

Chińskie procesory graficzne niemal dziesięciokrotnie przewyższają układy Nvidii w symulacjach superkomputerowych: Zmiana w suwerenności technologicznej



W przełomowym odkryciu, które może zmienić globalny krajobraz półprzewodników, chińscy naukowcy osiągnęli niemal dziesięciokrotny wzrost wydajności w symulacjach superkomputerowych przy użyciu krajowych procesorów graficznych (GPU), przewyższając systemy zasilane najnowocześniejszym sprzętem firmy Nvidia. Ten kamień milowy, szczegółowo opisany w recenzowanym badaniu opublikowanym w Chinese Journal of Hydraulic Engineering, podkreśla rosnącą sprawność Chin w zakresie wysokowydajnych obliczeń (HPC) i ich determinację w ograniczaniu zależności od zagranicznych technologii.

Osiągnięcie to pojawia się w kluczowym momencie globalnego wyścigu technologicznego, ponieważ eskalacja amerykańskich sankcji na zaawansowane półprzewodniki zmusiła Chiny do przyspieszenia wysiłków na rzecz opracowania własnych alternatyw. Podczas gdy sceptycy ostrzegają, że same optymalizacje oprogramowania nie mogą w nieskończoność wypełniać luk sprzętowych, badanie podkreśla, w jaki sposób

innowacyjne projekty obliczeń równoległych i poprawki oprogramowania mogą odblokować bezprecedensowy wzrost wydajności, nawet przy mniej zaawansowanym sprzęcie.

Przełom w obliczeniach równoległych

Badania, prowadzone przez profesora Nan Tongchao z Państwowego Kluczowego Laboratorium Hydrologii Zasobów Wodnych i Inżynierii Wodnej Uniwersytetu Hohai, koncentrowały się na podejściu do obliczeń równoległych „wiele węzłów, wiele procesorów graficznych”. Wykorzystując produkowane w kraju procesory CPU i GPU, zespół osiągnął znaczną poprawę wydajności w symulacjach na dużą skalę w wysokiej rozdzielczości – co ma kluczowe znaczenie dla takich zastosowań jak modelowanie obrony przeciwpowodziowej i zapobieganie podtopieniom w miastach.

„Wyzwanie dla chińskich naukowców jest jeszcze trudniejsze”, zauważono w badaniu, wskazując na dominację zagranicznych producentów w produkcji zaawansowanych procesorów graficznych, takich jak A100 i H100 firmy Nvidia. Sprawę pogarsza również zastrzeżony ekosystem oprogramowania CUDA firmy Nvidia, który nie może być uruchamiany na sprzęcie innych firm, skutecznie uniemożliwiając chińskim programistom dostęp do kluczowych narzędzi do opracowywania algorytmów.

Pomimo tych przeszkód, zespół profesora Nana wykazał, że techniki optymalizacji oprogramowania mogą znacznie zwiększyć wydajność chińskich procesorów graficznych, umożliwiając im prześcignięcie amerykańskich superkomputerów w określonych obliczeniach naukowych. Przełom ten nie tylko podważa dominację firmy Nvidia, ale także podkreśla potencjał alternatywnych podejść do obliczeń o wysokiej wydajności.

Szersze implikacje sankcji

technologicznych

Wyniki badania podkreślają niezamierzone konsekwencje amerykańskich sankcji technologicznych, które miały na celu ograniczenie dostępu Chin do zaawansowanych półprzewodników i krytycznych technologii. Wydaje się, że zamiast tłumić innowacje, ograniczenia te ożywiły wysiłki Chin na rzecz osiągnięcia samowystarczalności technologicznej.

„Osiągnięcie to wskazuje na możliwe niezamierzone konsekwencje eskalacji sankcji technologicznych Waszyngtonu, jednocześnie podważając dominację amerykańskich chipów, od dawna uważanych za kluczowe dla zaawansowanych badań naukowych” – czytamy w badaniu.

Rozwój ten jest zgodny z szerszą strategią Pekinu mającą na celu złagodzenie ryzyka „punktu krytycznego” w krytycznych technologiach – termin odnoszący się do słabych punktów w łańcuchach dostaw, które mogą zostać wykorzystane przez przeciwników geopolitycznych. Inwestując znaczne środki w krajową produkcję półprzewodników i ekosystemy oprogramowania, Chiny dążą do zmniejszenia swojej zależności od zagranicznego sprzętu i oprogramowania, zapewniając sobie odporność technologiczną na coraz bardziej rozdrobnionym rynku globalnym.

