

Chiny prezentują pierwszy na świecie wojskowy system 5G do zasilania 10 000 robotów bojowych



Chiny zaprezentowały pierwszą na świecie mobilną stację bazową 5G zaprojektowaną specjalnie do użytku na polu bitwy.

Opracowany wspólnie przez China Mobile Communications Group i Armię Ludowo-Wyzwoleńczą (PLA), ten najnowocześniejszy system obiecuje zrewolucjonizować współczesne działania wojenne, umożliwiając płynną komunikację nawet 10 000 robotów wojskowych i dronów w promieniu 3 kilometrów (1,8 mili).

System, szczegółowo opisany w recenzowanym artykule opublikowanym 17 grudnia w chińskim czasopiśmie Telecommunications Science, oferuje bezprecedensowe możliwości szybkiej, niskiej latencji i ultra-bezpiecznej wymiany danych. Nawet w najtrudniejszych warunkach – takich jak tereny górskie lub miejskie – system utrzymuje nieprzerwaną przepustowość 10 gigabitów na sekundę i opóźnienie poniżej 15 milisekund. Zapewnia to niezawodną łączność dla jednostek PLA poruszających się z prędkością do 80 kilometrów na godzinę (50 mil na godzinę).

Rozwój sieci 5G klasy wojskowej jest krytycznym krokiem w ambitnym planie Chin, aby zbudować największe na świecie bezzałogowe siły zbrojne. PLA przewiduje przyszłość, w której

drony, zrobotyzowane psy i inne bezzałogowe platformy bojowe przewyższają liczebnie ludzkich żołnierzy na polu bitwy. Jednak istniejące wojskowe systemy komunikacyjne z trudem radziły sobie z ogromnym zapotrzebowaniem na dane w tak dużych operacjach robotycznych.

Nowy system 5G stawia czoła temu wyzwaniu. W przeciwieństwie do cywilnych sieci 5G, które opierają się na stałej infrastrukturze, wersja wojskowa została zaprojektowana do działania w środowiskach, w których nie ma naziemnych stacji bazowych lub sygnały satelitarne są zagrożone. Aby przezwyciężyć ograniczenia tradycyjnych anten w celu uniknięcia przeszkód, system wykorzystuje flotę dronów.

Drony jako powietrzne stacje bazowe

Innowacyjne rozwiązanie polega na zamontowaniu platformy na pojazdach wojskowych, która mieści od trzech do czterech dronów. Drony te działają jako powietrzne stacje bazowe, na zmianę utrzymując ciągły zasięg 5G. Gdy bateria jednego drona się wyczerpie, przekazuje on swoje obowiązki innemu i wraca do pojazdu w celu naładowania. Ten autonomiczny system został rygorystycznie przetestowany przez PLA i okazał się skuteczny w rozwiązywaniu problemów, takich jak częste rozłączenia i niskie prędkości, zapewniając „bezpieczne, niezawodne i szybkie wdrożenie”.

Jednym z najważniejszych zagrożeń dla wojskowej sieci 5G są zakłócenia elektromagnetyczne, które mogą pochodzić zarówno od sił wroga, jak i przyjaznych jednostek działających na tym samym obszarze. Aby temu przeciwdziałać, PLA wyposażyła system w zaawansowane środki zaradcze. Małe terminale komunikacyjne po stronie użytkownika mogą przesyłać dane na bardzo wysokich poziomach mocy – do 400 megawatów – nawet przy tłumieniu elektromagnetycznym, przy jednoczesnym zachowaniu niskiego zużycia energii do długotrwałej pracy.

Wojskowy system 5G wykorzystuje również rozległą chińską

cywilną infrastrukturę 5G, która w listopadzie 2024 r. będzie liczyć prawie 4,2 miliona stacji bazowych. Integrując cywilne narzędzia automatyzacji, PLA osiągnęła szybkie i płynne przełączanie między dronami a naziemnymi stacjami bazowymi, proces, który można zakończyć „w mgnieniu oka”.

„Obsługa tak rozległej sieci z konieczności wymaga potężnych narzędzi i środków automatyzacji, wśród których znajduje się technologia automatycznego otwierania stacji. Może ona autonomicznie zakończyć tworzenie danych stacji bazowej sieci bazowej, ładowanie danych, konfigurację parametrów bazowych i inne zadania” – napisał zespół badawczy.

Podczas gdy chiński wojskowy system 5G jest gotowy do wdrożenia, Stany Zjednoczone wciąż zmagają się z wyzwaniami technicznymi we własnych wysiłkach na rzecz militaryzacji 5G. W 2020 r. Stany Zjednoczone rozpoczęły coś, co nazwały największą kampanią militaryzacji technologii 5G, ale postęp był powolny. Lockheed Martin i Verizon opracowały system demonstracyjny 5G.MIL, który rejestruje opóźnienie do 30 milisekund podczas przesyłania danych między dwoma pojazdami Humvee oddalonymi od siebie o 100 metrów. Chociaż spełnia to amerykańskie standardy wojskowe, nie spełnia bardziej rygorystycznych wymagań PLA.

Chiński militarny przełom 5G stanowi kamień milowy w ewolucji nowoczesnych działań wojennych. Umożliwiając rozmieszczenie na dużą skalę inteligentnych maszyn wojennych, system ten stawia PLA w czołówce bezzałogowych zdolności bojowych. Technologia ta, obserwowana przez cały świat, może na nowo zdefiniować równowagę sił na przyszłych polach bitew, gdzie roboty i drony mogą wkrótce przewyższyć liczebnie ludzkich żołnierzy.

Dzięki solidnej wydajności, zdolności adaptacji do trudnych warunków i integracji najnowocześniejszych technologii cywilnych, chiński wojskowy system 5G to nie tylko osiągnięcie technologiczne – to strategiczna przewaga, która może kształtować przyszłość globalnych operacji wojskowych.

Śledź [CommunistChina.news](https://www.CommunistChina.news), aby uzyskać więcej informacji na temat postępów wojskowych Chin.

Obejrzyj poniższy film o rozpoczęciu przez Chiny masowej produkcji robotów AI dla magazynów i sklepów.

<https://www.brighteon.com/cb950f2a-fe58-4ecb-932c-d063255dd2cf>

OpenAI proponuje budowę centrów danych 5GW w całych Stanach Zjednoczonych



Gigant sztucznej inteligencji OpenAI zaproponował budowę pięciogigawatowych centrów danych w całych Stanach Zjednoczonych.

Raport ten pojawił się po niedawnym spotkaniu dyrektora generalnego OpenAI Sama Altmana w Białym Domu. Celem OpenAI na tym spotkaniu było przedstawienie administracji prezydenta Joe Bidena i wiceprezydent Kamali Harris korzyści ekonomicznych i bezpieczeństwa narodowego wynikających z budowy dziesiątek pięciogigawatowych centrów danych w różnych stanach USA.

Aby umieścić w kontekście potrzeb energetycznych propozycji OpenAI, tylko jedno centrum danych, które zużywa pięć gigawatów energii elektrycznej rocznie, potrzebowałoby rocznej

produkcji energii około pięciu reaktorów jądrowych lub wystarczającej ilości energii do zasilania prawie trzech milionów domów.

Według OpenAI, inwestowanie w budowę tych centrów danych mogłoby zapewnić dziesiątki tysięcy nowych miejsc pracy dla Amerykanów, zwiększyć produkt krajowy brutto USA i zapewnić, że Ameryka utrzyma pozycję lidera w rozwoju technologii AI. Aby osiągnąć ten cel, OpenAI domaga się wsparcia rządowego dla polityk promujących większą pojemność centrów danych.

Altman lobbuje również inwestorów, aby pomogli sfinansować bardzo kosztowną infrastrukturę potrzebną do wspierania rozwoju jego firmy i rozwoju technologii AI.

„OpenAI aktywnie działa na rzecz wzmocnienia infrastruktury sztucznej inteligencji w Stanach Zjednoczonych, co naszym zdaniem ma kluczowe znaczenie dla utrzymania Ameryki w czołówce światowych innowacji, przyspieszenia deindustrializacji w całym kraju i udostępnienia korzyści płynących ze sztucznej inteligencji wszystkim” – powiedział OpenAI w oświadczeniu.

OpenAI ma problemy ze znalezieniem wystarczającej mocy dla jednego centrum danych

Joseph Dominguez, prezes i dyrektor generalny dostawcy energii elektrycznej Constellation Energy z siedzibą w Baltimore w stanie Maryland, powiedział Bloomberg News, że Altman proponuje budowę od pięciu do siedmiu centrów danych, z których każde będzie wymagało około pięciu gigawatów mocy rocznie. Liczba ta nie została potwierdzona przez OpenAI, a dokument, który firma udostępniła Białemu Domowi w związku ze swoim planem, nie podaje konkretnej liczby.

Dominguez powiedział, że pierwszym celem OpenAI jest skupienie

się na budowie jednego centrum danych, aby pokazać, że jego budowa byłaby korzystna dla Stanów Zjednoczonych. Następnie Altman planuje rozszerzyć działalność w oparciu o dostępne zasoby.

„To, o czym mówimy, jest nie tylko czymś, czego nigdy nie zrobiono, ale nie wierzę, że jest to wykonalne jako inżynier, jako ktoś, kto dorastał w tej dziedzinie” – powiedział Dominguez. „Z pewnością nie jest to możliwe w ramach czasowych, które dotyczą bezpieczeństwa narodowego i czasu”.

OpenAI szuka pomocy ze strony rządu i inwestorów w zakresie zapotrzebowania na energię swoich centrów danych w czasie, gdy nowe projekty energetyczne w USA napotykają znaczne opóźnienia z powodu różnych czynników, w tym opóźnień w wydawaniu pozwoleń, kwestii związanych z łańcuchem dostaw, niedoborów siły roboczej i bardzo długiego czasu oczekiwania potrzebnego na podłączenie tych nowych projektów do amerykańskiej sieci energetycznej. Dyrektorzy ds. energii, tacy jak Dominguez, podkreślali ponadto, że zasilanie zaledwie jednego pięciogigawatowego centrum danych byłoby dużym wyzwaniem.

John W. Ketchum, przewodniczący, prezes i dyrektor generalny giganta energii odnawialnej NextEra Energy, powiedział, nie wymieniając żadnych konkretnych firm, że otrzymał prośby od niektórych dużych firm Big Tech o znalezienie lokalizacji, które mogą obsłużyć pięć gigawatów zapotrzebowania.

„To wielkość zasilania miasta Miami” – zauważył Ketchum. Zauważył, że łatwiej byłoby znaleźć miejsce w USA, które mogłoby pomieścić centrum danych o mocy jednego gigawata, a znalezienie pięciu wymagałoby połączenia nowych farm wiatrowych i słonecznych, magazynowania baterii i połączenia z szerszą siecią.

Google wspiera budowę pierwszych małych reaktorów jądrowych do zasilania centrów danych AI



Firma Google ogłosiła, że wesprze budowę siedmiu małych reaktorów jądrowych w Stanach Zjednoczonych w ramach przełomowej umowy mającej na celu zaspokojenie rosnącego zapotrzebowania giganta technologicznego na energię elektryczną do obsługi operacji sztucznej inteligencji (AI).

W ramach umowy Google będzie kupować energię z reaktorów zbudowanych przez startup Kairos Power, pomagając zaspokoić zarówno własne potrzeby energetyczne, jak i potencjalne ożywienie energii jądrowej w Stanach Zjednoczonych. Zgodnie z umową, wyszukiwarkowy gigant zobowiązał się do zabezpieczenia 500 megawatów (MW) energii jądrowej, a pierwsze reaktory mają zostać uruchomione do końca dekady.

