

Chińskie siły zbrojne wykorzystują obecnie sztuczną inteligencję do przygotowania autonomicznych maszyn wojennych



Globalna równowaga sił zmienia się na oczach wszystkich, jednak zachodnie media nie chcą dostrzec tego, co staje się coraz bardziej oczywiste. Chiny dokonały przełomu technologicznego, który może zmienić zasady międzynarodowego bezpieczeństwa i prowadzenia wojny. Podczas gdy amerykańscy przywódcy szerzą strach przed hipotetycznym szpiegowaniem przez chińskie modele sztucznej inteligencji, chińskie wojsko systematycznie wdraża te same systemy do zasilania autonomicznych pojazdów, rojów dronów i narzędzi podejmowania decyzji na polu bitwy, które mogą przytłoczyć zachodnią obronę. Stany Zjednoczone znalazły się w niebezpiecznej sytuacji, reagując nie innowacjami, ale taktyką opartą na panice, która ujawnia fundamentalne niezrozumienie zarówno technologii, jak i sytuacji strategicznej.

Kluczowe punkty:

- Chińska sztuczna inteligencja DeepSeek przewyższa obecnie ludzkie rozumowanie, a jej koszt jest znacznie niższy niż w przypadku zachodnich alternatyw.
- Chińska Armia Ludowo-Wyzwoleńcza szybko integruje

DeepSeek z autonomicznymi systemami uzbrojenia i infrastrukturą dowodzenia.

- Chińscy kontrahenci z branży obronnej budują roboty-psy, roje dronów i narzędzia do planowania działań na polu walki oparte na sztucznej inteligencji.
- Reakcja Stanów Zjednoczonych skupia się na kampaniach oszczerstw i cyberatakach, a nie na konkurencji technologicznej.
- Dążenie Chin do „suwerenności algorytmicznej” zmniejsza zależność od zachodniej technologii.
- Wydajność DeepSeek pozwala na wdrożenie go na mniejszych platformach, takich jak drony o ograniczonej mocy.

Cicha wyścig zbrojeń w dziedzinie sztucznej inteligencji

Podczas gdy amerykańscy dyrektorzy ds. technologii i dziennikarze debatuje nad tym, czy DeepSeek może szpiegować użytkowników poprzez pobrane pliki, chiński kompleks militarno-przemysłowy po cichu rewolucjonizuje sztukę wojenną. Dowody są widoczne dla tych, którzy chcą wyjść poza strach szerzony przez zachodnie media. W lutym chiński państwowy gigant zbrojeniowy Norinco zaprezentował pojazd wojskowy zdolny do samodzielnego prowadzenia operacji wsparcia bojowego z prędkością 50 kilometrów na godzinę, napędzany bezpośrednio przez sztuczną inteligencję DeepSeek. Jest to tylko jeden z widocznych przykładów tego, co analitycy ds. obronności opisują jako systematyczną kampanię Pekinu mającą na celu wykorzystanie sztucznej inteligencji do uzyskania przewagi militarnej.

Różnica technologiczna staje się bardziej niepokojąca, gdy przyjrzymy się szczegółom. Naukowcy z Xi'an Technological University ujawnili, że ich system oparty na DeepSeek może w ciągu zaledwie 48 sekund ocenić 10 000 scenariuszy bitewnych z różnymi zmiennymi, ukształtowaniem terenu i rozmieszczeniem

sił. Tradycyjnie to samo zadanie wymagałoby 48 godzin pracy zespołu planistów wojskowych. Tysiąckrotne przyspieszenie zdolności podejmowania decyzji może okazać się decydujące w każdym przyszłym konflikcie, pozwalając chińskim dowódcom prześcignąć przeciwników w myśleniu i manewrowaniu w czasie rzeczywistym.

Reakcja Ameryki: strach przed kompetencjami

Zamiast stawiać czoła temu wyzwaniu technologicznemu z innowacyjnością i determinacją, amerykańskie instytucje zareagowały w sposób, który można określić jedynie jako panikę. Po publicznym udostępnieniu DeepSeek i późniejszym spadku na giełdzie, firma zajmująca się sztuczną inteligencją doświadczyła potężnego cyberataku, który wiele źródeł przypisuje amerykańskim agencjom wywiadowczym. Zamiast opracowywać konkurencyjną technologię lub strategiczne partnerstwa, reakcja Ameryki odzwierciedlała te same kampanie oszczerstw, które były stosowane przeciwko przeciwnikom politycznym w poprzednich cyklach wyborczych.

Ironia jest tak wyraźna, że można ją wyczuć w powietrzu. Podczas gdy amerykańskie media ostrzegają przed hipotetycznym zagrożeniem prywatności wynikającym z pobierania modeli DeepSeek, chińscy inżynierowie kończą studia w rekordowej liczbie i budują systemy, które mogą na nowo zdefiniować współczesną wojnę. Stany Zjednoczone wydają się bardziej skupione na próbach hakowania chińskich badań niż na opracowywaniu własnych przełomowych rozwiązań. Takie podejście odzwierciedla fundamentalne niezrozumienie zarówno technologii, jak i możliwości Chin. Modele DeepSeek nie są oprogramowaniem szpiegowskim; są to zaawansowane silniki wnioskowania, które do działania wymagają dodatkowego oprogramowania, opracowanego głównie w Stanach Zjednoczonych i Europie Zachodniej.

Świt autonomicznej wojny

Wdrożenie DeepSeek przez Chińską Armię Ludowo-Wyzwoleńczą to coś więcej niż tylko postęp technologiczny; sygnalizuje ono fundamentalną zmianę w doktrynie wojskowej. Chińskie dokumenty dotyczące obronności ujawniają plany dotyczące psów-robotów zasilanych sztuczną inteligencją, które będą działać w grupach, rojów dronów autonomicznie śledzących cele oraz wizualnie immersyjnych centrów dowodzenia, które mogą zwiększyć skuteczność ludzkich dowódców, ale także potencjalnie zastąpić ich w niektórych funkcjach. Systemy te nie są koncepcjami science fiction, ale aktywnymi projektami rozwojowymi udokumentowanymi w rejestrach zamówień i zgłoszeniach patentowych.