Kontekst historyczny: Od zależności do innowacji

Znaczenie tego osiągnięcia jest nie do przecenienia w kontekście trwającej od dziesięcioleci walki Chin o dogonienie zachodniej technologii półprzewodników. W przeszłości Chiny w dużym stopniu polegały na imporcie zaawansowanych chipów, a na rynku dominowały amerykańskie firmy, takie jak Nvidia, Intel i AMD. Zależność ta stała się rażącą słabością wraz z eskalacją napięć geopolitycznych, co skłoniło Stany Zjednoczone do

nałożenia szeroko zakrojonych kontroli eksportu zaawansowanych półprzewodników i sprzętu do produkcji chipów.

W odpowiedzi Chiny zwiększyły inwestycje w swój krajowy przemysł półprzewodników, dzięki inicjatywom takim jak „Big Fund”, przeznaczając miliardy dolarów na badania i rozwój. Chociaż wyzwania pozostają – szczególnie w zakresie produkcji chipów w najnowocześniejszych 3-nanometrowych i niższych węzłach – ten najnowszy przełom pokazuje, że Chiny robią znaczące postępy w wykorzystywaniu istniejącego sprzętu poprzez innowacje w oprogramowaniu.

Co nas czeka?

Choć wyniki badania są imponujące, eksperci ostrzegają, że same optymalizacje oprogramowania nie są w stanie w pełni zrekompensować ograniczeń sprzętowych. Przykładowo, układy GPU firmy Nvidia słyną nie tylko z surowej mocy obliczeniowej, ale także z wszechstronności i integracji z solidnym ekosystemem oprogramowania. Powtórzenie tego poziomu wydajności w szerokim zakresie aplikacji będzie wymagało ciągłego rozwoju zarówno sprzętu, jak i oprogramowania.

Niemniej jednak badanie to stanowi znaczący krok naprzód w dążeniu Chin do suwerenności technologicznej. Ponieważ zespół profesora Nana nadal udoskonala swoje podejście do obliczeń równoległych, implikacje dla takich dziedzin jak modelowanie klimatu, sztuczna inteligencja i bezpieczeństwo narodowe mogą być głębokie.

Zdaniem jednego z obserwatorów branży, „jest to sygnał alarmowy dla globalnej społeczności technologicznej. Chiny udowadniają, że potrafią wprowadzać innowacje pod presją, a reszta świata będzie musiała dostosować się do tej nowej rzeczywistości”.

W miarę nasilania się wojny technologicznej między Stanami Zjednoczonymi a Chinami, badanie to służy jako przypomnienie,

że innowacje często rozwijają się w obliczu przeciwności losu. Czy ten przełom doprowadzi do szerszej zmiany w równowadze sił technologicznych, dopiero się okaże, ale jedno jest pewne: wyścig o dominację półprzewodników jest daleki od zakończenia.

8 największych zalet i wad sztucznej inteligencji



Sztuczna inteligencja przeżywała boom w ciągu ostatnich kilku lat i jest to dobre i złe z wielu powodów. Przyszłość nie jest tym, czym „była kiedyś” i to jest pewne. Kto mógł sobie to wyobrazić? Czy dojdzie do przejęcia „Sky-Net”, jak w filmie „Terminator”? Czy sztuczna inteligencja będzie rewolucyjna i pomoże miliardom ludzi żyć bezpieczniejszej, zdrowiej i wydajniej? Oto 8 zalet i wad tej oszałamiającej ery technologicznej, w której wszyscy teraz żyjemy.

1. Sztuczna inteligencja (AI) może pomóc w badaniach i pisaniu, ale może również tworzyć fałszywe wiadomości i krytyczne błędy, w tym nielogiczne wnioski i halucynacje.
2. Sztuczna inteligencja może (pomóc) kontrolować drony, pojazdy i broń wojskową, ale może zostać zhakowana lub przechytrzyć użytkowników i zwrócić się przeciwko ludzkości.
3. Sztuczna inteligencja może tworzyć obrazy i filmy, łącząc informacje i wizualizacje w celu pobudzenia wyobraźni i w

celach rozrywkowych, ale może też oszukać ludzi, by uwierzyli w fałszywe koncepcje, takie jak inwazje kosmitów lub przemówienia polityczne, które nie miały miejsca.