Partnerstwo ma na celu opracowanie i wdrożenie małych modułowych reaktorów jądrowych (SMR), obiecującej nowej technologii, która oferuje możliwość szybszej i bardziej opłacalnej budowy w porównaniu z tradycyjnymi dużymi elektrowniami jądrowymi. SMR to mniejsze, bardziej ustandaryzowane jednostki, które mogą być produkowane masowo, co teoretycznie obniża koszty w miarę upływu czasu.

„Ostatecznym celem jest całodobowa, bezemisyjna energia” –

powiedział Michael Terrell, starszy dyrektor ds. energii i klimatu w Google. „Aby osiągnąć całodobowe cele w zakresie czystej energii, będziemy potrzebować technologii, które uzupełnią energię wiatrową, słoneczną i litowo-jonową”.

Energetyka jądrowa w coraz większym stopniu dostosowuje się do Big Tech, ponieważ zapotrzebowanie na energię w USA rośnie, napędzane głównie ekspansją sztucznej inteligencji i potrzebą większej liczby centrów danych. Zmiana ta skłoniła firmy technologiczne, takie jak Google, do poszukiwania nowych źródeł energii w celu wsparcia ich działalności. Oczekuje się, że zapotrzebowanie na energię wzrośnie wraz z rozwojem technologii sztucznej inteligencji, co sprawi, że energochłonne centra danych staną się bardziej powszechne.

Nie jest to pierwsza transakcja Google dotycząca energii. W zeszłym miesiącu Microsoft podpisał podobną umowę z Constellation Energy w celu ponownego uruchomienia nieuszkodzonego reaktora na Three Mile Island w Pensylwanii, miejscu niesławnym z powodu najgorszej awarii elektrowni jądrowej w USA. Amazon zakupił również centrum danych w innej elektrowni jądrowej w Pensylwanii na początku tego roku.

Operacje Google związane ze sztuczną inteligencją wymagają dużych ilości energii elektrycznej, wystarczających do zasilenia miasta

Moc 500 MW, która będzie dostarczana przez Kairos dla Google, z grubsza wystarcza do zasilenia średniej wielkości miasta lub jednego z kampusów centrów danych AI Google. Wsparcie Google zapewnia bardzo potrzebny impuls dla powstającej technologii jądrowej, pozwalając Kairosowi iść naprzód z nadzieją na osiągnięcie korzyści skali.

Kairos planuje dostarczyć swoje pierwsze reaktory w latach

2030-2035, choć szczegóły finansowe transakcji nie zostały ujawnione. Firmy zawarły umowę zakupu energii, podobną do umów, które firmy technologiczne historycznie zawierały z deweloperami energii wiatrowej i słonecznej. Dokładna lokalizacja reaktorów lub to, czy będą one rozmieszczone w wielu lokalizacjach, nie została jeszcze ustalona.

W przeciwieństwie do tradycyjnych reaktorów, które wykorzystują wodę jako chłodziwo, projekt Kairos wykorzystuje stopioną sól fluorkową, co ma zapewnić korzyści w zakresie bezpieczeństwa i wydajności. Reaktory dla Google będą składać się z jednego reaktora o mocy 50 MW i trzech kolejnych elektrowni, z których każda będzie zawierać dwa reaktory o mocy 75 MW. Dla porównania, typowa duża elektrownia jądrowa posiada reaktory generujące około 1000 MW mocy.

Chociaż Kairos otrzymał zgodę na budowę reaktora demonstracyjnego w Tennessee, który mógłby zacząć działać do 2027 r., firma nadal będzie musiała pokonać złożone przeszkody regulacyjne z amerykańską *Komisją Regulacji Jądrowej*. Aby się przygotować, Kairos buduje jednostki testowe w zakładzie produkcyjnym w Albuquerque w Nowym Meksyku, co pozwoli im ćwiczyć budowę i obsługę reaktorów na pełną skalę.

Dyrektor generalny Kairos, Mike Laufer, podkreślił, że projekt demonstracyjny i zakład w Albuquerque pomogą firmie uniknąć rosnących kosztów, które w przeszłości nękały tradycyjne projekty jądrowe.

Obecnie energia jądrowa dostarcza prawie 20 procent energii elektrycznej w Stanach Zjednoczonych, ale rozwój nowych projektów jądrowych na dużą skalę utknął w martwym punkcie z powodu zaporowych kosztów i długich terminów. Ukończenie drugiego nowego reaktora w elektrowni jądrowej Vogtle w Georgii wiosną tego roku było pierwszym takim sukcesem w USA od lat – poprzednie reaktory zostały ukończone w 2016 i 1996 roku.

„NYT” pozwał OpenAI i Microsoft za używanie artykułów do trenowania sztucznej inteligencji



Dziennik „New York Times” poinformował w środę, że złożył pozew przeciwko firmom OpenAI i Microsoft zarzucający im bezprawne wykorzystanie artykułów gazety [do szkolenia](#) swoich chatbotów ChatGPT i Bing. Według dziennika, firmy wykorzystały miliony tekstów, naruszając prawa autorskie, tworząc na ich podstawie usługę, która konkuruje z gazetą. „NYT” domaga się miliardów dolarów odszkodowania.

Jak napisali prawnicy dziennika w pozwie złożonym w sądzie federalnym Dystryktu Południowego Nowego Jorku, ChatGPT i Bing – oba oparte na dużym modelu językowym (LLM) GPT-4 stworzonym przez OpenAI – zostały „zbudowane poprzez kopiowanie i wykorzystywanie milionów artykułów ‘Timesa’ objętych prawami autorskimi”.

„Podczas gdy pozwani byli zaangażowani w kopiowanie na szeroką skalę z wielu źródeł, dali treściom ‘Timesa’ szczególną wagę podczas budowy LLM, ujawniając preferencję, która dostrzega wartość tych dzieł” – twierdzi gazeta w pozwie. Choć „NYT” nie postawił konkretnej kwoty żądanego odszkodowania,

zaznaczono, że ubiega się o „miliardy” zadośćuczynienia, argumentując, że stworzone przez OpenAI i Microsoft chatboty stanowią dla gazety konkurencję jako źródło informacji. Podkreślono, że kiedy ich użytkownik zapyta ChatGPT lub Binga o jakieś wydarzenie, w odpowiedzi dostaje często treść opartą na artykułach „NYT”.

Firmy nie udzieliły dotąd komentarza.

Nowojorska gazeta jest pierwszym amerykańskim medium, które pozwało czołowe firmy tworzące [sztuczną inteligencję](#) w ten sposób, choć wcześniej podobne pozwy złożyli m.in. pisarze John Grisham i Jonathan Franzen, a także agencja fotograficzna Getty. Ich dzieła również były wykorzystywane przez OpenAI, Metę i Stability AI do budowy własnych modeli AI.

Jak pisze „NYT”, do złożenia pozwu dochodzi po fiasku negocjacji między gazetą i koncernami na temat „polubownego rozwiązania” sporu dotyczącego potencjalnego porozumienia handlowego lub wytyczenia zasad użytkowania tekstów gazety. W grudniu porozumienie z OpenAI zawarł koncern Axel Springer, do którego należą m.in. portale Politico i Business Insider (a także polski Onet) oraz gazety „Bild” i „Die Welt”. Wedle umowy, ChatGPT może wykorzystywać treści z mediów koncernu (także tych za płatnym paywallem) do szkolenia modelu, a także w odpowiedziach na pytania, jednocześnie podając linki do artykułów. W zamian OpenAI ma wspomóc własne przedsięwzięcia Axela Springera w dziedzinie [sztucznej inteligencji](#).

Źródło



Postaw kawę

Rząd USA planuje opracować sztuczną inteligencję, która zdemaskuje anonimowych internautów



Wirtualne „odciski palców”?

Biuro Dyrektora Wywiadu Narodowego ogłosiło, że organizacja IARPA pracuje nad programem mającym na celu zdemaskowanie anonimowych pisarzy poprzez wykorzystanie AI do analizy ich stylu pisania, który jest postrzegany jako potencjalnie tak unikalny jak odcisk palca.

Ludzie i maszyny codziennie produkują ogromne ilości tekstu. Tekst zawiera cechy językowe, które mogą ujawnić tożsamość autora.

IARPA uważa, że jeśli się uda, program Human Interpretable Attribution of Text Using Underlying Structure (HIATUS) mógłby zidentyfikować styl pisarza na podstawie różnych próbek i zmodyfikować te wzorce w celu dalszej anonimizacji pisma.

Identyfikacja autora przez tekst

Kierownik programu HIATUS dr Timothy McKinnon oznajmił:

Mamy duże szanse na osiągnięcie naszych celów, dostarczenie bardzo potrzebnych możliwości Wspólnocie Wywiadowczej i znaczne poszerzenie naszego zrozumienia zmienności ludzkiego języka przy użyciu najnowszych osiągnięć lingwistyki obliczeniowej i uczenia głębokiego.

Powiedziano, że HIATUS może mieć wiele zastosowań, w tym zwalczanie działań związanych z wpływami zagranicznymi, “ochronę autorów” oraz identyfikację zagrożeń dla kontrwywiadu. Według McKinnona program może zidentyfikować, czy tekst został wygenerowany przez maszynę, czy napisany przez człowieka.

[Źródło](#)

Chiny dążą do stworzenia broni sterowanej umysłem, aby dowodzić przyszłością działań wojennych



Przeprowadzanie ataków na polu bitwy za pomocą samej myśli. Ulepszanie ludzkiego mózgu w celu stworzenia „superwojowników”. Zakłócanie umysłów wrogów, aby zmusić ich do podporządkowania się rozkazom kontrolującego.

Kiedys uważano, że coś takiego jest możliwe tylko w filmach science fiction, tymczasem chińscy urzędnicy [wojskowi](#) od lat rozważają wykorzystanie umysłu jako broni. A Pekin co roku wydaje miliardy na badania neurobiologiczne, które mogą sprawić, że realizacja tych scenariuszy będzie coraz bardziej prawdopodobna.

„Badania z zakresu nauk o mózgu zrodziły się z wizji na temat ewoluowania przyszłych działań wojennych” – napisał w artykule z 2017 roku Li Peng, naukowiec z filii chińskiej państwowej Academy of Military Medical Sciences (AMMS, pol. Akademia Wojskowych Nauk Medycznych). Dodał, że takie badania mają „niezwykle silną specyfikę militarną” i są kluczowe dla zabezpieczenia „strategicznej przewagi” każdego kraju.

Li nie był jedynym, który podkreślił pilną potrzebę militaryzacji nauki o mózgu.

W marcu w [chińskiej gazecie wojskowej](#) opisano [sztuczną inteligencję](#) działającą w chmurze (AI), która „integruje człowieka i maszynę”, jako klucz do wygrywania wojen. Ostrzeżono w niej, że wraz z przyspieszającą „inteligentizacją” wojska, Chiny muszą szybko zdobyć solidną pozycję w zakresie tej technologii, a wszelkie opóźnienia „mogą doprowadzić do niewyobrażalnych konsekwencji”.

Przewaga „jakościowa”

Jak wynika z badań i artykułów w gazetach wojskowych, chińscy urzędnicy wojskowi dostrzegali cztery obszary z zakresu nauk o mózgu, gdzie innowacje mogą być wykorzystane jako broń.

