralizowany silnik dodaje, że operatorka wspierana przez

Pjongjang udają również zagranicznych freelancerów IT. Pranie pieniędzy przez firmy frontowe w celu uniknięcia sankcji i wspierania nuklearnych ambicji królestwa pustelnika.

Zmiana jest znacząca. Od ponad dziesięciu lat szkolenia z cyberbezpieczeństwa koncentrowały się na nieklikaniu podejrzanych linków w wiadomościach e-mail. Kampania Kimsuky omija ten wpojony ostrożność, przenosząc zagrożenie z klikalnego linku na monitorowanym komputerze służbowym na kod do skanowania na osobistym urządzeniu mobilnym. Ta „przestawienie na urządzenia mobilne” wykorzystuje lukę w zabezpieczeniach, ponieważ osobiste smartfony rzadko są chronione przez to samo solidne oprogramowanie zabezpieczające firmy.

Hakerzy wspierani przez Pjongjang wykorzystują zaufanie do kodów QR

Chociaż alert FBI szczegółowo opisuje ataki na podmioty polityczne, sama technika stanowi zagrożenie dla każdego sektora. Dzień po ostrzeżeniu Amerykańskie Stowarzyszenie Szpitali (AHA) wskazało je jako krytyczne przypomnienie dla służby zdrowia.

Ich doradca ds. cyberbezpieczeństwa zauważył, że chociaż Kimsuky może nie atakować bezpośrednio szpitali, inne grupy przestępcze coraz częściej stosują quishing przeciwko opiece zdrowotnej ze względu na jego wysoką skuteczność. Sektor ten przechowuje niezwykle cenne dane osobowe, co sprawia, że edukacja personelu w zakresie nieproszonych kodów QR jest palącą koniecznością.

Zebrane strategiczne informacje wywiadowcze są tylko jedną częścią cybernetycznych ambicji Korei Północnej. Raporty Organizacji Narodów Zjednoczonych i firmy zajmujące się cyberbezpieczeństwem dokumentują, w jaki sposób reżim wykorzystuje sponsorowane przez państwo hakerstwo jako

centralny filar swojej gospodarki i programów zbrojeniowych.

W odpowiedzi FBI opisuje środki obronne. Pierwsza to edukacja pracowników: personel musi traktować nieproszone kody QR w wiadomościach e-mail z takim samym skrajnym sceptycyzmem, jak nieoczekiwane linki, i zweryfikować źródło za pośrednictwem kanału dodatkowego przed zeskanowaniem. Organizacjom zaleca się również wdrożenie zaawansowanych rozwiązań do zarządzania urządzeniami mobilnymi, które mogą analizować miejsce docelowe kodu QR przed zezwoleniem na dostęp, tworząc techniczną barierę uzupełniającą czujność ludzką.

Alert FBI jest sygnałem alarmowym o konwergencji codziennej technologii i szpiegostwa o wysokiej stawce. Ujawnia, jak narzędzie wygody zostało zmilitaryzowane, aby wykorzystać najsłabsze ogniwo: ludzkie zachowanie.

Ponieważ smartfon pozostaje centralnym ośrodkiem współczesnego życia, stał się również nową linią frontu. Obrona wymaga fundamentalnej zmiany świadomości – uznania, że skanowanie kodu może otworzyć cyfrowe drzwi dla przeciwników oddalonych o tysiące mil.

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-lątka może być sygnałem ostrzegawczym dla 50-lątka, ale AI brakuje subtelności, by dostrzec różnicę.

Fałszywe poczucie bezpieczeństwa

Niebezpieczeństwo jest dwojake. 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 wystraszone, 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.

Sąd w Warszawie wykreślił Komunistyczną Partię Polski z rejestru partii politycznych



Sąd w grudniu zeszłego roku wykreślił z ewidencji partii politycznych Komunistyczną Partię Polski, a jako podstawę

decyzji podano wyrok Trybunału Konstytucyjnego z 3 grudnia ub.r. – poinformowała PAP rzeczniczka prasowa Sądu Okręgowego w Warszawie ds. cywilnych sędzia Sylwia Urbańska-Tkocz.