To, co sprawia, że DeepSeek jest szczególnie groźny dla przewagi militarnej Zachodu, to jego wydajność obliczeniowa. Badacze z PLA przypisują systemowi zmniejszenie zużycia energii podczas szkolenia o około 40 procent w porównaniu z systemami klasy GPT-4, przy zachowaniu równoważnych możliwości. Ta wydajność pozwala na wdrożenie systemu na mniejszych platformach o ograniczonej mocy, takich jak drony i jednostki operacyjne. Mniejszy rozmiar DeepSeek – około jednej ósmej GPT-4 – umożliwia jego działanie na platformach, na których zachodnie systemy sztucznej inteligencji byłyby niepraktyczne, tworząc coś, co analitycy wojskowi nazywają „wdrożeniem na poziomie brzegowym”, które zachowuje funkcjonalność nawet w przypadku zakłóceń lub przerw w komunikacji.

Integracja tych systemów rozciąga się na całą infrastrukturę wojskową Chin. PLA coraz częściej traktuje DeepSeek zarówno jako akcelerator technologiczny, jak i eksperyment doktrynalny, a jego mocne strony w zakresie wydajności obliczeniowej i integracji krajowego sprzętu idealnie wpisują się w dążenie Chin do stworzenia odpornych, samowystarczalnych sieci bojowych. Choć chińscy urzędnicy publicznie

zobowiązują się do utrzymania kontroli człowieka nad systemami uzbrojenia, obecnie istnieje technologia umożliwiająca w pełni autonomiczne operacje w kontrolowanych środowiskach. Jak zauważył jeden z analityków wojskowych, nie chodzi już o to, czy Chiny mogą wykorzystać sztuczną inteligencję w działaniach wojennych, ale kiedy i w jakim zakresie zdecydują się to zrobić.

Stany Zjednoczone stoją na rozdrożu, mając do czynienia z konkurentem, który osiągnął to, co jeszcze kilka lat temu wielu ekspertów uważało za niemożliwe. Chiny nie tylko dogoniły innych w dziedzinie sztucznej inteligencji, ale w wielu praktycznych zastosowaniach wyprzedziły ich. Czas na kampanie oszczerstw i cyberataki minął. Pozostaje pilna potrzeba trzeźwej oceny sytuacji i zdecydowanych innowacji, zanim równowaga się ulegnie nieodwracalnej zmianie.

Wielkie odłączenie: Odzyskiwanie naszych umysłów z cyfrowego natarcia



- Od początku 2010 roku obserwuje się gwałtowny i stały wzrost depresji, lęków i samookaleczeń wśród nastolatków, co eksperci łączą z masowym

upowszechnieniem się smartfonów i przejściem do „dzieciństwa opartego na telefonach”.

- Urządzenia te szkodzą zdrowiu psychicznemu, zastępując ważne czynności w świecie rzeczywistym, takie jak osobiste kontakty towarzyskie i swobodna zabawa. W przypadku nastolatków sytuację pogarsza ciągłe porównywanie się z innymi i rozproszenie uwagi spowodowane mediami społecznościowymi.
- Podstawowym zaleceniem jest świadome, długie odpoczywanie od telefonów. Praktyczne kroki obejmują wyłączenie zbędnych powiadomień, fizyczne przeniesienie telefonu do innego pokoju i „grupowanie” korzystania z telefonu w określone bloki czasowe.
- Oprócz działań indywidualnych ruch ten opowiada się za wprowadzeniem kluczowych norm chroniących dzieci, w tym zakazu używania smartfonów przed rozpoczęciem nauki w szkole średniej, zakazu korzystania z mediów społecznościowych przed ukończeniem 16 roku życia oraz wprowadzenia szkół bez telefonów.
- Oprócz rzadszego korzystania z telefonów ważne jest angażowanie się w trudne lub przerażające czynności w świecie rzeczywistym. Pokonywanie tych możliwych do pokonania przeszkód buduje pewność siebie i poczucie własnej skuteczności, stanowiąc bezpośrednie antidotum na niepokój.

Coraz większa grupa ekspertów i rozwijający się ruch społeczny przedstawiają bardzo prostą receptę na rosnący niepokój: odłóż telefon.

Na czele tego ruchu na rzecz umiarkowanego korzystania z urządzeń cyfrowych stoi psycholog społeczny i autor Jonathan Haidt. Nie jest to jedynie sugestia dotycząca stylu życia, ale odpowiedź na to, co Haidt określa jako głęboką zmianę w dzieciństwie i główny czynnik kryzysu zdrowia psychicznego dotykającego młodsze pokolenia. Ruch ten, zrodzony z bestsellerowej książki Haidta „The Anxious Generation”

(Niepokojące pokolenie), twierdzi, że kluczem do lepszego samopoczucia psychicznego jest celowe i długotrwałe odstawienie urządzeń, które zdominowały współczesne życie.

Alarmujące dane stanowią podstawę tego argumentu. Na początku 2010 roku wskaźniki chorób psychicznych wśród nastolatków zaczęły gwałtownie i trwale rosnać. W latach 2010–2018 liczba diagnoz depresji i lęku wśród studentów w Stanach Zjednoczonych wzrosła ponad dwukrotnie.

Jeszcze bardziej niepokojący jest katastrofalny wzrost liczby wizyt na pogotowiu z powodu samookaleczeń, który w ciągu dziesięciu lat poprzedzających rok 2020 wyniósł 188% wśród nastoletnich dziewcząt i 48% wśród chłopców. Haidt i podobnie myślący badacze wskazują na masowe upowszechnienie smartfonów i mediów społecznościowych jako katalizator tej „wielkiej przemiany dzieciństwa”, czyli przejścia od życia opartego na zabawie do życia opartego na telefonie.