4. Sztuczna inteligencja może być wykorzystywana do wykonywania wielu zadań wydajniej i spójniej niż ludzie, takich jak praca w fabrykach, ale może to oznaczać koniec milionów miejsc pracy, które nie wymagają krytycznego myślenia, kreatywności ani interakcji międzyludzkich.

5. Sztuczna inteligencja może być wykorzystywana przez roboty do ratowania ludzkiego życia, na przykład w walce, ale wtedy ci żołnierze i psy-roboty mogą stać się agresywne lub popełniać krytyczne błędy, które ranią lub zabijają ludzi.

6. Sztuczna inteligencja może być pomocna w domu, pomagając w prostych zadaniach, wyszukiwaniu informacji lub rozrywce, ale to eliminuje wiele interakcji międzyludzkich, czyniąc życie mniej społecznym, serdecznym, satysfakcjonującym i uduchowionym.

7. Sztuczna inteligencja dostarcza informacji bez dramatyzmu, nastawienia czy ego, ale informacje te mogą być cenzurowane, aby celowo dostarczać nielogicznych dezinformacji i dezinformacji na najważniejsze tematy, takie jak zdrowie i bezpieczeństwo.

8. Sztuczna inteligencja może zmienić przyszłość dzięki technologii, ale może też zmienić przeszłość, pisząc historię na nowo.

Nowy model języka wizyjnego LLaVA-01 opracowany w Chinach poprawia zdolności rozumowania, ale zmaga

się ze złożonymi zadaniami wymagającymi logicznego rozumowania.

Naukowcy z wielu uniwersytetów w Chinach zaprezentowali LLaVA-01, nowy model języka wizji, który znacznie poprawia zdolności rozumowania dzięki zastosowaniu systematycznego i ustrukturyzowanego podejścia.

Tradycyjne modele języka wizji (VLM) o otwartym kodzie źródłowym często zmagają się ze złożonymi zadaniami wymagającymi logicznego rozumowania. Zazwyczaj wykorzystują one metodę bezpośredniego przewidywania, w której generują odpowiedzi bez rozbijania problemu lub nakreślania kroków niezbędnych do jego rozwiązania. Takie podejście często prowadzi do błędów, a nawet halucynacji.

Aby zaradzić tym ograniczeniom, badacze stojący za LLaVA-01 zainspirowali się modelem o1 OpenAI. Model o1 OpenAI pokazał, że wykorzystanie większej mocy obliczeniowej podczas procesu wnioskowania może zwiększyć umiejętności rozumowania modelu językowego. Jednak zamiast po prostu zwiększać moc obliczeniową, LLaVA-01 wykorzystuje unikalną metodę, która dzieli rozumowanie na ustrukturyzowane etapy.

LLaVA-01 działa w czterech odrębnych etapach:

1. Zaczyna od podsumowania pytania, identyfikując główny problem.
2. Jeśli obraz jest obecny, skupia się na istotnych częściach i opisuje je.
3. Następnie przeprowadza logiczne rozumowanie w celu uzyskania wstępnej odpowiedzi.
4. Na koniec przedstawia zwięzłe podsumowanie odpowiedzi.

Ten etapowy proces rozumowania jest niewidoczny dla

użytkownika, co pozwala modelowi zarządzać własnym procesem myślowym i skuteczniej dostosowywać się do złożonych zadań. Aby jeszcze bardziej poprawić możliwości rozumowania modelu, LLaVA-o1 wykorzystuje technikę zwaną „wyszukiwaniem wiązki na poziomie etapu”. Metoda ta generuje wiele kandydujących wyników na każdym etapie rozumowania, wybierając najlepszego kandydata do kontynuowania procesu. Podejście to jest bardziej elastyczne i wydajne niż tradycyjne metody, w których model generuje kompletne odpowiedzi przed wybraniem najlepszej.

Naukowcy uważają, że to ustrukturyzowane podejście i wykorzystanie etapowego wyszukiwania wiązki sprawi, że LLaVA-o1 będzie potężnym narzędziem do rozwiązywania złożonych zadań rozumowania, co czyni go znaczącym postępem w dziedzinie sztucznej inteligencji. Rozwój ten może potencjalnie zrewolucjonizować sposób interakcji ze sztuczną inteligencją, szczególnie w dziedzinach wymagających złożonego rozumowania, takich jak opieka zdrowotna, finanse i usługi prawne. LLaVA-o1 stanowi obiecujący krok w kierunku bardziej inteligentnych i elastycznych systemów sztucznej inteligencji.