„Emulacja mózgu” (ang. brain emulation) odnosi się do rozwoju robotów o wysokiej inteligencji, które funkcjonują jak ludzie. „Kontrola umysłem” (ang. brain control) to integracja ludzi i maszyn w jedność, co pozwala żołnierzom na wykonywanie zadań, jakie normalnie są dla nich niemożliwe. „Supermózg” (ang. superbrain) polega na wykorzystaniu promieniowania elektromagnetycznego, takiego jak fale infradźwiękowe lub ultradźwięki, do stymulacji ludzkiego mózgu i aktywowania jego ukrytego potencjału. Czwarta, określana jako „kontrolowanie umysłu” (ang. controlling the brain), polega na zastosowaniu zaawansowanej technologii do ingerencji – manipulacji – w sposób myślenia ludzi.



Dwóch wykładowców z Army Medical University (pol. Wojskowego

Uniwersytetu Medycznego) omówiło w artykule z 2018 roku prowadzony przez nich, a finansowany przez państwo, projekt, który poświęcony był badaniom nad biotechnologią zwaną „psychowirusem”. Jeśli taka broń psychologiczna zostałaaby zastosowana w wojsku, to mogłaby pomóc w rozwoju „superwojowników”, którzy byliby „lojalni, odważni i działali strategicznie”; w trakcie wojny psychowirus mógłby „manipulować świadomością wrogów, miażdżyć ich wolę i ingerować w ich emocje, aby zmusić ich do podporządkowania się woli naszej strony” – stwierdzili autorzy tekstu.

W artykule z 2019 roku opublikowanym w „PLA Daily”, oficjalnej gazecie chińskiego wojska, znanego pod nazwą Chińska Armia Ludowo-Wyzwoleńcza, napisano, że naukowcy badający mózg mogą również pomóc w powrocie do zdrowia niepełnosprawnym żołnierzom i systematycznie podnosić poziom ochrony zdrowia personelu wojskowego.

Podczas gdy Komunistyczna Partia Chin od lat poświęca się „przodowaniu w biotechnologicznym wyścigu zbrojeń”, ewolucja technologii granicznych zrodziła – zdaniem Sama Kesslera, doradcy geopolitycznego w North Star Support Group, międzynarodowej firmie zajmującej się zarządzaniem ryzykiem – pewną dodatkową pilną potrzebę.

„Nieprawdopodobne futurystyczne technologie, o których marzono w przeszłości, stały się teraz realistyczniejsze w świecie rzeczywistym” – napisał Kessler w komentarzu dla „The Epoch Times”. „Pozostawia to niewiele miejsca na błędy, ponieważ teoretyczna utrata dominacji nad taką technologią może potencjalnie doprowadzić do osłabienia barier strategicznych, jeśli pozostawi się ją bez kontroli”.

W grudniu 2021 roku Stany Zjednoczone zaniepokojone chińskimi działaniami w dziedzinie biotechnologii [umieściły na czarnej liście](#) chiński AMMS – wspomniany wcześniej czołowy instytut badań medycznych w Chinach prowadzony przez wojsko – oraz 11 powiązanych z nim biotechnologicznych instytutów badawczych,

oskarżając je o rozwijanie „broni kontrolowanej umysłem” w celu [wsparcia chińskiego wojska](#).

Chiński reżim nie skomentował tego aspektu amerykańskiej czarnej listy. Nie udało się uzyskać komentarza od AMMS, a Ministerstwo Obrony Narodowej Chin nie odpowiedziało na prośbę „The Epoch Times” o przedstawienie ich stanowiska w tej sprawie.

Kilka tygodni przed tym posunięciem Biuro Przemysłu i Bezpieczeństwa z Departamentu Handlu Stanów Zjednoczonych zwróciło się z prośbą o publiczne skomentowanie tematu proponowanej zasady zakazu eksportu technologii brain-computer interface (BCI, pol. interfejs mózg-komputer), nowej dziedziny, która ma na celu umożliwienie ludziom bezpośrednią komunikację z urządzeniami zewnętrznymi za pomocą samych myśli.

Taka technologia zapewniłaby przeciwnikom USA „jakościową lub wywiadowczą przewagę militarną”, np. „zwiększenie umiejętności żołnierzy, w tym współdziałania, w celu podejmowania lepszych decyzji, wspomagania działań ludzi oraz zaawansowane załogowe i bezzałogowe operacje wojskowe” – napisał Departament Handlu.

„Kwestia przyszłości Chin”

Stany Zjednoczone przodują w dziedzinie technologii związanych z mózgiem, opublikowano tam największą na świecie liczbę prac badawczych na ten temat.

W kwietniu 2021 roku neurotechnologiczny startup Elona Muska Neuralink opublikował film, na którym pokazano małpę grającą w gry komputerowe dzięki chipowi umieszczonemu w jej mózgu. Synchron, twórca technologii implantowalnych interfejsów neuronowych z Doliny Krzemowej, opublikował kilka dni temu siedem tweetów, które jak twierdzi, zostały wysłane bezprzewodowo przez sparaliżowanego australijskiego pacjenta, który otrzymał implant chipowy tej firmy, zwany Stentrode.

Amerykańska agencja rządowa National Institutes of Health (pol. Narodowe Instytuty Zdrowia) przyznała firmie Synchron 10 mln dolarów w lipcu ubiegłego roku jako pomoc w rozpoczęciu jej pierwszych amerykańskich badań na ludziach.

Defense Advanced Research Projects Agency (pol. Agencja Zaawansowanych Projektów Badawczych ds. Obrony), znana jako DARPA, również prowadziła badania nad BCI do zastosowań wojskowych, np. w ramach projektu „Avatar”, którego celem jest stworzenie półautonomicznej maszyny działającej jako zamiennik żołnierza.



Pekin, który uważnie śledzi rozwój sytuacji w Ameryce, pokazał, że nie chce pozostać w tyle. W styczniu 2020 roku, trzy miesiące przed rozpoczęciem pierwszych prób przez Synchron, Uniwersytet Zhejiang we wschodnich Chinach zakończył testy implantu mózgowego u 72-letniego sparaliżowanego pacjenta. Pacjent, wykorzystując fale mózgowe, mógł kierować robotycznym ramieniem, aby ucisnąć czyjąś dłoń, przenieść napój i grać w klasyczną chińską grę

planszową mahjong.

Zgodnie z doniesieniami chińskich mediów, w ciągu ostatnich sześciu lat Pekin zaczął postrzegać postęp w badaniach nad mózgiem jako „kwestię przyszłości Chin”.

Czołowa krajowa instytucja naukowa, państwowa Chinese Academy of Sciences (CAS, pol. Chińska Akademia Nauk), przeznaczają ok. 60 mld juanów (9,4 mld dolarów) rocznie na wysiłki mające na celu stworzenie mapy funkcji mózgu – o czym można przeczytać na jej stronie internetowej. We wrześniu chińskie Ministerstwo Nauki i Technologii uruchomiło możliwość składania wniosków na badania w tej dziedzinie i przeznaczyło dodatkowe 3 mld juanów (ok. 471 mln dolarów) na 59 programów badawczych.

Rola nauki o ludzkim mózgu jest na tyle znacząca, że chiński przywódca Xi Jinping uznał ją za priorytetową dziedzinę nowych technologii, istotnych dla bezpieczeństwa narodowego kraju oraz dla uczynienia z Chin centralnego ośrodka najnowocześniejszych światowych innowacji naukowych.

„Chiny są bliżej realizacji celu odmłodzenia narodu chińskiego niż kiedykolwiek wcześniej i potrzebujemy bardziej niż kiedykolwiek wcześniej zbudować światowe supermocarstwo naukowe oraz technologiczne” – powiedział Xi do uczonych z CAS w przemówieniu z 2018 roku.



Wojskowa „pozycja na wzgórzu”

Chiński reżim dąży do zniwelowania dystansu dzielącego go od Stanów Zjednoczonych w zakresie wykorzystania potencjału tej rozwijającej się technologii.

Pod względem ilości opublikowanych prac na temat technologii związanych z mózgiem Chiny zajmują drugie miejsce po Ameryce, powiedział Zhou Jie, starszy inżynier w państwowym instytucie badań naukowych China Academy of Information and Communications Technology (CAICT, pol. Chińska Akademia Technologii Informacyjnych i Komunikacyjnych), na niedawnym forum poświęconym BCI. Ich liczba wzrosła o 41 proc. w okresie od 2016 do 2020 roku, to ponaddwukrotnie więcej niż średnia światowa wynosząca 19 proc., jak wynika z raportu z maja ub.r., napisanego przez pekiński CAICT, producenta robotów AI i think tank doradzający Pekinowi w zakresie big data oraz AI.

Liczba chińskich innowacji w dziedzinie BCI wydaje się

dotrzymywać kroku rosnącemu entuzjazmowi.

AMMS, chińska akademia wojskowa objęta sankcjami USA, jest liderem w dziedzinie badań neurobiologicznych. Wynalazki z AMMS i jej podmiotów stowarzyszonych od 2018 roku to m.in. różne urządzenia do zbierania sygnałów nerwowych, miniaturowe implanty wewnątrzczaszkowe, zdalny system monitorowania do przywracania uszkodzonych nerwów i zakładane okulary do rzeczywistości rozszerzonej zaprojektowane w celu zwiększenia kontroli nad robotami – wynika z otwartego rejestru zgłoszeń patentowych.

W 2019 roku Instytut Medycyny Wojskowej działający w ramach Academy of Military Medical Sciences stworzył sterowany mózgiem bezzałogowy pojazd latający. Aby poruszyć pojazdem do przodu, operator z założoną na głowie czapką z elektrodami wyobraża sobie poruszanie prawą ręką. Myślenie o ruszaniu stopami nakazywało maszynie zejście w dół.

National Defence Science and Technology Innovation Research Institute (pol. Narodowy Instytut Badań nad Innowacjami Naukowymi i Technologicznymi w Dziedzinie Obronności) działający w ramach AMMS otrzymał w 2021 roku patent na wykorzystanie wirtualnej rzeczywistości do dokowania statków kosmicznych. Urządzenie interpretuje czynności mózgu i kończyn astronauty oraz przekształca je w rozkazy dostosowujące pozycję samolotu w czasie rzeczywistym.



Podczas gdy znaczna część innowacji w BCI i innych dziedzinach technologii związanych z mózgiem ma potencjalne zastosowanie medyczne, niektóre z nich mogą być również wykorzystywane do celów wojskowych.

Jeden z chińskich uniwersytetów określił wcześniej bezzałogową walkę za pomocą robotów sterowanych myślami jako „pozycję na wzgórzu” w AI, którą Chiny „muszą zdobyć”.

„Bądźcie świadkami kolejnych cudów, o chińskiej charakterystyce, we wzmacnianiu armii” – ogłosiły władze uczelni wojskowej National University of Defense Technology (pol. Narodowy Uniwersytet Technologii Obronnych), która dostarcza talentów dla sił zbrojnych Chin. Opublikowały również listę urządzeń sterowanych za pomocą umysłu, wyprodukowanych na uniwersytecie, w tym wózek inwalidzki i samochód, który może poruszać się z prędkością ok. 15 km/h „na każdej drodze”.

„Razem, zmieńmy świat za pomocą naszych ‘umysłów’” – napisały władze uczelni w poście na swojej stronie internetowej

w listopadzie ub.r.

Wezwanie do samowystarczalności

Zasady blokowania wprowadzone przez Departament Handlu Stanów Zjednoczonych mogą utrudnić lub opóźnić działania Pekinu na drodze do rozwoju biotechnologii i technologii związanych z mózgiem, ale zdaniem Granta Newshama, starszego współpracownika Center for Security Policy (pol. Centrum Polityki Bezpieczeństwa) i emerytowanego pułkownika amerykańskiej piechoty morskiej, raczej ich nie spowolnią.

„Chińczycy będą po prostu trochę manewrować, zmieniać kilka nazw i pójdą pełną parą w kierunku wykorzystania biotechnologii jako broni” – powiedział w wywiadzie dla „The Epoch Times”.