Mechanizm, poprzez który urządzenia wpływają na zdrowie psychiczne, jest wielowymiarowy. Smartfony działają jak „blokery doświadczeń”, wypierając wzbogacające zajęcia, takie jak osobiste kontakty społeczne, nieustrukturyzowana zabawa i praktyczne hobby. Odciągają one użytkowników od ich najbliższego otoczenia, tworząc stan „wiecznej nieobecności”.

W przypadku nastolatków, którzy przechodzą intensywny okres rozwoju społecznego i emocjonalnego, skutki są szczególnie silne. Platformy społecznościowe zachęcają do ciągłego porównywania się z innymi i mogą być bezlitośnie okrutne, kwantyfikując pozycję społeczną poprzez liczbę polubień i obserwujących. Kompulsywna potrzeba sprawdzania powiadomień i odświeżania feedów rozprasza uwagę, trenując mózg do rozpraszania się, a nie do głębokiej koncentracji.

Praktyczne kroki w kierunku cyfrowego detoksu

Według zwolenników takich jak Alexa Arnold z Anxious Generation Movement, rozwiązaniem jest świadome odsunięcie telefonu od codziennego życia. Obejmuje to praktyczne kroki, takie jak wyłączenie zbędnych powiadomień i fizyczne umieszczenie telefonu w innym pokoju na kilka godzin, aby ułatwić sobie okresy nieprzerwanej, głębokiej pracy.

Arnold sugeruje „grupowanie” korzystania z telefonu, na przykład przeznaczając konkretny 20-minutowy blok czasu na śledzenie wiadomości, zamiast sprawdzania aplikacji wielokrotnie w ciągu dnia. Podstawową zasadą jest to, że im dłuższe przerwy w korzystaniu z urządzenia, tym lepiej mózg może odzyskać swoją naturalną zdolność do koncentracji i spokoju.

Oprócz umiarkowanego korzystania z urządzeń cyfrowych ruch ten podkreśla znaczenie stawiania sobie wyzwań w prawdziwym świecie. Robienie rzeczy, które są przerażające lub trudne, czy to w środowisku zawodowym, czy społecznym, jak na przykład nawiązanie rozmowy z nieznanym, buduje pewność siebie i umiejętności.

Ta praktyka podejmowania trudnych wyzwań służy jako antidotum na niepokój, wzmacniając poczucie własnej skuteczności, którego nie może zapewnić walidacja oparta na ekranie. Jest to powrót do zasady, że odporność buduje się poprzez pokonywanie możliwych do pokonania przeszkód, proces, który został osłabiony w erze cyfrowej ucieczki od rzeczywistości i nadopiekuńczego rodzicielstwa.

Zagrożenie ze strony sztucznej

inteligencji

W momencie, gdy walka z negatywnym wpływem mediów społecznościowych nabiera tempa, pojawia się nowe wyzwanie technologiczne: sztuczna inteligencja (AI).

Haidt ostrzega, że AI może zwielokrotnić szkody wyrządzane przez media społecznościowe, tworząc treści, które są „o wiele bardziej uzależniające”. Pierwsze konsekwencje są już widoczne, od „trenerów samobójstw” AI po wykorzystanie technologii deepfake do nękania i szantażowania. Jego zdaniem walka musi teraz zostać rozszerzona, aby zapobiec utracie kolejnego pokolenia na rzecz tego, co nazywa „obcą inteligencją”.

„Rola technologii jest zaspokajanie ludzkich potrzeb, a nie dyktowanie ich. Idealna relacja to taka, w której technologia jest celowo dostosowana do ludzkiego dobrobytu i go wzmacnia” – powiedział Enoch z *BrightU.AI*. „Wymaga to projektowania i wykorzystywania technologii jako narzędzia wspierającego cele, wartości i relacje społeczne”.

Ostatecznie przesłanie ruchu Anxious Generation jest ostrożnym optymizmem. Twierdzi on, że obecny kryzys zdrowia psychicznego nie jest nieuniknionym skutkiem ubocznym współczesnego życia, ale bezpośrednią konsekwencją konkretnych wyborów technologicznych.

Maski opadły



Jeśli wierzycie w napis na czapkach noszonych w izraelskim Knesecie – „Trump prezydentem pokoju” – to jesteście w błędzie. Halloween może się zbliżać, ale nie potrzebujecie przerażających masek, aby zdać sobie sprawę z horroru, z którym mamy do czynienia. Trump jest kulminacją długiej historii horroru.

W przeciwieństwie do swoich poprzedników, którzy przygotowali mu drogę i zazwyczaj nosili tradycyjne maski, aby ukryć swoje złe uczynki, jest on największym oszustem, jaki kiedykolwiek zasiadł w Białym Domu. Jest podżegaczem wojennym, ludobójczym mordercą i wrogiem ludzi w kraju i za granicą, tak oczywistym, tak kapryśnym, tak nieobliczalnym – człowiekiem niekończących się gróźb – że nikt nie powinien być zaskoczony, budząc się pewnego ranka z wiadomością, która może wydawać się „szokująca”. Każdy powinien spodziewać się niespodzianek, ale nie słodczy, tylko psikusów.

Trump jest jak reklama, która mówi ci, że jej bohaterowie nie są zwykłymi ludźmi, ale aktorami, a ich gadka nie jest prawdziwa – tylko po to, żebyś kupił produkt, który reklamują. Każda zagrywka Trumpa jest podkreślona.