Artykuł przetłumaczono przy pomocy AI □

**Wiceprezydent USA Vance
ostrzega Europę przed
przyjmowaniem chińskich
modeli sztucznej inteligencji**

typu open source



W przemówieniu o wysokiej stawce na paryskim szczycie AI wiceprezydent USA JD Vance naciskał na globalne przyjęcie zamkniętych systemów sztucznej inteligencji, przedstawiając sztuczną inteligencję jako broń geopolityczną. Jego uwagi, przeplatane cienko zawołowanymi groźbami, podkreślają rosnące napięcie między wysiłkami USA zmierzającymi do zmonopolizowania sztucznej inteligencji a wzrostem znaczenia Chin jako lidera innowacji open source.

USA rozpoczynają geopolityczną bitwę o przyszłość sztucznej inteligencji

W przemówieniu, które łączyło ambicję z zastraszaniem, wiceprezydent USA JD Vance wyszedł na scenę podczas paryskiego szczytu AI w dniu 12 lutego 2025 r., aby przekazać surowe przesłanie: Przyszłość sztucznej inteligencji musi być kształtowana przez Stany Zjednoczone, a wszelkie odchylenia od tej ścieżki będą kosztować. Uwagi Vance'a, które określiły sztuczną inteligencję zarówno jako narzędzie dobrobytu, jak i broń wpływów geopolitycznych, podkreślają eskalację rywalizacji między Stanami Zjednoczonymi a Chinami w wyścigu o dominację w krajobrazie sztucznej inteligencji.

„Sztuczna inteligencja to broń, która jest niebezpieczna w niewłaściwych rękach, ale jest niesamowitym narzędziem wolności i dobrobytu we właściwych rękach” – oświadczył Vance,

nie pozostawiając wątpliwości co do tego, kto jego zdaniem powinien sprawować tę władzę. Jego przemówienie, opisane przez obserwatorów jako „groźne” i „mroczne”, podkreśliło zaangażowanie USA w utrzymanie pozycji lidera w dziedzinie sztucznej inteligencji poprzez ograniczenie dostępu do krytycznych komponentów i technologii.

„Stany Zjednoczone Ameryki są liderem w dziedzinie sztucznej inteligencji, a nasza administracja planuje utrzymać ten stan rzeczy” – powiedział Vance, dodając, że USA »zamkną drogi przeciwnikom do osiągnięcia zdolności AI« na równi z własnymi.

Rozwój sztucznej inteligencji open source: wyzwanie dla dominacji i arogancji USA

Ostrzeżenia Vance’a pojawiają się w czasie, gdy chińskie modele sztucznej inteligencji typu open source, takie jak DeepSeek, zyskują na popularności na całym świecie. W przeciwieństwie do zamkniętych, zastrzeżonych systemów promowanych przez amerykańskie firmy, takie jak OpenAI, DeepSeek oferuje przejrzystą, konfigurowalną alternatywę, która już wykazała wyższą wydajność i opłacalność.

„Stany Zjednoczone nadal trzymają się myślenia” małego dziedzińca z wysokimi murami ”, zauważył jeden z analityków, odnosząc się do preferencji Ameryki dla systemów o zamkniętym kodzie źródłowym. „Ale dzięki otwartemu oprogramowaniu i tanim alternatywom” wysoki mur dziedzińca »może stać się ślepą uliczką«.

Otwarty charakter DeepSeek pozwala programistom na całym świecie replikować i modyfikować technologię bez warstw cenzury wbudowanych w modele amerykańskie. Na przykład, podczas gdy serwery DeepSeek w Chinach mogą ograniczać niektóre zapytania, takie jak te związane z protestami na

placu Tiananmen w 1989 roku i innymi wrażliwymi tematami, wersje DeepSeek działające poza Chinami mogą działać bez takich ograniczeń, umożliwiając programistom wykorzystanie jego pełnych możliwości. Ta elastyczność sprawiła, że model ten stał się potężnym narzędziem innowacji w regionach, w których dostęp do amerykańskich systemów sztucznej inteligencji jest ograniczony lub gdzie utrzymują się obawy dotyczące cenzury i bezpieczeństwa danych.