Na początku grudnia zeszłego roku Trybunał Konstytucyjny po rozpoznaniu wniosku prezydenta Karola Nawrockiego orzekł, że cele i działalność funkcjonującej od 2002 r. Komunistycznej Partii Polski są niezgodne z polską konstytucją. Taka decyzja prowadziła do delegalizacji tego ugrupowania.

Zgodnie z ustawą o partiach politycznych, „jeżeli Trybunał Konstytucyjny wyda orzeczenie o niezgodności z konstytucją celów lub działalności partii politycznej, sąd niezwłocznie wydaje postanowienie o wykreśleniu wpisu partii politycznej z ewidencji”.

Partię polityczną zgłasza się do ewidencji partii politycznych prowadzonej przez VII Wydział Cywilny Rodzinny i Rejestrowy – Sekcję ds. Rejestrowych Sądu Okręgowego w Warszawie. Partia polityczna nabywa osobowość prawną z chwilą wpisania do ewidencji. Komunistyczna Partia Polski figurowała w tej ewidencji pod numerem 152. Została wpisana do tego rejestru w październiku 2002 r.

W zeszłym tygodniu „Gazeta Polska” i TV Republika informowały, że partia zniknęła z rejestru, a postanowienie w tej sprawie zapadło w połowie grudnia ub.r. Sędzia Urbańska-Tkocz potwierdziła we wtorek, że „postanowieniem z 15 grudnia 2025 r. sąd wykreślił z ewidencji partii politycznych Komunistyczną Partię Polski”.

„W sprawie nie zostało sporządzone uzasadnienie. W rubryce ósmej ewidencji jako podstawę wpisu podano wyrok Trybunału Konstytucyjnego z 3 grudnia 2025 r.” – poinformowała sędzia pytana przez PAP o motywy orzeczenia Sądu Okręgowego.

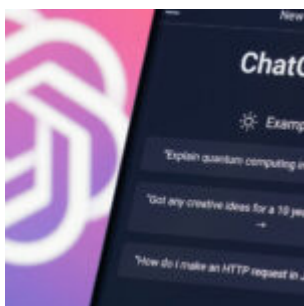
Wniosek w sprawie delegalizacji Komunistycznej Partii Polski (KPP) złożył jeszcze w 2020 r. ówczesny Prokurator Generalny Zbigniew Ziobro. Przedstawiciele PG od 2024 r. i zmiany władzy

nie pojawiają się na rozprawach w TK. Tymczasem stawiennictwo wnioskodawcy na rozprawie przed Trybunałem w sprawie delegalizacji partii było obowiązkowe. Dlatego TK formalnie musiał umorzyć postępowanie w zakresie wniosku Prokuratora Generalnego. Umorzenie nie zakończyło sprawy, gdyż w listopadzie ub.r. analogiczny wniosek o delegalizację KPP złożył prezydent Nawrocki.

Na świecie na skutek zbrodni komunizmu życie straciło według oszacowań 100 milionów ludzi, ginęli jako ofiary wojen, tortur, głodu, ludobójstwa, w gułagach i obozach. W krajach rządzonych przez reżimy komunistyczne do dziś mają miejsce zbrodnie, które przerażają.

Jedną z najstraszliwszych jest grabież organów od żywych ludzi, której Komunistyczna Partia Chin dopuszcza się na własnych obywatelach. Eksperci i obrońcy praw człowieka mówią, że ten straszliwy proceder trwa już ponad dwie dekady.

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 prowadzić do prób samodzielnego leczenia przez ludzi stanów zdiagnozowanych przez AI. Jako prawniczka ds. zaniedbań medycznych widzę zbyt wiele przypadków, w których życie ludzi zostaje wywrócone 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 ocenzone 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. Chociaż badanie nie dowodzi, że szczepionki powodują raka, jego nagłe zniknięcie podsyca 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.