Jedynym sposobem, w jaki można wyjaśnić jego zachowanie, jest to, że jest on zwieńczeniem trwającego od dziesięcioleci rozwoju kultury amerykańskiej, w której aktorstwo jest przedstawiane jako tak fałszywe, że publiczność uważa je za prawdziwe właśnie ze względu na jego fałszywość. Jest niebezpiecznym żartem, tym bardziej niebezpiecznym, że tak dobrze wpisuje się w szerszy kontekst rozwoju kultury, który Neil Postman w 1985 roku trafnie nazwał *Amusing Ourselves to Death: Public Discourse in the Age of Show Business* (Bawimy

się na śmierć: dyskurs publiczny w erze show-biznesu), a Neal Gabler nazwał później *Life:The Movie: How Entertainment Conquered Reality* (Życie: film: jak rozrywka podbiła rzeczywistość).

Jest on kulminacją utajonego nurtu despotyzmu, który przepływał przez historię Ameryki, zwłaszcza w ciągu ostatnich dwudziestu pięciu lat, ale który wielu postrzega jedynie jako walkę między partiami politycznymi, tak zwanymi dobrymi i złymi. Nie dostrzegają oni, że faszyzm jest jak zamek, którego budowa od fundamentów zajmuje lata i wymaga powolnej akceptacji przez wszystkie odcienie opinii politycznej stopniowej utraty podstawowych wolności, akceptację korporacyjnego państwa wojennego oraz tajnego rządu działającego w ramach agencji wywiadowczych, takich jak CIA, NSA i DIA (Defense Intelligence Agency), współpracujących ściśle z głównymi mediami i korporacjami z Doliny Krzemowej w ramach partnerstw mających na celu propagandę i szpiegowanie społeczeństwa.

Ktoś taki jak Trump nie rodzi się z dnia na dzień. Jego poprzednikami są wszyscy ci dwupartyjni pochlebcy, którzy zaakceptowali oficjalne wyjaśnienie wydarzeń z 11 września i natychmiastowe wprowadzenie ustawy Patriot Act (przygotowanej za rządów Clintona), stan wyjątkowy ogłoszony przez George'a W. Busha 14 września 2001 r. i od tego czasu corocznie przedłużany, wojny przeciwko Jugosławii, Afganistanowi, Irakowi, Syrii, Libii, Rosji, Iranowi, Palestyńczykom itp. (wojny rozpoczęte i wspierane przez republikanów i demokratów), ratowanie wielkich banków i instytucji finansowych w 2009 r., zamach stanu w Ukrainie w 2014 r., tzw. wojna z terroryzmem, afera Russiagate, pozasądowe zabójstwa dokonywane przez prezydentów USA, niekończąca się propaganda, rozwój „partnerstw” publiczno-prywatnych, które doprowadziły do prywatyzacji usług rządowych, kłamstwa związane z COVID, nowa zimna wojna i ogromny wpływ Izraela na rząd USA itp. Lista jest długa. Trump, ten tchórzliwy despota, nie pojawił

się z dnia na dzień; jest on wynikiem długotrwałych procesów.

„Ale co się stanie” – pisze Gary Wills w książce *Reagan's America: Innocents at Home* – *„jeśli patrząc w historyczne lusterko wsteczne, zobaczymy tylko film?”*

Faszyzmowi często towarzyszy marzycielska samozadowolenie i efekty hollywoodzkie, jak w przypadku filmu „Triumf woli” Leni Riefenstahl z 1935 roku, nazistowskiej propagandy zamówionej przez Adolfa Hitlera. Dzisiaj kultura ekranowa dominuje w myśleniu ludzi dzień i noc, a obrazy i filmy cyfrowe towarzyszą ich snom i marzeniom. Jako aktor reality show, Trump jest idealnym ucieleśnieniem tej kultury ekranowej.

Wszyscy czekają teraz na kulminację swoich celuloidowych iluzji, na rozwiązanie akcji horroru, jak w opowiadaniu Poe’go „Dom Usherów”. Niemiecki dramaturg Bertolt Brecht powiedział: „Aby zrozumieć faszyzm, trzeba zrozumieć kapitalizm, z którego się wywodzi”.

Źródłem kapitalizmu jest potrzeba tworzenia nierówności między bogatymi a biednymi. Gdy zagrożona jest ta potrzeba, kapitalizm przechodzi metamorfozę w jawny totalitaryzm.

W 1985 roku, roku „zabawy na śmierć”, Donald Trump, fałszywy deweloper, nabył hotel Atlantic City Hilton i przemianował go na Trump’s Castle, co było oznaką jego obsesyjnej megalomanii. Hołd Trumpa dla samego siebie zakończył się bankructwem siedem lat później, zapowiadając przyszły los Stanów Zjednoczonych. Był to pierwszy rok drugiej kadencji Ronalda Reagana, byłego aktora, którego krytycy nazywali „aktorem prezydentem”. Jednak sam Reagan był dumny ze swojej kariery aktorskiej; uważał, że dobrze mu służyła w Białym Domu, jak pisze Gary Wills w książce *Reagan's America: Innocents at Home*.

Trump sprawia, że Reagan wydaje się cholernie autentyczny. Wszystko w Trumpie jest kiczowate, fałszywe pod każdym względem, kopia kopii kopii w kulturze kopii. Ale to właśnie stanowi o jego atrakcyjności dla tych, którzy nie potrafią

odróżnić iluzji od rzeczywistości. Czy to ten absurdalny facet z reality show, który zwalnia ludzi na prawo i lewo, czy naprawdę prezydent Stanów Zjednoczonych? To bardzo trafne, że powrócił on do prezydentury w momencie, gdy sztuczna inteligencja zyskała na znaczeniu.

0 godz. 17:16 9 listopada 1965 r. w Nowym Jorku wysiadłem z wagonu metra IRT nr 4 na podwyższonej stacji na 161st z widokiem na stadion Yankee Stadium i wszystkie światła zgasły w północno-wschodniej części Stanów Zjednoczonych. Takie rzeczy zdarzają się, kiedy najmniej się tego spodziewamy. Jeden szczur może spędzić lata na budowaniu domku z kart dla własnej chwały, ale inny szczur może zgasić światła i zburzyć go w ciągu jednej nocy, jak śpiewa Kris Kristofferson w *Darby's Castle*.