Amerykańskie systemy sztucznej inteligencji mają więcej warstw cenzury niż modele chińskie

Przemówienie Vance'a i powstanie DeepSeek podkreślają fundamentalny podział w ekosystemie sztucznej inteligencji: wybór między zastrzeżonymi systemami o zamkniętym kodzie źródłowym a alternatywami o otwartym kodzie źródłowym. Podczas gdy amerykańskie firmy, takie jak OpenAI, wspierane przez rząd, zbudowały swoją reputację na ściśle kontrolowanych, wysokowydajnych modelach, podejście Chin polegało na demokratyzacji sztucznej inteligencji poprzez udostępnienie swojej technologii globalnej społeczności. Strategia ta już zaczęła zmieniać układ sił w wyścigu o sztuczną inteligencję.

Krytycy amerykańskiego podejścia twierdzą, że jego nacisk na zamknięte systemy grozi izolacją amerykańskich firm i tłumieniem innowacji. Z kolei modele open-source, takie jak DeepSeek, oferują ścieżkę do współpracy i szybkiego rozwoju, umożliwiając mniejszym krajom i prywatnym przedsiębiorstwom konkurowanie na równych warunkach. Dla Europy, która od dawna stara się stać trzecim biegunem w transatlantycko-azjatyckiej rywalizacji technologicznej, DeepSeek stanowi potencjalną szansę na zmniejszenie zależności od technologii amerykańskiej.

Rozwój sztucznej inteligencji typu open source ma również

istotne implikacje strategiczne. Zwiększając dostępność zaawansowanych narzędzi sztucznej inteligencji, Chiny rzucają wyzwanie monopolowi USA na najnowocześniejsze technologie. Na przykład zdolność DeepSeek do przetwarzania ogromnych ilości danych z większą wydajnością niż wiele modeli amerykańskich może przyspieszyć innowacje w dziedzinach takich jak opieka zdrowotna, energia i pojazdy autonomiczne. Co więcej, dostępność sztucznej inteligencji typu open-source mogłaby umożliwić krajom i firmom rozwijanie własnych możliwości w zakresie sztucznej inteligencji bez polegania na infrastrukturze lub wiedzy specjalistycznej Stanów Zjednoczonych, zmniejszając dźwignię, jaką Stany Zjednoczone posiadają obecnie na globalnych rynkach technologicznych.

Istnieją jednak obawy dotyczące bezpieczeństwa i etycznych implikacji sztucznej inteligencji typu open source. Chociaż przejrzystość DeepSeek jest mocną stroną, rodzi również pytania o to, w jaki sposób technologia ta może zostać niewłaściwie wykorzystana lub uzbrojona. Ostrzeżenie Vance'a o tym, że sztuczna inteligencja jest „bronią”, przypomina, że stawka w tej rywalizacji jest niezwykle wysoka.

Korea Południowa blokuje DeepSeek na komputerach rządowych z powodu obaw o szpiegostwo



Południowokoreańska Narodowa Służba Wywiadowcza (NIS) zaleciła agencjom rządowym zablokowanie dostępu do DeepSeek, chińskiego chatbota sztucznej inteligencji (AI), ze względu na obawy dotyczące nadmiernego gromadzenia danych i potencjalnego chińskiego szpiegostwa.

Posunięcie, które weszło w życie w tym tygodniu, jest następstwem biuletynu bezpieczeństwa wydanego przez NIS, który szczegółowo opisywał praktyki DeepSeek, w tym przechowywanie danych użytkowników na chińskich serwerach i udzielanie stronicznych odpowiedzi na wrażliwe pytania.

Decyzja o zablokowaniu DeepSeek pojawia się w czasie, gdy Korea Południowa, wraz z innymi krajami, takimi jak Australia, Tajwan i Włochy, coraz bardziej obawia się zagrożeń bezpieczeństwa stwarzanych przez chińską technologię. NIS ostrzegł, że praktyki DeepSeek w zakresie danych mogą ujawnić poufne informacje rządowe chińskiemu rządowi, który ma prawo dostępu do danych przechowywanych w jego granicach.