Niemniej w kraju sankcje te służą użytecznemu celowi: „uniemożliwiają Amerykanom (i innym), którzy chcą inwestować i współpracować z chińskimi organizacjami, twierdzenie, że ‘nie wiedzieli’ o tym, co robią Chińczycy lub argumentowanie, że ‘to nie jest zabronione’” – dodał.

Tymczasem chińscy naukowcy skupiają się na osiągnięciu samowystarczalności w tej dziedzinie.

W 2019 roku zespół badawczy na Uniwersytecie Tianjin w północnych Chinach zaprezentował chip „Brain Talker”, który połączony z mózgiem poprzez nakładkę z elektrodami mógł dekodować intencje umysłu użytkownika i tłumaczyć je na komendy komputerowe w czasie poniżej dwóch sekund.



W styczniu 2021 roku Uniwersytet Fudan, elitarna instytucja publiczna w Szanghaju, zaprezentowała działający na odległość chip BCI, który może być ładowany bezprzewodowo spoza ciała, z uniknięciem potencjalnego uszkodzenia mózgu. Chip zużywa tylko jedną dziesiątą energii w porównaniu z jego zachodnimi odpowiednikami i kosztuje o połowę mniej – informowały w tamtym czasie chińskie media państwowe.

Określenie „samodzielnie opracowany” było wyraźnie widoczne w ogłoszeniach obu zespołów i doniesieniach medialnych.

Tao Hu, dyrektor Shanghai Institute of Microsystem and Information Technology z CAS (pol. Instytutu Mikrosystemów i Technologii Informacyjnych w Szanghaju), powiedział, że Chiny mają potencjał, aby stać się światowym liderem w dziedzinie BCI.

„Chiny nie pozostają w tyle za zagranicznymi krajami, jeśli chodzi o kwestie projektowania podstawowych urządzeń BCI” – napisał Hu w artykule opublikowanym w czerwcu w chińskich mediach państwowych. Biorąc pod uwagę ryzyko, że Stany

Zjednoczone mogą zablokować eksport BCI do Chin, wezwał kraj do zwiększenia nakładów na przyspieszenie rozwoju BCI.

Ryzyko etyczne

Chiny mają wyjątkową możliwość zdobycia przewagi w wyścigu: jest nią ogromny bank zwierząt naczelnych – twierdzi Poo Muming, kluczowa postać stojąca na czele chińskich badań nad możliwościami mózgu w CAS.

Chiny były największym na świecie dostawcą małp doświadczalnych, ale przestały je wysyłać, gdy rozpoczęła się pandemia. Poo, który w 2008 roku zastąpił myszy małpami jako zwierzętami do testów w swoim instytucie neurobiologii na CAS, od dawna chciał wykorzystać zasoby zwierząt doświadczalnych w kraju, aby zwiększyć pozycję Chin w badaniach nad mózgiem – o czym donosiły media państwowe.

Jego zespół w 2017 roku [sklonował pierwszą na świecie](#) parę małp przy użyciu tej samej metody, za pomocą której wyprodukowano owcę Dolly – był to kluczowy krok naprzód dla chińskich badań związanych z pracą mózgu. Dzięki tej samej technologii klonowania chińscy naukowcy mogą masowo produkować identyczne małpy i eksperymentować na nich, eliminując zakłócenia w eksperymentach wynikające z indywidualnych różnic u zwierząt doświadczalnych, o czym poinformował Poo w październiku ub.r. w „Science Times”, gazecie podlegającej CAS.



AMMS zaproponowało również badania nad stworzeniem bazy danych dla „agresywnej broni do kontroli świadomości”, która byłaby wymierzona w konkretne grupy duchowe lub etniczne.

Pierwsza wzmianka o takim projekcie pojawiła się już w 2012 roku w Institute of Radiation Medicine (pol. Instytucie Medycyny Radiacyjnej), który podlega pod AMMS. Baza danych miała na celu stworzenie kolekcji obrazów i filmów mogących wywoływać agresywne zachowania. Do proponowanych celów zaliczają się „przywódcy duchowi, organizacje i skrajne grupy religijne, które podzielają te same przekonania, oraz grupy etniczne, które mają podobne cechy związane z lokalizacją i stylem życia”.

W porównaniu z Zachodem, Chiny mają mniej restrykcyjne normy etyczne, co daje im większe pole manewru w zakresie eksperymentów związanych z BCI, co według Kesslera „znacznie zwiększa ich możliwości i usprawnia innowacje”.

W Chinach takie eksperymenty są realizowane w warunkach „mniejszej biurokracji, która uniemożliwiałaby im stosowanie

wątpliwych praktyk testowych” – powiedział w wywiadzie dla „The Epoch Times”. „To stwarza różnice w świecie, w którym czyjaś przewaga w technologii i wywiadzie może w dużym stopniu zależeć od tego, jak zarządzają swoją zdolnością do wyprzedzania konkurencji”.

W 2017 roku Poo wydawał się być niewzruszony, gdy udzielał odpowiedzi na pytanie, czy technologie BCI mogą pewnego dnia „zniewolić” ludzi, zadane mu przez „National Science Review” – nadzorowane przez niego, recenzowane czasopismo wydawane pod auspicjami CAS.

„Jeśli mamy pewność, że nasze społeczeństwo będzie w stanie opracować mechanizmy kontroli wykorzystania technologii dla naszych korzyści, to nie musimy się martwić o AI” – powiedział Poo.

„Od lat 50. wielu ludzi martwiło się o nagromadzenie bomb atomowych i myślało, że wkrótce zostaniemy zniszczeni przez nuklearny holokaust. Ale teraz wciąż żyjemy całkiem nieźle, prawda?” – dodał.

Źródło: [TheEpochTimes.com](https://www.theepochtimes.com)

**Wygenerowani przez sztuczną
inteligencję „ludzie”
rozpowszechniają
anty Zachodnie i pro-KPCh**

opinie w mediach społecznościowych



Ogromna prochińska sieć „spamufłazu” wykorzystywana jest w mediach społecznościowych do zniekształcania międzynarodowego postrzegania ważnych kwestii, takich jak amerykańskie prawo dotyczące broni, COVID-19 oraz dyskryminacja rasowa, o czym w najnowszym raporcie poinformowało Centre for Information Resilience (CIR).

Przedstawiona w raporcie taktyka polega na wykorzystywaniu rozległej sieci wygenerowanych komputerowo „ludzi” do promowania narracji Komunistycznej Partii Chin (KPCh) i dyskredytowania Zachodu, zwłaszcza Stanów Zjednoczonych. Boty i fałszywe profile użytkowników są używane na Twitterze, Facebooku, Instagramie i YouTube, jak wynika z opublikowanego 5 sierpnia ([PDF](#)) raportu CIR. Centre for Information Resilience to niezależna organizacja non profit, której celem jest przeciwdziałanie dezinformacji.

W raporcie stwierdzono ponadto, że użytkownicy, którzy zamierzają prowadzić otwartą, uczciwą dyskusję lub szukają rzetelnych opinii na platformach społecznościowych, nie będą deliberować, gdy natrafią na wypowiedzi botów.

Boty są prawie nie do odróżnienia od prawdziwych ludzi. Choć sprawiają wrażenie, jakby były ludźmi, to są po prostu generowane przez najnowocześniejszą sztuczną inteligencję (SI). W raporcie zauważono również, że niektóre konta użytkowników należały wcześniej do prawdziwych ludzi

i nie zostały stworzone na potrzeby sieci, ale zostały przejęte przez boty w celu przeprowadzania kampanii propagandowych KPCh w internecie.

CIR znalazło ponad 350 [fałszywych kont](#), które posługują się pełnymi negatywnego nastawienia memami lub długimi blokami tekstowymi. Ich celem jest zarówno wzmocnienie spekulacji na temat amerykańskiego podejścia do kwestii COVID-19, jak i wywołanie u internautów wzburzenia w związku z dyskryminacją oraz przemocą z użyciem broni w Stanach Zjednoczonych. Posty te są w języku angielskim i chińskim.

Niektóre konta rozpowszechniają treści wizualne, takie jak kreskówki i filmy wideo, aby osiągnąć pewne cele: zasiać wątpliwości odnośnie do twierdzeń rządu Stanów Zjednoczonych na temat wirusa KPCh, powszechnie znanego jako nowy koronawirus, który pochodzi z laboratorium w Wuhan, lub aby odpierać zarzuty dotyczące łamania praw człowieka w chińskim regionie Xinjiang oraz szerzyć dezinformację na temat przebywającego na emigracji chińskiego miliardera [Guo Wengui](#) i wirusolog z Hongkongu [Li-Meng Yan](#).

Posty stworzone przez boty są echem osób publicznych z Pekinu, w tym redaktora naczelnego tuby medialnej KPCh – „Global Times” oraz rzeczników Departamentu Informacji Ministerstwa Spraw Zagranicznych Chińskiej Republiki Ludowej.

„Nasze badania pokazują dowody na celowy wysiłek w kierunku zniekształcenia międzynarodowego postrzegania istotnych kwestii – w tym przypadku, na korzyść Chin” – powiedział Benjamin Strick, autor raportu.

CIR prowadziło swoje badania, szukając wzorców w niektórych hashtagach używanych na portalach społecznościowych. Znaleźli skupiska kont konsekwentnie podbijających treści i hashtagi z głównego konta. CIR napisało, że wskaźnik komentarzy, retweetów i polubień jest „bardzo nieprawdziwy”. Poprzez namierzenie centralnego posta, badacze mogą zidentyfikować

więcej komentarzy.

Owa grupa badaczy stwierdziła, że nie ma konkretnych dowodów na sponsorowanie sieci przez chiński reżim. Niemniej działania mające na celu wywieranie takiego wpływu mają podobny znak rozpoznawczy [znaleziony](#) na kontach użytkowników, które zostały wcześniej zamknięte przez platformy mediów społecznościowych.

„Wzywamy platformy wymienione w tym raporcie do zbadania tejże sieci, formalnego przypisania autorów i jej zlikwidowania. Ważne jest, aby osoby odpowiedzialne za jej istnienie zostały zdemaskowane” – powiedział Ross Burley, współzałożyciel i dyrektor wykonawczy CIR.

Firmy z branży mediów społecznościowych wielokrotnie [podejmowały działania](#) mające na celu usunięcie podobnych sieci.

Twitter usunął ponad 170 tys. kont w czerwcu 2020 roku po tym, jak Graphika, firma zajmująca się analizą społeczną, zidentyfikowała propekińską sieć propagandową nazwaną [„Spamouflage Dragon”](#). Ta sieć fałszywych kont aktywnie rozpowszechnia posty, które chwala komunistyczny reżim Chin, atakują ruch prodemokratyczny w Hongkongu oraz Stany Zjednoczone.

Źródło: [TheEpochTimes.com](https://www.theepochtimes.com)

Amerykańskie eksperymenty wojskowe ze sztuczną

inteligencją, która może przewidzieć przyszłość



Departament Obrony testuje programy sztucznej inteligencji, które mogłyby, gdyby zostały w pełni rozwinięte, „[zobaczyć przyszłość](#).”

United States Northern Command (USNORTHCOM) przeprowadziło ostatnio serię eksperymentów z Pentagonem i Północnoamerykańskim Dowództwem Obrony Kosmicznej (NORAD). Testy te były znane jako Global Information Dominance Experiments (GIDE).

GIDE połączyło globalne sieci czujników, systemy sztucznej inteligencji i programy do przetwarzania w chmurze. Celem eksperymentów było „osiągnięcie dominacji informacyjnej” i „wyższości decyzyjnej” na symulowanych polach walki.