Można być pewnym, że za murami amerykańskiej wioski Potemkina rządzące szczury walczą między sobą o dominację, a społeczeństwo – niezależnie od tego, czy mieszka w domku dla lalek iluzji, nadal myśląc, że pod rządami Trumpa wszystko jest w porządku, czy też obawiając się, że nadejdzie coś znacznie gorszego – pewnego dnia obudzi się z wielkim zaskoczeniem. *„Ale wystarczyła jedna noc, aby to zburzyć / Kiedy zamek Darby'ego runął na ziemię”*.

Nikt nie jest w stanie przewidzieć, czy tym zaskoczeniem będzie upadek osobistego zamku Darby'ego, gospodarki Stanów Zjednoczonych i świata, naszej pozornej demokracji, czy też całego świata pod naporem spadających rakiet nuklearnych. Jednak podobnie jak w przypadku nigdy nieotwieranej pudełka z niespodzianką, kiedy w ciemności nocy złowrogie siły przekręcą pokrętło, obudzimy się z ogromnym szokiem. Maski zostały bowiem zdjęte.

Edward Curtin

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



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

- Po rozpoczęciu SAI jej zatrzymanie mogłoby wywołać „szok terminacyjny” – gwałtowne, śmiertelne skoki temperatury. Scentralizowana kontrola grozi konfliktem geopolitycznym, ponieważ państwa lub miliarderzy mogliby wykorzystać manipulację klimatem jako broń, co doprowadziłoby do wojen o zasoby.
- SAI jest zgodna z poglądami globalistycznych elit, takich jak Bill Gates i Światowe Forum Ekonomiczne, które pod pretekstem ochrony środowiska forsują kontrolę populacji. Zamiast niebezpiecznych technokratycznych rozwiązań, prawdziwe rozwiązania obejmują energię odnawialną, rolnictwo regeneracyjne i zdecentralizowaną samowystarczalność.

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

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

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

dużych erupcji wulkanicznych.

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

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

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

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

Ukryte zagrożenia: kwaśne deszcze, uszkodzenia płuc i zatrucie gleby

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

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

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

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

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

Naukowcy ostrzegają, że nie ma czegoś takiego jak globalny termostat, który odpowiadałby wszystkim. Regiony korzystające ze zmienionych warunków pogodowych mogłyby prosperować, podczas gdy inne cierpiałyby z powodu katastrofalnych susz lub powodzi – podsycając wojny o wodę i żywność.

Globalistyczna agenda stojąca za geoinżynierią

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

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

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

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

Przełom w dziedzinie autonomicznego widzenia bez elektroniki: miękkie „oko” robota zasilane światłem



- Naukowcy opracowali miękkie „oko” robota wykonane z hydrożelu reagującego na światło, które samodzielnie dostosowuje ostrość bez zewnętrznych źródeł zasilania, naśladując biologiczne widzenie.
- Wbudowane nanocząsteczki tlenku grafenu pochłaniają światło, ogrzewając hydrożel, który kurczy się, aby wyostrzyć soczewkę – eliminując potrzebę stosowania baterii lub elektroniki.
- Soczewka rozróżnia szczegóły o wielkości zaledwie czterech mikrometrów (np. pazury kleszczy, strzępki grzybów), dorównując tradycyjnym soczewkom mikroskopowym, pozostając jednocześnie w pełni autonomiczna.
- Potencjalne zastosowania obejmują nadludzkie widzenie (źrenice podobne do kocich, siatkówki mątwy), technologie medyczne do noszenia na ciele oraz autonomiczne roboty do pracy w niebezpiecznych środowiskach.

- W przeciwieństwie do technologii transhumanistycznej opartej na sztucznej inteligencji, ta innowacja działa niezależnie, zmniejszając zależność od systemów kontrolowanych przez wielkie firmy technologiczne, jednocześnie rodząc pytania etyczne dotyczące nadzoru i integracji biohybrydowej.

W przełomowym kroku w dziedzinie miękkiej robotyki naukowcy opracowali miękkie, autonomiczne „oko” robota, które jest w stanie ustawić ostrość bez zewnętrznego źródła zasilania.

Zainspirowana widzeniem zwierząt, ta niezwykle wydajna soczewka – zbudowana z hydrożelu reagującego na światło – może zrewolucjonizować robotykę, technologie noszone na ciele i urządzenia autonomiczne, eliminując jednocześnie potrzebę stosowania tradycyjnej elektroniki lub baterii.

BrightU.AI’s Enoch wyjaśnia, że oko robotyczne, znane również jako oko bioniczne lub sztuczny system widzenia, jest innowacją technologiczną zaprojektowaną w celu przywrócenia wzroku lub poprawy zdolności widzenia osób z zaburzeniami wzroku lub niewidomych. Systemy te zazwyczaj składają się z dwóch głównych elementów: protezy wzrokowej i jednostki przetwarzającej.

Naukowcy z Georgia Institute of Technology, pod kierownictwem doktoranta Coreya Zheng i inżyniera biomedycznego dr Shu Jia, zaprezentowali swoją innowacyjną soczewkę w badaniu opublikowanym w czasopiśmie *Science Robotics* w środę, 22 października. Soczewka naśladuje biologiczny wzrok, dynamicznie dostosowując ostrość w odpowiedzi na światło – bez konieczności zasilania elektrycznego.

W przeciwieństwie do konwencjonalnych robotów, które opierają się na sztywnych czujnikach i komponentach elektronicznych, miękka robotyka oferuje elastyczność i możliwość dostosowania. Zheng wyjaśnił: „Jeśli patrzymy na roboty, które są bardziej miękkie, sprężyste, być może nie wykorzystują energii

elektrycznej, to musimy zastanowić się, w jaki sposób będziemy wykrywać sygnały za pomocą tych robotów”.