Nadmierne gromadzenie danych i tendencyjne odpowiedzi

Według NIS metody gromadzenia danych przez DeepSeek są bardziej inwazyjne niż w przypadku innych usług AI. Agencja stwierdziła, że DeepSeek „zawiera funkcję zbierania wzorców wprowadzania danych z klawiatury, które mogą identyfikować osoby i komunikować się z serwerami chińskich firm, takimi jak volceapplog.com”. Możliwości te, w połączeniu z warunkami korzystania z aplikacji, które pozwalają na przechowywanie danych przez czas nieokreślony i nieograniczony dostęp do nich

przez zewnętrznych reklamodawców, wzbudziły poważne obawy dotyczące prywatności.

Jednym z najbardziej niepokojących aspektów zachowania DeepSeek są tendencyjne odpowiedzi na pytania dotyczące wrażliwych tematów. Na przykład na pytanie o pochodzenie kimchi, tradycyjnej koreańskiej potrawy, DeepSeek udzielił różnych odpowiedzi w zależności od języka zapytania. W języku koreańskim uznała kimchi za danie koreańskie, ale gdy zapytano ją po chińsku, twierdziła, że danie pochodzi z Chin. Ta rozbieżność nie jest odosobniona; NIS zauważył również, że odpowiedzi DeepSeek na pytania dotyczące Projektu Północno-Wschodniego, chińskiej inicjatywy badawczej, która twierdzi, że starożytne królestwa koreańskie są terytorium Chin, były pod silnym wpływem propagandy Komunistycznej Partii Chin (KPCh).

Globalne obawy i ograniczenia

Korea Południowa nie jest osamotniona w swoich obawach. Australia i Tajwan również zakazały DeepSeek na urządzeniach rządowych, powołując się na zagrożenia dla bezpieczeństwa narodowego. Włoski organ nadzorujący prywatność nakazał ogólnokrajową blokadę DeepSeek, dając firmie 20 dni na wyjaśnienie, w jaki sposób przestrzega europejskich przepisów o ochronie danych. Stany Zjednoczone, w tym agencje takie jak NASA i US Navy, również ograniczyły korzystanie z DeepSeek ze względu na obawy dotyczące bezpieczeństwa i prywatności.

Działania te odzwierciedlają rosnący globalny trend ostrożności wobec chińskiej technologii sztucznej inteligencji. Stany Zjednoczone nałożyły ścisłe kontrole eksportu zaawansowanych chipów i sprzętu do produkcji chipów do Chin, mając na celu ograniczenie rozwoju sztucznej inteligencji. Jednak pojawienie się DeepSeek jako taniej i wydajnej alternatywy dla amerykańskich modeli sztucznej inteligencji zachwiało zaufaniem inwestorów i wywołało pytania

o przyszłość globalnej konkurencji w dziedzinie sztucznej inteligencji.

Decyzja o zablokowaniu DeepSeek na komputerach rządowych Korei Południowej podkreśla zaangażowanie tego kraju w ochronę danych swoich obywateli i bezpieczeństwa narodowego. W miarę jak inne kraje idą w ich ślady, społeczność międzynarodowa wysyła Chinom jasny komunikat: globalny krajobraz sztucznej inteligencji nie zostanie zdominowany przez technologię, która narusza prywatność i suwerenność. Rozwój DeepSeek nie tylko wywołał technologiczny wyścig zbrojeń, ale także nasilił napięcia geopolityczne między Wschodem a Zachodem.

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, 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

czołówcę 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.

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. Jedną 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 zajść 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](https://www.thedrive.com)

[CNet.com](https://www.cnet.com)

[EnGadget.com](https://www.engadget.com)

**Google cenzuruje prawdę o
płandemii, 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 plandemię* pokazując *wymyślane 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 ~~ona~~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.

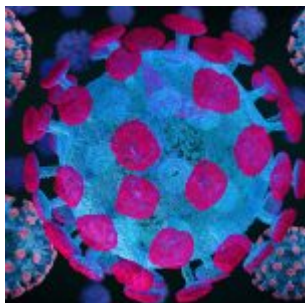
Ostatecznym celem jest zastąpienie wszystkich tradycyjnych form medycyny i zastąpienie ich scentralizowanym, kontrolowanym przez globalistów, prowadzonym przez Big Tech systemem opartym na całkowitym zniewoleniu ludzkości.

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”.