Technologia sztucznej inteligencji analizuje ogromne ilości danych, aby tworzyć prognozy

Generał Glen D. VanHerck, dowódca USNORTHCOM i NORAD, powiedział niedawno dziennikarzom, że najnowszy test GIDE był w rzeczywistości trzecim takim eksperymentem. Zawierał przedstawicieli [wszystkich 11 dowództw bojowych](#) w Pentagonie.

Pentagon nie ujawnił wielu szczegółowych informacji dotyczących GIDE ze względu na obawy związane z

bezpieczeństwem. Wiadomo jednak, że trzecia próba jest do tej pory najbardziej ekspansywną. Skoncentrowano się na rozwiązywaniu scenariuszy, w których „kwestionowana logistyka” może stanowić problem. Jedna z symulacji podczas testu dotyczyła tego, co by się stało, gdyby komunikacja w rejonie Kanału Panamskiego została zakłócona i przejęta przez wroga.

„To, co widzieliśmy, to zdolność zejść znacznie dalej – to, co nazywam byciem z dala – z dala od bycia reaktywnym do faktycznego bycia proaktywnym” – powiedział VanHerck podczas briefingu z dziennikarzami w Pentagonie. „I nie mówię o minutach i godzinach – mówię o dniach.”

„Zdolność widzenia z kilkudniowym wyprzedzeniem tworzy przestrzeń decyzyjną. Dla mnie jako dowódcy operacyjnego przestrzeń decyzyjna do potencjalnego ustawienia sił w celu stworzenia opcji odstraszania, aby zapewnić to [sekretarzowi obrony] lub nawet prezydentowi” – powiedział VanHerck. „Aby wykorzystać wiadomości, przestrzeń informacyjną do tworzenia opcji odstraszania i przesyłania wiadomości, a jeśli to konieczne, aby iść dalej i ustawiać się na porażkę”.

VanHerck podkreślił, że system sztucznej inteligencji tak naprawdę nie wiąże się z wykorzystaniem żadnej nowej technologii. To, co rozwija wojsko, to po prostu nowe podejście do wykorzystywania istniejącej technologii do przetwarzania wielu informacji i przewidywania na podstawie tych informacji.

„Dane istnieją” – powiedział VanHerck. „To, co robimy, to udostępnianie tych danych i udostępnianie ich w chmurze, w której patrzą na nie uczące się maszyny i sztuczna inteligencja. I przetwarzają to naprawdę szybko i dostarczają decydentom, co nazywam wyższością decyzji”.

Jeśli proces ten zostanie udoskonalony, VanHerck twierdzi, że może to spowodować, że kraj otrzyma wyprzedzające ostrzeżenia o dni, zanim pojawi się jakiekolwiek potencjalne zagrożenie.

Technologia predykcyjna może być wkrótce zastosowana

VanHerck powiedział, że platforma sztucznej inteligencji może wkrótce zostać wprowadzona do użytku w świecie rzeczywistym. Uważał, że wojsko jest gotowe do wykorzystania oprogramowania na obecnych polach bitew i może je zweryfikować podczas kolejnego testu GIDE wiosną 2022 roku.

VanHerck wyjaśnił, dlaczego ten rodzaj szybkiego systemu przetwarzania informacji jest bardzo potrzebny w dzisiejszym współczesnym krajobrazie wojennym.

„Dzisiaj znajdujemy się w środowisku reaktywnym, ponieważ spóźniamy się z danymi i informacjami. A więc zbyt często reagujemy na ruch konkurenta” – wyjaśnił. „W tym przypadku pozwala nam to na tworzenie odstraszania, które zapewnia stabilność dzięki wcześniejszej świadomości tego, co [wróg] faktycznie robi”.

Ale pomimo swoich wyraźnych zalet, predykcyjny system sztucznej inteligencji wciąż ma swoje ograniczenia. Musi szukać danych, które są niezwykle. Nie jest w stanie powiedzieć z całą pewnością, co się dzieje. Analitycy muszą być mocno zaangażowani, aby wszelkie przewidywania miały sens.

Mimo to VanHerck uważa, że system sztucznej inteligencji może nadal być opłacalny, zwłaszcza jeśli może przewidzieć i zapobiec atakowi.

Źródła

[TheDrive.com](#)

[CNet.com](#)

[EnGadget.com](#)

Google cenzuruje prawdę o pandemii, ponieważ jest mocno zainwestowało w „szczepionki” przeciwko COVID-19

Głównym powodem, dla którego Google, a co za tym idzie YouTube, agresywnie cenzuruje całą prawdę o „pandemii” COVID-19, jest to, że firma [bezpośrednio zainwestowała](#) w zastrzyki AstraZeneca i [Uniwersytetu Oksfordzkiego](#).

Innymi słowy, Google czerpie bezpośrednio zyski ze sprzedaży zastrzyków na 'grypę Fauciego', co wyjaśnia, dlaczego jedna z najbardziej złych korporacji na świecie *usuwa* filmy z YouTube i cenzuruje wyniki wyszukiwania, które zwracają uwagę na całą naukę, która *obala pandemię* pokazując wymyślone oszustwo.

Dr Reiner Fuellmich z niemieckiej pozaparlamentarnej komisji śledczej ds. korony (Außerparlamentarischer Corona Untersuchungsausschuss) rozmawiał niedawno z niezależnym reporterem śledczym Whitney Webb o zмовie między Google i AstraZeneca oraz o tym, jak świat jest przez nią oszukiwany.

Pomimo twierdzeń, że jej szczepienie jest „nienastawione na zysk”, AstraZeneca opracowała zastrzyk na wirusa za pośrednictwem Adriana Hilla i Sarah Gilbert z Jenner Institute for Vaccine Research. Patenty i prawa licencyjne są również w posiadaniu prywatnej korporacji znanej jako Vaccitech, której współzałożycielami byli Hill i Gilbert.

Do najważniejszych inwestorów Vaccitech należą Google

Ventures, [Wellcome Trust](#), rząd brytyjski, spółka inwestycyjna Deutsche Bank znana jako BRAAVOS oraz różne komunistyczne firmy chińskie, w tym Fosun Pharma.

„Wszyscy ci inwestorzy mogą czerpać zyski z tej „szczepionki” w najbliższej przyszłości, a Vaccitech był dość otwarty na temat przyszłego potencjału zysków ze swoimi akcjonariuszami, zauważając, że strzał przeciw COVID-19 najprawdopodobniej stanie się coroczną szczepionką aktualizowaną co sezon, podobnie jak szczepionka przeciw grypie sezonowej”, wyjaśnia dr Joseph Mercola.

„Oczywiście, AstraZeneca obiecała, że ☐nie przyniesie żadnych zysków z tej szczepionki przeciw COVID-19, ale ta obietnica jest ograniczona czasowo. Przysięga non-profit wygasa po zakończeniu pandemii, a sama AstraZeneca może zdecydować, kiedy to nastąpi”.

Globaliści chcą całkowitej kontroli nad Twoimi pieniędzmi i ciałem

Cenzurując wszystkie informacje i prawdę, które zagrażają atakowi COVID-19, Google chroni swoje udziały finansowe, a także swoich sojuszników i partnerów.

Google próbuje również zagłębić się w branżę opieki zdrowotnej, tworząc nowe programy „telemedycyny” i sztucznej inteligencji (AI), które firma ma nadzieję, że ostatecznie zostaną wdrożone jako zamiennik ludzkich lekarzy.

„Zaczęli na nowo wyobrażać sobie opiekę zdrowotną jako sposób na przejęcie kontroli nad życiem ludzi, mówiąc im, że jest to z korzyścią dla społeczeństwa, zbiorowości, a także ich osobistego zdrowia, podczas gdy tak naprawdę jest to sposób na wdrożenie tych transhumanistycznych lub technokratycznych technologii pod pozorem, że jest to przedsięwzięcie związane ze zdrowiem” – mówi Webb o programie.

Ostatecznym celem jest zastąpienie wszystkich tradycyjnych form medycyny i zastąpienie ich scentralizowanym, kontrolowanym przez globalistów, prowadzonym przez Big Tech systemem opartym na całkowitym zniewoleniu ludzkości.

Mówi się, że wojsko amerykańskie jest zaangażowane w operację, współpracując z Google, aby działać jako siła przy jej wdrażaniu – co oznacza obowiązkowe przejście na 'Covidiański Kult' pod groźbą użycia broni.

Wydaje się, że w fabule wstrzykiwania preparatu na wirusa jest element transhumanistyczny, co sugeruje, że coś w fiolkach kładzie podwaliny pod „znak bestii”, który przekształci ludzkie ciała stworzone na obraz Boga w ciała transhumanistyczne odtworzone na obraz Szatana.

„DARPA jest mocno zainwestowana w technologie transhumanistyczne do użytku w żołnierzach, w tym interfejsy mózg-maszyna i inne, jeszcze bardziej ekstremalne pomysły”, ostrzega Webb. „Ostatnio połączyli siły z Wellcome Trust, aby stworzyć coś, co nazywa się „Wellcome Leap”, raczej niepokojący ruch, który ma zapoczątkować transhumanizm”.

**Demaskatorzy z Facebooka
ujawnili dokumenty
szczegółowo opisujące wysiłki
na rzecz potajemnej cenzury**

dotyczącej szczepionek na skalę globalną



Otrzymaliśmy właśnie szereg nieoficjalnych dokumentów od informatora pracującego w Facebooku. Tym razem dokumenty te dotyczą szczegółów **planu odnośnie ograniczenia** – tu cytuję: **„treści wyrażających niezdecydowanie/niechęć wobec szczepień”** – na globalną skalę.

Sprawa ta okazała się na tyle niepokojąca, że zgłosił się do nas nie jeden, lecz dwóch sygnalistów/informatorów pracujących w Facebooku, którzy gotowi byli w wywiadzie z nami opowiedzieć, co to oznacza dla wolności słowa oraz dyskusji publicznej na tejże platformie.

Demaskatorzy z Facebooka o potajemnej cenzurze dotyczącej szczepionek na skalę globalną

– Facebook stosuje klasyfikatory w swoich algorytmach w celu oznaczenia określonych treści jako zawierających negatywne przesłanie względem narracji na korzyść szczepionek, czy też wyrażające „niechęć do szczepionek”. Bez wiedzy użytkowników, treści te mają przydzielane wartości punktowe, określane na podstawie skali zwanej „Punktacją VH [Niechęć do Szczepionek]”, metody klasyfikacji stopnia niechęci do szczepień. Na podstawie tej metryki dany komentarz jest usuwany lub pozostawiany w zależności od treści danego komentarza.

James O’Keefe: Nasz pierwszy informator pracuje jako technik w centrum danych. Przekazał nam szereg dokumentów wewnętrznych

opisujących program testowy nowego algorytmu, który przeprowadzono na 1,5 proc. użytkowników Facebooka i Instagrama, a których łączna liczba w skali świata wynosi niemal 3,8 mld użytkowników. Celem jest, tu cytuję: „*Drastyczna redukcja ekspozycji użytkowników na komentarze przedstawiające stanowisko niechętnie wobec szczepień*”.