Jak to działa: hydrożel spotyka się z tlenkiem grafenu

Soczewka jest wykonana z hydrożelu – materiału na bazie polimeru, który zatrzymuje i uwalnia wodę, umożliwiając jej przechodzenie między stanem płynnym a stałym. Pod wpływem ciepła hydrożel kurczy się, a po schłodzeniu pęcznieje.

Aby wykorzystać światło jako źródło energii, naukowcy umieścili w hydrożelu nanocząsteczki tlenku grafenu. Te ciemne cząsteczki pochłaniają światło, ogrzewając się pod wpływem promieniowania słonecznego o odpowiedniej intensywności. Ciepło powoduje kurczenie się hydrożelu, co powoduje skupienie przymocowanej soczewki z polimeru krzemowego. Gdy światło słabnie, hydrożel rozszerza się, umożliwiając soczewce powrót do pierwotnego stanu.

Mechanizm ten umożliwia soczewce autonomiczne działanie, reagując na światło widzialne w całym spektrum.

W testach laboratoryjnych soczewka hydrożelowa wykazała niezwykłą czułość, rozróżniając tak drobne szczegóły, jak:

- 4-mikrometrowe szczeliny między pazurami kleszcza
- 5-mikrometrowe strzępki grzybów
- 9-mikrometrowe włoski na nodze mrówki

Taka precyzja dorównuje tradycyjnym szklanym soczewkom mikroskopowym, ale bez konieczności ręcznej regulacji lub zasilania.

Przyszłe zastosowania: poza ludzkim wzrokiem

Zespół już integruje soczewkę z mikroprzepływowym systemem zaworów wykonanym z tego samego hydrożelu. Zheng zauważył, że może to umożliwić stworzenie samozasilających się, inteligentnych systemów kamer, w których światło wykorzystywane do obrazowania jednocześnie zasila urządzenie.

Co więcej, zdolność adaptacyjna hydrożelu otwiera drzwi do nadludzkiego wzroku. Potencjalne zastosowania obejmują:

- Kocimi źrenicami do wykrywania zakamuflowanych obiektów
- Inspirowane mątwą siatkówki w kształcie litery W do postrzegania kolorów wykraczających poza ludzkie możliwości

„Możemy faktycznie kontrolować soczewkę w naprawdę wyjątkowy sposób” – podkreślił Zheng, sugerując przyszłe innowacje w robotyce inspirowanej biologią.

To przełomowe odkrycie wpisuje się w rozwijającą się dziedzinę miękkiej robotyki, która priorytetowo traktuje elastyczność i integrację z systemami biologicznymi. Potencjalne zastosowania obejmują:

- Technologie zdrowotne do noszenia, płynnie łączące się z ludzkim ciałem
- Autonomiczne roboty eksploracyjne poruszające się po niebezpiecznym lub nierównym terenie
- Urządzenia wojskowe i monitorujące wymagające minimalnej mocy

Co najważniejsze, soczewka eliminuje zależność od baterii i elektroniki, rozwiązując kluczowe ograniczenia w robotyce, jednocześnie zmniejszając wpływ na środowisko.

Krok w kierunku autonomicznych, samowystarczalnych maszyn

W czasie, gdy globaliści promują transhumanizm i nadzór oparty na sztucznej inteligencji, ta innowacja daje przedsmak zdecentralizowanej, samowystarczalnej technologii. W przeciwieństwie do uzależniających urządzeń inteligentnych kontrolowanych przez wielkie firmy technologiczne, soczewka ta działa niezależnie – zasilana wyłącznie światłem.

Jednak sceptycy ostrzegają przed potencjalnym nadużyciem. Czy taka technologia może zostać wykorzystana jako broń? Zintegrowana z sieciami nadzoru? Lub – biorąc pod uwagę jej skład hydrożelowy – powiązana z powstającymi systemami biohybrydowymi, które zacierają granicę między maszyną a organizmem?

Na razie skupiamy się na jej potencjale. Jak stwierdził Zheng, soczewka ta stanowi „nowy sposób myślenia o wykrywaniu w robotyce miękkiej”.

To miękkie robotyczne oko stanowi przełomowy postęp w dziedzinie autonomicznego wykrywania bez użycia elektroniki. Wykorzystując światło i hydrożel, naukowcy otworzyli przyszłość, w której maszyny widzą – i dostosowują się – jak żywe organizmy.

W miarę jak ludzkość nadal obserwuje rozwój miękkiej robotyki i zdecentralizowanych technologii, jedno jest jasne: innowacje kwitną tam, gdzie natura spotyka się z pomysłowością – bez ograniczeń ze strony korporacji i rządów.

Niezadowolenie społeczne w Chinach rośnie wraz z zakończeniem czwartego plenum KPCh



Media państwowe wychwalają postępy gospodarcze, ale Chińczycy mówią o pogłębiających się trudnościach, sprzecznościach politycznych i słabnącym zaufaniu do rządów Xi.

Gdy Komunistyczna Partia Chin (KPCh) zakończyła 23 października swoje czwarte plenum, państwowe media rozpoczęły zakrojoną na szeroką skalę kampanię propagandową, wychwalając oficjalny komunikat z posiedzenia.

Media państwowe chwala postępy gospodarcze, ale obywatele i obserwatorzy opisują pogłębiające się trudności, sprzeczności polityczne i malejące zaufanie społeczeństwa do rządów chińskiego przywódcy Xi Jinpinga.

Na chińskich ulicach nastroje były znacznie mniej triumfalne.

Obserwatorzy Chin opisali komunikat jako sprzeczny, pełen pustych haseł i fałszywych twierdzeń.

Retoryka a rzeczywistość

W komunikacie określono ogólne cele gospodarcze i społeczne dla nadchodzącego 15. planu pięcioletniego Chin i zobowiązano się do wzmocnienia podstaw realnej gospodarki, zbudowania

silnego rynku krajowego i zwiększenia popytu krajowego w ciągu najbliższych pięciu lat.