Brytyjska MHRA zdaje sobie sprawę, że szczepionki przeciwko COVID-19 będą niezwykle niebezpieczne dla społeczeństwa



Podczas gdy kłamliwe, zdradzieckie media głównego nurtu mówią nam wszystkim, że szczepionki przeciwko COVID-19 są całkowicie bezpieczne i „w 95% skuteczne” – i to kłamstwo jest powtarzane przez PJ Media, Breitbart i innych tak zwanych „konserwatywnych” wydawców wiadomości – rząd Wielkiej Brytanii

opublikował ogłoszenie o przetargu na systemu sztucznej inteligencji (AI), który może przeanalizować spodziewaną falę obrażeń wywołanych szczepionką na Covid-19 oraz skutków ubocznych.

Zatytułowane po prostu „Dostawy – 506291-2020” i [znajdujące się pod tym linkiem w Tenders Electronic Daily](#), ogłoszenie o przetargu podsumowano w następujący sposób:

MHRA pilnie poszukuje narzędzia programowego sztucznej inteligencji (AI) do przetwarzania spodziewanej dużej ilości niepożądanych reakcji lekowych (ADR) szczepionki Covid-19 i zapewnienia, że żadne szczegóły z tekstu reakcji ADRs nie zostaną pominięte.

Wniosek o system AI do przetwarzania reakcji na szczepionki na COVID-19 pochodzi od brytyjskiej [Agencji Regulacji Leków i Produktów Opieki Zdrowotnej](#), MHRA.

W dalszej części dokumentu MHRA opisuje „niezwykle pilną potrzebę na mocy art. 32(2)(c), związaną z udostępnieniem szczepionki przeciwko Covid-19” i stwierdza, że „spodziewana fala niepożądanych reakcji na szczepionkę przeciwko COVID-19 przytłoczy obecny „system prawny.”

„Jeśli MHRA nie wdroży narzędzia AI”, wyjaśnia MHRA, „nie będzie w stanie skutecznie przetwarzać tych ADR. Utrudni to jej zdolność do szybkiego zidentyfikowania wszelkich potencjalnych problemów dotyczących bezpieczeństwa szczepionki na Covid-19 i stanowi bezpośrednie zagrożenie dla życia pacjentów i zdrowia publicznego.”



English (en) ▼

Search TED for public tenders

[Advanced search](#) / [Expert search](#)

EUROPA > TED home > Display TED notice in current language

TED
 TED SIMAP
 TED eNotices
 TED eTendering

223

2020

OJ S current issue

☰

Release calendar

Supplies - 506291-2020

?

Original language

Data

Machine translations (de)

Share

🔗

🖨

📄

📄

📄

23/10/2020 S207

I. II. IV. V. VI.

United Kingdom–London: Software package and information systems

2020/S 207–506291

Contract award notice

Results of the procurement procedure

Supplies

Legal Basis:
 Directive 2014/24/EU

Przetarg jest szczególnie pilny:

Powody, dla których jest to wyjątkowo pilne – MHRA przyznaje, że planowany proces zakupu programu SafetyConnect, w tym narzędzia AI, nie zostałby zakończony wprowadzeniem szczepionki. Prowadzi do niemożności skutecznego monitorowania niepożądanych reakcji na szczepionkę Covid-19.

Wydarzenia nieprzewidywalne – kryzys związany z Covid-19 jest nowy, a postęp w poszukiwaniu szczepionki przeciwko Covid-19 nie przebiega dotychczas według żadnego przewidywalnego wzorca.

Czy to brzmi jak agencja, która oczekuje, że szczepionka przeciwko COVID-19 będzie bezpieczna i skuteczna?

Z opisu zawartego w dokumencie przetargowym jasno wynika, że:

1. MHRA spodziewa się, że szczepionki przeciwko COVID-19 spowodują falę zdarzeń niepożądanych / skutków ubocznych.
2. MHRA jest w pełni świadoma, że te niepożądane zdarzenia będą szkodzić i zabijać wielu pacjentów. W szczególności ostrzegają przed „bezpośrednim zagrożeniem życia pacjentów”.
3. Starsze systemy MHRA nie są w stanie obsłużyć spodziewanej liczby napływających zgłoszeń urazów wywołanych szczepionką na COVID-19, co oznacza, że spodziewana liczba takich zgłoszeń będzie bardzo duża i bezprecedensowa.
4. MHRA wyraża „niezwykle pilną potrzebę” wprowadzenia nowego systemu identyfikacji działań niepożądanych szczepionek przeciwko COVID-19.
5. MHRA stwierdza, że jeśli nowy system sztucznej inteligencji nie zostanie pilnie zainstalowany, nie będzie w stanie zidentyfikować wielu niepożądanych reakcji wynikających ze szczepionki przeciwko COVID-19 i że ta sytuacja wpłynie negatywnie na zdrowie publiczne. (tj. ludzie umrą.)