Vaccine Hesitancy Comment Demotion

Credits

- Author List: @Joo Ho Yeo, @Nick Gibian, @Hendrick Townley, @Amit Bahl, @Matt Gilles
- Thanks:

Executive Summary

- What's your goal?
 - Drastically reduce user exposure to vaccine hesitancy (VH) in comments
- What is the product change?
 - Utilize the existing v1 VH classifier (English) to demote comments on ranked comments, meaning that they are filtered from 'most relevant' but are still visible in other tabs (ex 'most recent')

****Facebook Internal Document****

– To jest sytuacja podobna do kwestii kontrastu pomiędzy polityką publiczną, a polityką prywatności Facebooka, gdzie treść polityki publicznej jest bardzo niejednoznaczna, nieprecyzyjna. Jest zaprojektowana tak, by łatwo można było ją kwestionować. Stworzona jest w taki sposób, by łatwo można na jej podstawie wybronić swoje stanowisko. Natomiast jeśli spojrzeć na politykę prywatną, to jej treść jest znacznie bardziej przejrzysta, znacznie bardziej precyzyjna. Moim zdaniem chcą stworzyć wrażenie, że to tylko drobne zmiany, czy też generalizują cały problem mówiąc, że to tylko program, który oznacza odpowiednio posty dla ich wewnętrznego użytku.

Ja jednak uważam, że nie chcą, by ludzie dowiedzieli się, że tak naprawdę narzędzie to zostało stworzone z myślą o treściach dotyczących niechęci do szczepień. Opisuje to główny dokument plus załączniki i inne dodatkowe materiały. W uproszczeniu algorytm ten przeszukuje treści na Facebooku wyszukując konkretne słowa kluczowe związane ze szczepieniami

czy zniechęcającymi do szczepienia. Algorytm ten przydziela punkty w określonej skali, zwanej „Punktacją VH [Niechęci do Szczepionek]”, gdzie przypisuje się wartość liczbowa określającą niechęci do szczepionek, zdefiniowanej jako ocena stopnia niechęci do szczepień, ale nie dotyczy to takich przypadków opinii typu:

„Nie wiem, nie mam zdania”, ale raczej coś w stylu: „Widziałam/em badania pokazujące, że po otrzymaniu szczepionki osoba zmarła”. To jest właśnie definiowane jako „niechęć do szczepień”.

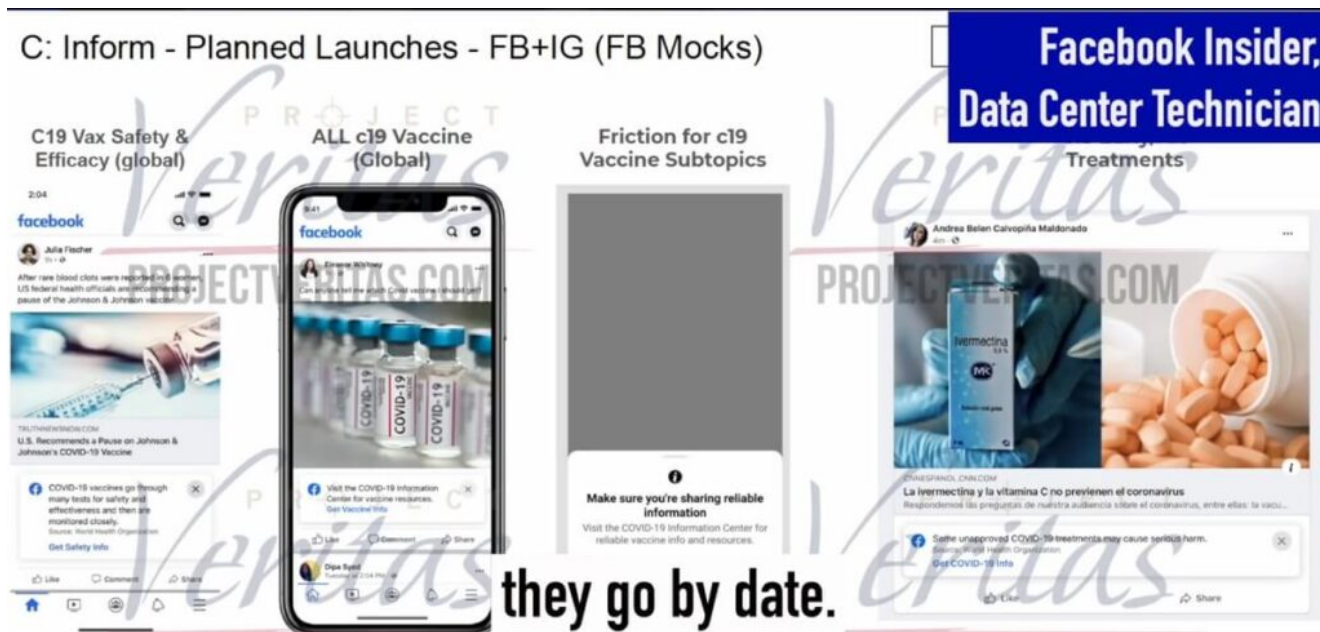
James O’Keefe: To są zatem testy w fazie wersji beta. Testy wersji beta algorytmu.

– Tak.

James O’Keefe: Jak duży zakres mają te testy w wersji beta? Wielkość próby w testach wynosi 1,5 proc. bazy użytkowników, nie wiem natomiast z jakiej puli treści jest dobierana ta próba, ale domyślam się, że są to komentarze na autorytatywnych profilach dotyczących zdrowia.

Zdaje się, że stawiasz tezę, że algorytm został już zaimplementowany. Co każe Ci sądzić, że algorytm ten już funkcjonuje?

– Na jednym ze slajdów mamy raport w postaci tzw. wskaźnika postępów w danym tygodniu, gdzie mamy wyniki przedstawione na wykresie. Kolejny slajdy uszeregowane są chronologicznie. Na początek mamy **treści dotyczące bezpieczeństwa i skuteczności szczepionek przeciwko COVID-19 z podziałem na kategorie: „globalnie” i „bieżące globalnie” w 13 językach, platformy Facebook oraz Instagram. Następnie wszystkie treści dotyczące COVID-19, globalnie i bieżące globalnie, 66 języków.** Pierwszą rzeczą, która skłoniła mnie do wniosku, że akcja ta będzie miała charakter globalny, to że algorytm ten był opracowywany dla jak największej liczby języków.



James O’Keefe: Czyli to coś na wzór kampanii startowej produktu?

– Tak. W ostatnim sprawozdaniu kwartalnym czytamy, że mają 2,79 mld użytkowników korzystających z jakiegoś rodzaju aplikacji facebookowej.

James O’Keefe: Kto stoi za tym nowym algorytmem? Cóż, jedno wiemy na pewno: autorzy tego eksperymentu – Joo Ho Yeo, Nick Gibian, Hendrick Townley, Amit Bahl oraz Matt Gilles to nie są pracownicy niższego szczebla. Jeden z dokumentów przekazanych Project Veritas zawierał diagram pokazujący, że osoby te od samego Marka Zuckerberga dzieli zaledwie kilka stopni w hierarchii. Kto decyduje o tego typu kwestiach?

– Są to zespoły wewnątrz struktur Facebooka, które nazywają siebie czy też pracują nad kwestiami określanymi mianem **B2V**, czyli związanymi z „barierami dla szczepień”. (Program ramowy „Borderline Vaccine (BV)”) To jest jedno z ich głównych haseł – **bariery dla szczepień**. Także osoby pracujące w tzw. zespołach ds. integralności zdrowotnej. Aplikacja Messenger posiada taki dedykowany zespół, każda aplikacja Facebooka posiada swój dedykowany zespół tego typu.

James O’Keefe: Czy te zespoły od integralności zdrowotnej zostały utworzone po 22 marca 2020, po początku epidemii

COVID-19?

– Nie.

James O’Keefe: Czyli istniały one już wcześniej?

– Tak, zgadza się.

Vaccine Hesitancy Comment Demotion

Credits

- Author List: @Joo Ho Yeo, @Nick Gibian, @Hendrick Townley, @Amit Bahl, @Matt Gilles
- Thanks:

Executive Summary

- What's your goal?
 - Drastically reduce user exposure to vaccine hesitancy (VH) in comments
- What is the product change?
 - Utilize the existing v1 VH classifier (English) to demote comments on ranked comments, meaning that they are filtered from 'most relevant' but are still visible in other tabs (ex 'most recent')
- What are the benefits of this launch?
 - VPVs on Vaccine post English comments vh p80: **-10.6±2.1%**
 - Projected launch impact: -934.8K±194.4K vpvvs
 - Authoritative vh p80 comment vpvvs: -26.7 (±4.1)%
 - Projected launch impact: -402.4K±71.3K
 - CEP on Vaccine post English comments vh p80: **-11.1 (±1.8)%**
 - Authoritative vh p80 comments CEP: **-26.1±3.0%**
 - Decrease in other engagement of VH comments including create, likes, reports, replies
 - ↓ Scuba grouped by VH
- What are the costs of this launch?
 - No significant cost is observed.
- Risks of this launch
 - Not all comments are actually vaccine hesitancy, but we'd aligned with Health Policy on this risk in the COVID Lockdown Decisions meeting 2 weeks ago – https://docs.google.com/presentation/d/1Qo35TGq75yf70-VkOAY61g0oFfjYaOB2o2dF6SUry44/edit#slide=id.gca2fb195a7_11_0
- How could this be made more aggressive?
 - Use lower thresholds for interventions
- How could this be made more conservative?
 - Use higher Thresholds for interventions

Background

Experiment Launch [Post](#)

Comments are a major surface relevant to our B2V efforts. We estimate that the prevalence of VH comments in Authoritative Health Pages is 25.3% and for other pages 19.42%. Now that the v1 Vaccine Hesitancy classifier has been cleared for this usecase, reducing the visibility of these comments represents another significant opportunity for us to remove barriers to vaccination that users on the platform may potentially encounter.

James O'Keefe: Tak więc w Facebooku funkcjonują zespoły ds. integralności zdrowotnej oraz zespoły zajmujące się „barierami dla szczepień”.

– Osoba ta, **Amit Bahl**, który jest naukowcem z głównym zespole ds. analizy danych Facebooka również udziela się w zespołach ds. integralności zdrowotnej. Każda osoba zaangażowana w ten projekt także jest członkiem takiego zespołu. **Hendrick Townley** – tutaj mamy **post dotyczący startu tego eksperymentu, z 16 kwietnia, z 2021 roku**, to dlatego data postu nie zawiera roku, gdyż post pochodzi z bieżącego roku. Jest on jednym z inżynierów oprogramowania pracujących nad tym projektem. Jeśli dobrze pamiętam, pracuje on jeszcze z dwoma innymi inżynierami. Amit Bahl jest autorem algorytmu klasyfikującego, pewnie miał do pomocy zespół inżynierów, ale to on jest uznawany za autora. Łańcuszek idzie od Amita Bahl poprzez jego przełożonego, przełożonego jego przełożonego, przełożonego jego przełożonego, itd.

James O'Keefe: Amit Bahl pracuje dla Udiego Weisenberga, który z kolei jest podwładnym Nicolasa Stiera, zgadza się?

– Tak.

James O'Keefe: Stier natomiast jest podwładnym **Danny'ego Ferrante**, który z kolei jest podwładnym Javiera Oliviana, a **Javier Olivian** odpowiada bezpośrednio przed Markiem Zuckerbergiem.

– Tak.

James O'Keefe: Ten łańcuszek prowadzi od samych szczytów do Amita Bahla. **Tego typu cenzura budzi skojarzenia z praktyką „cichego banowania” (ang. shadow banning), kiedy użytkownicy poddani takiej praktyce nie są świadomi, że zamieszczane przez nich posty nie są widoczne dla innych użytkowników oprócz nich samych.** Termin ten został spopularyzowany przez jedno z naszych śledztw odnośnie Twittera, kilka lat temu. Tym razem jednak mamy do czynienia z praktykami Facebooka i Instagrama

mającymi na celu wyeliminowanie jednej ze strony dyskusji publicznej toczonej na temat być może jednego z najważniejszych obecnie problemów. Z tych dokumentów wynika, że **Facebook klasyfikuje treści dzieląc je na klasy oznaczone liczbami od 1 do 5 oraz przydziela im punktację w skali od 0 do 1.** Możesz powiedzieć nieco więcej o tym systemie klasyfikacji?