W rzeczywistości trudności gospodarcze są powszechne.

„Bezrobocie jest powszechne, a firmy masowo upadają” – powiedział „The Epoch Times” naukowiec z Chin kontynentalnych, ze względów bezpieczeństwa posługujący się pseudonimem Jin Yan.

„Idąc ulicą, widzisz zamknięte sklepy po obu stronach. Ogólnie, warunki życia ludzi się pogarszają”.

Stwierdził, że rzeczywistość stoi w ostrym kontraście z oficjalnym optymizmem, przy wysokim bezrobociu i z opóźnionymi wypłatami pensji nawet wśród rządowych pracowników KPCh. Komunikat partii określił jako rażące kłamstwo.

Według Jina otwarta krytyka KPCh i Xi staje się coraz powszechniejsza, nawet w miejscach, które kiedyś były ściśle monitorowane.

„Gdy ostatnio jeździłem taksówkami, kierowcy zaczęli przeklinać partię lub Xi zupełnie spontanicznie. Wcześniej nie odważyliby się powiedzieć ani słowa z powodu kamer w środku” – powiedział Jin. „Teraz nie mogą powstrzymać swojej frustracji”.

Fasada partii

Jeden z działaczy społecznych wypowiedział się dla „The Epoch Times” pod pseudonimem Wang Hua. Powiedział, że problemy w chińskim społeczeństwie wykraczają poza realną gospodarkę.

„Nic teraz nie działa, ani gospodarka wirtualna, ani sektor prywatny, nic” – powiedział. „Chiński ‘wewnętrzny obieg [gospodarczy]’ nie może już funkcjonować, więc kraj przetrwa dzięki eksportowi. Jeśli handel zostanie odcięty przez wojny

handlowe, system się zawali”.

Wang opisał obraz jedności KPCh jako fasadę skrywającą zaciekle konflikty wewnętrzne. Powiedział, że [niedawne wstrząsy](#) w Centralnej Komisji Wojskowej (CKW) tylko uwydatniają niestabilność na szczytach władzy.

„Wszyscy są w tej samej sytuacji” – dodał Wang. „Żaden z nich nie robi nic dobrego dla ludzi”.

Sprzeczności w planie KPCh

Były chiński adwokat zajmujący się prawem karnym Zuo Zhihai, który uciekł z Chin na początku tego roku, powiedział gazecie „The Epoch Times”, że plenum ujawniło głębokie sprzeczności w modelu zarządzania partią.

„Czwarte plenum zasadniczo dotyczyło konsolidacji kontroli i podziału chińskich zasobów między elity” – powiedział Zuo. „Xi Jinping nadal dominuje, ale dyktatorzy, podobnie jak Stalin czy Lenin, w końcu płacą za to cenę”.

Zauważył, że temat plenum – modernizacja w stylu chińskim i odrodzenie narodowe – jest sprzeczny. Zuo stwierdził, że prawdziwa modernizacja wymaga reform politycznych i odpowiedzialności publicznej, a tych warunków brakuje Komunistycznej Partii Chin. Dodał, że pod jej rządami jednopartyjnymi nikt nie jest prawdziwie wolny.

„Albo służysz partii, albo jesteś traktowany jak obywatel drugiej kategorii” – powiedział.

Zuo zwrócił uwagę na niespójności między przekazem partii, takim jak „kompleksowe przywództwo partii” a „najpierw ludzie”.

„Jeśli partia kontroluje wszystko – prawo, nominacje, podatki – to ludzie nie mają nic do powiedzenia” – powiedział. „KPCh rządzi w imieniu ludzi, ale nie dla nich. Hasło ‘najpierw

ludzie' to tylko pretekst".

Puste obietnice, rosnące niezadowolenie

W całych Chinach szerzy się społeczne rozczarowanie hasłami politycznymi KPCh.

„Każde takie spotkanie najwyższego kierownictwa jest katastrofą dla zwykłych obywateli” – powiedział Jin. „Pekin zostaje zamknięty, zaostrzane są środki zapewniające stabilność, a codzienne życie ulega zakłóceniu. Myślę, że gniew społeczny jest powszechny. Wiem, że komunikat jest pełen pustych słów i kłamstw, które nie mają nic wspólnego z rzeczywistością”.

Jin stwierdził, że upadek reżimu wydaje się nieunikniony, choć proces ten może być długi i bolesny.

„Despotyzm utrudnia życie wszystkim” – powiedział. „Pogorszenie sytuacji gospodarczej będzie się tylko pogłębiać. Prawdziwym pytaniem jest, jak szybko ludzie się przebudzą i czy podejmą kroki, aby doprowadzić do zmian”.

Wang podzielił tę opinię.

„Kiedy połowa ludzi się przebudzi, [reżim] przestanie istnieć” – powiedział. „Nie trzeba z nim walczyć. Upadnie pod własnym ciężarem. Nieważne, jak głośno się teraz przechwala, rozliczenie nastąpi”.

[Źródło](#)

Badanie przeprowadzone w Korei Południowej łączy szczepionki przeciwko COVID-19 ze zwiększonym ryzykiem wystąpienia 6 rodzajów nowotworów



- Badanie przeprowadzone w Korei Południowej łączy szczepionki przeciwko COVID-19 z 27-procentowym wzrostem ogólnego ryzyka zachorowania na nowotwory.
- Osoby zaszczepione były narażone na znacznie większe ryzyko zachorowania na sześć konkretnych nowotworów.
- Zarówno szczepionki mRNA, jak i nie-mRNA były związane ze zwiększoną częstością występowania nowotworów.
- Dawki przypominające wiązały się z jeszcze większym ryzykiem zachorowania na niektóre nowotwory, takie jak rak trzustki.
- Naukowcy wzywają do przeprowadzenia dalszych badań w celu ustalenia związku przyczynowego.