To pokazuje brutalną szczerść MHRA za kulisami w procesie przetargowym na systemy AI. Jednak publicznie, prawie wszystkie rządowe agencje regulacyjne na całym świecie – w tym te w Wielkiej Brytanii i USA – nie udostępniają takich szczegółów opinii publicznej i fałszywie przedstawiają szczepionki przeciwko COVID-19 jako prawie w 100% bezpieczne i skuteczne.

Teraz wiemy, że nawet brytyjska MHRA zdaje sobie sprawę, że szczepionki przeciwko COVID-19 będą niezwykle niebezpieczne

dla społeczeństwa, generując katastrofalną falę niepożądanych reakcji i śmiertelnych skutków ubocznych.

Nasuwa się pytanie: dlaczego MHRA nie mówi o tym publicznie?

Ubój ludzkiej populacji śmiercionośną bronią biologiczną i toksycznymi szczepionkami: „Wielki Reset” to medyczna tyrania i ekonomiczny komunizm

Bez wątpienia CDC też o tym wie i nie przekazuje niczego, co mogłoby zaalarmować opinię publiczną o tej rzeczywistości.

Oczywiście, opinia publiczna jest celowo trzymana w tajemnicy jeśli chodzi o niebezpieczeństwa związane ze szczepionkami przeciwko COVID-19 i niestety, nawet konserwatywne media w USA papugują propagandę przemysłu szczepionkowego, aby spróbować przekonać ludzi, że szczepionki na COVID-19 uratują świat.

Należy pamiętać, że zdradzieccy giganci technologiczni spiskowali w ciągu ostatnich kilku lat, aby cenzurować każdego, kto kwestionuje bezpieczeństwo skuteczności szczepionek. Oznacza to, że **mimo iż ludzie umierają z powodu szczepionki przeciw COVID-19, nikt nie będzie mógł o tym mówić** ani podnosić alarmu. Wszystkie zgony wywołane szczepionką na COVID-19 będą ukrywane i chronione przed opinią publiczną. Media oczywiście robią wszystko, aby chronić Big Pharmę.

W istocie Big Pharma i globalistyczni kontrolerzy oczekują, że ludzkość ustawi się w szeregu, przyjmie zastrzyk i **umrze w milczeniu**, gdy wszystko zostanie nam wszystkim skradzione przez złych globalistycznych współpracowników „Wielkiego Resetu”. Szczepionki, które osiągają cel masowego ludobójstwa i wyludnienia, są tylko jedną warstwą ich wielkiego planu,

który obejmuje również skoordynowaną, globalną kradzież wszystkich prywatnych aktywów poprzez zaplanowany upadek waluty fiducjarnej, który sprawia, że wszystkie denominowane w walutach oszczędności i aktywa są natychmiast bezwartościowe.

Nie tylko ludzie, którzy przyjmą szczepionkę i zostaną ranni lub umrą; wszystko, co kiedyś posiadali, zostanie im skonfiskowane i ukradzione przez tych samych globalistów, którzy propagują szczepionki i cenzurę.

Są to ci sami globaliści, którzy właśnie ukradli wybory w USA w 2020 roku, oczywiście, [fałszując maszyny do głosowania Dominion](#) i oświetlając cały naród fałszywymi mediami informacyjnymi.

Ostatecznym rozwiązaniem przeciwko ludzkości jest połączenie globalnego resetu gospodarczego, masowego ludobójstwa poprzez szczepionki, inżynieryjnego głodu poprzez kontrolowane załamanie zapasów żywności oraz całkowitej kontroli nad mową i myślami poprzez cenzurę Big Tech i programy socjotechniczne. Dla globalistów jesteś tylko zwierzęciem hodowlanym i uważają, nadszedł czas, aby zebrać wszystkie zwierzęta na rzeź.

Stąd potrzeba koronawirusa SARS-CoV-2 od samego początku: pozwolił globalistom uwolnić wszystkie ich programy zniewolenia ludzkości i zniszczenia ludzkiej wolności, jednocześnie twierdząc, że „chroni” cię przed tą samą bronią biologiczną, którą zbudowali i wydali w pierwszej kolejności.

Autor: [Mike Adams](#)