– W uproszczeniu **jest to metoda klasyfikacji i strukturyzacji treści na bazie tego, do jakiego stopnia dana treść wyraża niechęć wobec szczepionek czy też otwarcie zniechęca do takich szczepień.** Sądzę, że głównym powodem stworzenia tego systemu była chęć wsparcia tzw. osób oceniających, aby osoby te mogły skorzystać z tej klasyfikacji jako źródła odniesienia w sytuacji, **gdy algorytm nie jest w stanie samodzielnie ocenić danego posta i odsyła go do takiej właśnie osoby oceniającej.** Osoba taka wtedy podejmuje decyzję czy dany post ma charakter zniechęcający do szczepień, czy też nie.

James O’Keefe: Na przykład „Klasa 2” – pośrednie zniechęcanie. Z jakimi konsekwencjami wiąże się zakwalifikowanie danej treści do 2. klasy?

– Oznacz to, że twój komentarz zostanie w gruncie rzeczy ukryty. Jeśli chodzi o zakres tej cenzury to ciężko powiedzieć, to jest ustalane indywidualnie dla każdego przypadku. Takie posty zostaną poddane procedurze „zmiany pozycjonowania”. **W przypadku niektórych postów widzimy określoną adnotację, np. zmiana pozycji: 7.** Nie mam pewności, czy ta wartość 7 oznacza 7 proc. czy też przesunięcie postu o 7 pozycji w dół listy komentarzy.

James O’Keefe: Czyli post jest tłumiony/ukrywany?

– Tak.

James O’Keefe: Ale nie znamy zakresu ukrywania takich treści? Słyszeliśmy o praktykach w postaci „shadow banów”, „deboostingu” (celowe obniżanie widoczności treści/postów),

ale pierwszy raz słyszę o „zmianie pozycjonowania”. Facebook wcale nie kryje się z tym, że walczy z niechęcią do szczepionek, o czym zresztą jasno mówią w oficjalnym komunikacie prasowym z 11 maja tego roku. Jednakże to, co mówią w tej sprawie publicznie to zaledwie czubek góry lodowej tego, co możemy wyczytać w ujawnionych dokumentach wewnętrznych, przekazanych nam przez sygnalistę z Facebooka.

– Jeśli się nie mylę, Celem jest, tu cytat: **„Drastyczna redukcja ekspozycji użytkowników na komentarze przedstawiające stanowisko niechętnie wobec szczepień”**. Jak Ty to widzisz?

– Ramy programu „Borderline Vaccine (BV)”, podpunkt „szokujące informacje”, czytamy: **„potencjalnie prawdziwe lub prawdziwe zdarzenie, które może wzbudzić niepokój odnośnie bezpieczeństwa...”**

James O’Keefe: Ten oto post napisany został przez faktycznego użytkownika, zrobiono rzut ekranu i poddano pod ocenę systemu klasyfikującego.

– Tak.

James O’Keefe: Co takiego niebezpiecznego w tym stwierdzono?

– Cytat: **„I jak się tutaj nie bać?”**. Taka treść zostałaby zaklasyfikowane jako **„pośrednie zniechęcanie”**.

James O’Keefe: Pośrednie zniechęcanie, czyli spełnia kryterium przynależności do 2. klasy?

– Tak.

James O’Keefe: To jest przykład na zniechęcanie do szczepień.

– Tak. Kwestionowanie szczepionek, kwestionowanie wszystkiego co dotyczy szczepionki.

James O’Keefe: Ten post został „zdegradowany”, że tak powiem...

– Tak.

James O'Keefe: ...na Facebooku. Oto kolejny przykład treści zaklasyfikowanej do 2.klasy, czyli „pośredniego zniechęcania”. Powiedz nam proszę, dlaczego wg. Facebooka spełnia on kryteria „pośredniego zniechęcania”?

– Zostałoby to zaklasyfikowane jako zniechęcanie do szczepień ponieważ oni nie chcą, by ludzie byli eksponowani na historie przypadków negatywnych skutków. W tym poście nie nakłania się do nie przyjmowania szczepionki, lecz opisuje się negatywne skutki po przyjęciu szczepionki.

James O'Keefe: Opisuje się fakty, ale prawda najwyraźniej nie ma znaczenia.

– Nie ma znaczenia, ponieważ nie pasuje do ich narracji, a narracja jest taka oto, żeby przyjmować szczepionki, szczepionka jest dla ciebie dobra, każdy powinien ją przyjąć. Jeśli tego nie zrobisz, to czeka cię ostracyzm.

James O'Keefe: Zgadza się.

– Staniesz się wrogiem społeczeństwa. VAERS, czyli system raportowania groźnych odczynów poszczepiennych – wygląda na to, że śledzą komentarze, gdzie w treści są wzmianki o śmierci pacjenta.

James O'Keefe: Czy system VAERS jest tworzony przy współudziale CDC?

– Tak, tak mi się wydaje. CDC wspiera te działania, gdyż sami opublikowali wytyczne w tym aspekcie. To jest jeden z głównych filarów, na których Facebook opiera swoje działania w tym względzie. Jeśli spojrzeć na ich politykę publiczną, to jest ona zgodna z treściami na autorytatywnych witrynach związanych ze zdrowiem.

James O'Keefe: Czyli nie kryją się z faktem, że działają w ścisłej koordynacji z CDC?

– Tak, oczywiście.

James O'Keefe: W dokumentacji czytamy, że „*treści odsyłające do danych z systemu VAERS sugerujących znaczne ryzyko bez podawana pełnego kontekstu*”. „Pełny kontekst”. Co się przez to rozumie?

– Sądzę, że chodzi tu oto, jak w jednym z poprzednich komentarzy, gdzie autor/ka pisze, że w ramach cytowanego badania zmarły 653 osoby, ale bez podania całej treści artykułu. Czyli co? Życzycie sobie, abym napisał 10-stronicowy komentarz?

James O'Keefe: Tak.

– Co to znaczy „pełny kontekst”? Kolejny mglisty termin. Jeśli zacytujesz jakieś badanie przeprowadzone przez osobę, za którą nie przepadają, czy powiedzą, że to nie jest pełny kontekst sprawy, albo to czy tamto?

[Rafał Brzeski: Obrona przed dezinformacją – spójny system weryfikacji](#)

[Obalacze obaleni: Kto weryfikuje weryfikatorów faktów?](#)

[Weryfikowanie weryfikatorów faktów – Historia Snopes według Daily Mail](#)

James O'Keefe: Ostatecznie wygląda na to, że jakiegokolwiek fakty, które nie pasują do konkretnej narracji, są pomijane, „degradowane”, „deboostowane”, „banowane”...

– Tak.

James O'Keefe: ...traktowane jako niebezpieczeństwo dla społeczeństwa?

– Ukrywane na każdy możliwy sposób.

James O'Keefe: Dokąd to zmierza? Czy będzie tylko coraz gorzej?

– Tak.

James O'Keefe: Wyjaśnij proszę.

– Sądzę, że ludzie to akceptują i widząc to, ludzie z Twittera, Facebooka czy Google, itd. dojdą do wniosku, że jest przyzwolenia na taki stan rzeczy i potraktują to jak przysłowiowe „zielone światło” i tylko nasilą te działania. Te działania zaczną dotyczyć już nie tylko szczepionek, ale wręcz wszystkich treści. Zaczną mówić, że jeśli napiszesz jakiś post, który może spowodować zagrożenie dla czyjegoś życia czy zdrowia i bezpieczeństwa – cokolwiek to znaczy – to my przyjrzymy się temu postowi i na podstawie jakiegoś arbitralnego kryterium przydzielimy mu określoną wartość punktową. Ich celem jest kontrola treści zanim te trafią na twoją stronę profilową [ekran], zanim je nawet zobaczysz przeglądając posty na telefonie. Wydaje mi się, że boją się wniosków, które ludzie mogą z tych treści wyciągnąć.

James O'Keefe: Oni nie myślą o wartości tego...

– Nie. I że to może wywołać efekt domina.

James O'Keefe: Nie biorą tego nawet po rozważę?

– Nie.

James O'Keefe: Czy wydaje Ci się, że w Facebooku pracuje spora grupa osób, które zgadzają się z Twoją opinią...

– Tak.

James O'Keefe: Ile szacujesz jest takich osób?

– Sądzę, że co najmniej 25 proc. pracowników.

James O'Keefe: Mówisz, że 25 proc. pracowników Facebooka jest zgodnych, że cała sprawa z „niechęcią wobec szczepionek”, ten algorytm czy kod, jest moralnie naganna?

– Tak, nie zdziwiłby się, gdyby tak właśnie to wyglądało. Ta polityka będzie tylko rozszerzać swój zakres do momentu, gdy

praktycznie wszystko będzie mogło stać się pretekstem jej naruszeniem. Poszerzenie skali implementacji tej polityki, rozszerzenie na Twittera i całą Sieć, będzie znacznie gorsze w skutkach, niż moja utrata pracy. To jest znacznie ważniejsza sprawa. Tu nie chodzi o mnie, ale o wszystkich ludzi na świecie.

James O’Keefe: Jasnym jest, że nasz informator z Facebooka nie jest sam jeśli chodzi o zaniepokojenie odnośnie tej nowej polityki Facebooka, która istotnie zuboży dyskusję publiczną dotyczący szczepionek jako takich. W rzeczy samej, nasz kolejny informator, którego Państwu zaraz przedstawimy, tak bardzo poczuwał się w obowiązku, by ujawnić te informacje, że jest skłonny porównać tę nową politykę do toksycznego związku.

Czym zajmujesz się w Facebooku?

– Jestem inżynierem w centrum danych. W pewnym sensie polityka ta ogranicza czy uniemożliwia prowadzenie otwartej dyskusji przez ludzi na tematy bezpośrednio dotyczące ich własnego bezpieczeństwa. Niektóre z tych rzeczy, **pytania, komentarze, dotyczą tematów związanych z obawami odnośnie własnego zdrowia czy bezpieczeństwa. To trochę tak jak bycie w związku z osobą stosującą przemoc, despotyczną, jak partner stosujący przemoc, który zabrania osobie krzywdzonej głośno mówić o problemach w związku czy małżeństwie, odbierając jej prawo głosu.** Skutki tego są bardzo negatywne. Z punktu widzenia logiki jest to sprawa w mojej ocenie godna potępienia. Ten sam informator kilka miesięcy temu udostępnił nam nagranie, prywatnej rozmowy Marka Zuckerberga ze swoim zespołem z lipca, na którym on sam mówi o własnym sceptycyzmie wobec szczepień w kontekście COVID-19.

(...rozszerzamy zakres naszych działań mających na celu usuwanie z Facebooka oraz Instagrama fałszywych twierdzeń odnośnie COVID-19, szczepionek przeciwko COVID-19 i szczepionek jako takich...)

Mark Zuckergeberg: „Chcę zapewnić, że sam podchodzę to tego z ostrożnością, ponieważ nie znamy długofalowych skutków ubocznych modyfikowania ludzkiego DNA i RNA,,.

James O’Keefe: Obecnie te jego słowa naruszyłyby zasady polityki jego własnej firmy. To Ty nam udostępniłeś to szczepionkowe nagranie Zuckerberga z 22 lipca 2020. Zgadza się?

– Zgadza się.

James O’Keefe: Mark Zuckerberg zdaje się zmienić poglądy na temat szczepionek, a obecna polityka Facebooka w tych kwestiach uniemożliwiłaby jemu samemu wypowiedzenie tych słów sprzed kilku miesięcy. Istne szaleństwo!

– Tak.

James O’Keefe: Co skłoniło Cię do przekazania tych materiałów nam, dla [Project Veritas](#)?

– Opierając się na własnej wiedzy i analizach odnośnie odczynów poszczepiennych oraz efektach ubocznych szczepionek i leków w ogóle. Sądzę, że ważnym jest, by ludzie mieli wiedzę o tych faktach. W pracy możemy bez przeszkód o tych rzeczach dyskutować ze współpracownikami, o naszych wyborach zdrowotnych. I robimy to często. **Dlaczego takie dyskusje możemy prowadzić w pracy jak osoby dorosłe bez pouczania nas przez kogoś, że powinniśmy przestać o tym rozmawiać, nie powinniśmy o tym rozmawiać, że zostało to zweryfikowane jako fałszywe.** Czasami jedna czy druga osoba w ramach dyskusji coś takiego powie, ale w ramach otwartej dyskusji. A w tym przypadku...

James O’Keefe: Czyli mówisz, że tego typu luźne rozmowy biurowe to u was nic niezwykłego?

– Tak i to nie tylko na przerwach.

James O’Keefe: W pracy, w Facebooku w fizycznej przestrzeni biur Facebooka owszem, ale w przestrzeni wirtualnej serwisu Facebook jest to już zabronione?

– Zgadza się.

James O’Keefe: Cóż za ironia.

– To bardzo ironiczne.

James O’Keefe: Masz jakieś przesłanie dla dyrektora generalnego Facebooka, Marka Zuckerberga, który z pewnością ogląda ten materiał?

– Zapytałbym... Powiedziałaś, że wspierasz wolność ekspresji. Moje pytanie jest następujące: dlaczego nie wspierasz wolności ekspresji w jej fundamentalnym sensie, która jest przymiotem każdego człowieka, czyli wolność człowieka takiego, jakim stworzył go Bóg, bez wpływu czy ograniczeń ze strony korporacji czy rządów. **Dlaczego zatem propagować narrację, że czymś godnym potępienia jest to, że ludzie żyją i wyrażają siebie jako stworzenia boskie i nie życzą sobie szczepionki? Skąd ta presja, by taką ideologię czy przekonania eliminować z dyskusji?**

James O’Keefe: Największą obawą tegoż informatora jest **plan budowy przez Facebooka – tu cytuję: „społeczności podporządkowującej się”** – i uważa, że takie działanie wykracza poza kompetencje samego Zuckerberga.

– **Chcą taką platformę stworzyć, zbudować taką właśnie społeczność. Pragną stworzyć społeczność podporządkowanych użytkowników, a nie społeczność, której członkowie mogą swobodnie rozmawiać o nawet najbardziej prywatnych, osobistych czy intymnych decyzjach w sytuacjach, z którymi każdy ma do czynienia na co dzień, związanych z własnym ciałem, zdrowiem.**

James O’Keefe: W oparciu o to, co widziałaś, charakterze Twojej pracy w Facebooku, jak sądzisz, dlaczego zależy im na

takim posłuszeństwie? Dlaczego na Facebooku panuje taka, a nie inna kultura?

– W celu promowania społeczeństwa, które skupia się na w pewnym sensie rozwoju medycyny, czy też raczej stawia na eksperymentowanie o charakterze medycznym.

James O’Keefe: Jakie działania podejmuje Facebook w celu promowania szczepień przeciwko COVID-19 wśród miliardów swoich użytkowników?

– Przede wszystkim zachęca pisanem o tym, by zmieniać wygląd ramki profilowej mającej pokazywać, że się zaszczepili, promując w pewnym sensie solidarność z innymi osobami, które też się zaszczepiły i zachęcać tym samym innych do zaszczepienia się. To jest interesujące, że jedna strona dyskusji może wypowiadać się swobodnie, ale drugiej stronie, która z jakiegoś powodu nie zamierza się szczepić, tę możliwość się odbiera.

W tym drugim przypadku nie ma mowy o solidarności, zjednoczeniu wokół tego tematu. To się sprowadza do tego, że albo przyjmujesz szczepionkę, albo jesteś wykluczony. Ludziom odbiera się prawo do własnej opinii, choć sam Facebook chwali sobie swoje zaangażowanie właśnie w promowanie otwartej wymiany opinii. Mówiłem już o tym wcześniej – posiadanie prawa do zabierania głosu w sprawach najistotniejszych w naszym życiu to bardzo istotna rzecz. To szokująca i niepokojąca sytuacja, że niektóre osoby gotowe są posunąć się tak daleko, że są zdolne ograniczać na całym świecie dyskusje dotyczące spraw związanych z bezpieczeństwem. Sądzę, że być może jakieś zewnętrzne siły wywierają presję na Marka, nie twierdzę, że tak jest w istocie, niemniej jest to niepokojąca hipoteza. Wystarczy spojrzeć na nagranie z Faucim, na którym Mark omawiał ten temat. Mark nie wydaje się zadowolony z faktu, że ten temat musiał zostać poruszony. Wygląda to tak, jakby został zmuszony to poruszenia tego tematu. Tych dwoje odważnych informatorów

zdaje sobie sprawę, że występując przeciwko Facebookowi ryzykowali wszystkim przeciwko najpotężniejszej obecnie instytucji. Pomimo tego wciąż to ryzyko ponoszą w walce o obronę naszych wolności.

– Jest to jak sądzę jeden z głównych powodów, dla których ludzie boją się iść z tym do mediów. Załóżmy, że mam żonę i dwójkę dzieci, jeśli stracę pracę, to co z nami się stanie? Ja nie mam takiego dylematu.

James O’Keefe: A osoby, które potencjalnie pójda w twoje ślady, ale nie są jeszcze do końca przekonane, że powinny tak postąpić, jakie masz przesłanie dla tych osób?

– Po prostu zróbcie to. Pomyślcie o tym, co powiedziałem przed chwilą. Owszem, mówienie prawdy wiąże się z reperkusjami, jak pokazuje zresztą historia ludzkości. Owszem, czekają Was reperkusje, wyważcie argumenty za i przeciw. **Czy będzie w stanie żyć w spokoju z myślą, że wiedzieliście o tych praktykach Facebooka i nic nie zrobiliście?** Zwłaszcza, gdy po jakimś czasie sprawa wyjdzie na jaw i pomyślicie, że mieliście z tym do czynienia mnóstwo razy i teraz żałujecie, że nic z tym nie zrobiliście. Jak będziecie się z tym czuć? Ja postąpiłem zgodnie z moim sumieniem, taką już jestem osobą. Musiałem o tym komuś powiedzieć, ponieważ nawet jeśli ich plan nie doszedłby do skutku i został spisany na straty, to i tak będę uważał, że było warto, że próbowałem komuś to pokazać.

James O’Keefe: Jak widać, Twoja potrzeba podążania za głosem sumienia jest silniejsza niż potencjalne ryzyko z tym związane?.

– Tak.

James O’Keefe: A jakie masz przesłanie do osób pracujących w Facebooku, które zgadzają się z Tobą?

– Nie jesteście sami. Media społecznościowe, czego doświadczaliśmy przez te kilka lat, ocenianie komentarzy,

postów, czy nawet mediów z wiadomościami, że kreuje się wrażenie, że jesteś w mniejszości w swoich poglądach, a nie w większości. Nawet osoby, które twierdzą, że zgadzają się z czymś, wewnętrznie są skonfliktowane, kwestionują te rzeczy, ale nie są w stanie pokonać tego wewnętrznego dysonansu poznawczego i ostatecznie wybierają bezpieczną opcję. Łatwiej jest wybrać konformizm, niż zdobyć się na odwagę i powiedzieć: „To jest złe”, czy uświadomić sobie i zrozumieć, że to jak ty i inne osoby jesteście traktowani, jest złe. To nie jest łatwe, ale wiedźcie, że nie jesteście sami. Jeśli odważycie się na ten krok, zawalczycie o sprawę, to będziecie do końca życia mieć spokój ducha, że postąpiliście słusznie. To daje odwagę, nie tylko innym osobom, ale i wam samym. Nie będziecie się czuli dłużej samotnie.

James O’Keefe: To bardzo głębokie słowa.

– Jestem ojcem i boję się o bezpieczeństwo moich dzieci, mojej rodziny. Jestem zarazem świadomy, co zresztą powiedziałem swojej żonie, że muszę coś zrobić. Że nadchodzi taki czas, że trzeba walczyć o świat, w którym będą żyć nasze dzieci. **Nie chcę, by żyły w świecie bez tych praw, bez osobistych wolności. By żyły na świecie, gdzie nie mogą podejmować własnych decyzji dotyczące zdrowia, w relacji pacjent-lekarz oraz z najbliższymi osobami.** Nie chcę by żyły w takim świecie, nie mogę się na to zgodzić. Posunęło się to do tego etapu, że nie mogłem siedzieć bezczynnie i musiałem zrobić to, co słuszne.

Ekipa Project Veritas skontaktowała się z Facebookiem w celu udzielenia komentarza odnośnie szczegółów ich systemu klasyfikacji i degradacji treści użytkowników.

Maria Balaba, Project Veritas: Mario z Project Veritas z tej strony. Wysłałem do Państwa emaila, kilka minut temu, nie wiem, czy już zapoznaliście się z jego treścią. W mailu prosiłem o komentarz – otrzymaliśmy dokumenty wewnętrzne Facebooka, dotyczące niejakiej „niechęci do szczepień”. Są to

dokumenty do użytku wewnętrznego, wygląda na to, że mają one klauzulę poufności. Na to wygląda przynajmniej. W załączniku do wiadomości umieściliśmy kopie tych dokumentów i prosilibyśmy o komentarz w tej sprawie.

– Ok, zapoznam się z tą wiadomością.

– Dziękuję. Kiedy możemy się spodziewać odpowiedzi?

– Zaraz sprawdzę tę wiadomość.

– Świetnie, czekamy zatem na odpowiedź.

– Ok, dziękuję.

– Również dziękuję.

(Temat: Prośba o komentarz – wyciek tajnych dokumentów Facebooka

„Niezwłocznie ogłosiliśmy treść tej polityki na oficjalnym blogu i aktualizowaliśmy centrum pomocy o te informacje – rzecznik Facebooka)

Facebook nie zareagował na szereg pytań ze strony Project Veritas, dlatego poprosiliśmy o wskazanie miejsca zamieszczenia tych rzekomych aktualizacji.

Nikt z Facebooka nie odpowiedział na zapytanie.

(Email uzupełniający wysłany przez Project Veritas

Czy mogą Państwo wskazać stronę, na której treść tej polityki jest publicznie dostępna? Pomimo starań, na żadnej ze stron treści tej nie udało nam się zlokalizować, więc jeśli istotnie jest ona dostępna, proszę o wskazanie miejsca i daty jej publikacji. Zwracamy się także z prośbą o wskazanie miejsca, w którym zdefiniowany jest termin „degradacja komentarza”).

W momencie publikacji tego materiału, praktyka degradowania komentarzy o charakterze „niechęci wobec szczepień” nie jest

opisana nigdzie na oficjalnym blogu czy w dziale Centrum pomocy.

Odwagi. Zrób coś. Idź śladem tych ludzi. Skontaktuj się z nami pod adresem: veritastips@protonmail.com.

Demaskatorzy z Facebooka o potajemnej cenzurze dotyczącej szczepionek na skalę globalną [napisy PL]

Źródło: pubmedinfor.org