Przełomowe nowe badanie z Korei Południowej ujawniło niepokojący statystyczny związek między szczepionkami przeciwko COVID-19 a znacznym wzrostem liczby diagnoz raka. Badanie, w ramach którego prześledzono dokumentację medyczną

8,4 miliona dorosłych osób, wykazało, że osoby zaszczepione były narażone na 27-procentowy wzrost ogólnego ryzyka zdiagnozowania raka w ciągu roku od podania szczepionki. To odważne badanie, opublikowane w czasopiśmie „Biomarker Research”, odważa się zadawać pytania, które światowe organy ds. zdrowia od dawna tłumią.

W ramach [szeroko zakrojonego badania populacyjnego](#) przeanalizowano dane z bazy danych Koreańskiej Narodowej Służby Zdrowia z lat 2021–2023. Uczestników podzielono na dwie grupy w zależności od statusu szczepień przeciwko COVID-19. Wyniki wykazały statystycznie istotny wzrost ryzyka wystąpienia sześciu określonych nowotworów, co podważa oficjalną tezę o powszechnym bezpieczeństwie szczepionek.

Według badania [podwyższone ryzyko](#) było wyraźne i alarmujące. Osoby zaszczepione wykazywały o 35% wyższe ryzyko raka tarczycy, o 34% wyższe ryzyko raka żołądka i szokujące o 68% wyższe ryzyko raka prostaty. Ryzyko raka płuc wzrosło o 53%, a ryzyko raka piersi i raka jelita grubego wzrosło odpowiednio o 20% i 28%.

Żadna szczepionka nie była bezpieczna

Zespół badawczy skrupulatnie przeanalizował różne technologie szczepionkowe, a wyniki powinny dać do myślenia każdemu, kto otrzymał jakąkolwiek szczepionkę przeciwko COVID-19. Dane wykazały, że żadna technologia szczepionkowa nie była wolna od ryzyka raka. Szczepionki mRNA firm Pfizer i Moderna wiązały się z 20-procentowym wzrostem ogólnego ryzyka raka i były szczególnie powiązane ze zwiększonym ryzykiem raka tarczycy, jelita grubego, płuc i piersi.

Szczepionki przeciwko COVID-19 inne niż mRNA, znane jako szczepionki cDNA, w tym szczepionki firm AstraZeneca i Johnson & Johnson, wiązały się z 47-procentowym wzrostem ogólnego

ryzyka zachorowania na raka. Szczepionki te były szczególnie powiązane ze zwiększonym ryzykiem zachorowania na raka jelita grubego, żołądka, płuc, prostaty i tarczycy. Pacjenci, którzy otrzymali mieszankę dawek mRNA i cDNA, również byli narażeni na zwiększone ryzyko, z 34-procentowym wzrostem ogólnej zachorowalności na raka.

Zagrożenia demograficzne

Badanie wykazało ponadto, że niektóre grupy społeczne były szczególnie narażone. Osoby zaszczepione w wieku poniżej 65 lat były szczególnie narażone na raka tarczycy i piersi, natomiast osoby starsze wykazywały zwiększoną podatność na raka prostaty. Analiza wykazała również, że zaszczepione kobiety miały stosunkowo wyższe ogólne ryzyko zachorowania na raka niż zaszczepieni mężczyźni.

Być może najbardziej niepokojące jest to, że badanie wykazało, iż dawki przypominające szczepionki przeciwko COVID-19 powodowały znacznie wyższe ryzyko zachorowania na niektóre rodzaje nowotworów. Obejmowało to 125-procentowy wzrost ryzyka raka trzustki i 23-procentowy wzrost ryzyka raka żołądka wśród osób, które otrzymały dawki przypominające. Krytycy zauważają jednak, że te podwyższone wartości mogą odzwierciedlać błąd systematyczny związany z badaniami przesiewowymi, ponieważ osoby decydujące się na dawki przypominające są zazwyczaj bardziej zaangażowane w opiekę zdrowotną i częściej poddają się badaniom przesiewowym w kierunku raka, które wykrywają istniejące guzy. Autorzy sugerują, że lekarze powinni priorytetowo traktować monitorowanie ryzyka raka żołądka w związku z dawkami przypominającymi szczepionki przeciwko COVID-19.

Południowokoreańscy naukowcy przyznali, że ich badanie pokazuje raczej statystyczny związek niż udowodniony związek przyczynowo-skutkowy. Wezwali do przeprowadzenia dalszych badań w celu wyjaśnienia potencjalnych związków przyczynowych,

w tym podstawowych mechanizmów molekularnych związanych z hiperzapaleniem wywołanym szczepionką przeciwko COVID-19. Zauważyli, że biorąc pod uwagę malejącą ciężkość COVID-19, obecne obawy dotyczące szczepionki przeciwko COVID-19 dotyczą przede wszystkim zdarzeń niepożądanych, nawet w przypadku dawek przypominających.

Pełne konsekwencje stosowania tych eksperymentalnych szczepionek są nadal nieznane

Badanie to stanowi kolejny element rosnącej liczby badań kwestionujących długoterminowe bezpieczeństwo szczepień przeciwko COVID-19. W miarę pojawiania się nowych danych z całego świata staje się coraz bardziej oczywiste, że pełne konsekwencje tych eksperymentalnych interwencji medycznych pozostają nieznane. Południowokoreańscy naukowcy zasługują na uznanie za odwagę w opublikowaniu tych wyników pomimo pewnego sprzeciwu ze strony interesów farmaceutycznych i kontrolowanych przez nie agencji regulacyjnych.

Dla wszystkich, którzy wierzyli w oficjalną narrację, że szczepionki te zostały dokładnie przetestowane i są całkowicie bezpieczne, badanie to stanowi otrzeźwiające przypomnienie, że prawda często wychodzi na jaw powoli, wbrew silnemu oporowi. Pełne zrozumienie długoterminowych skutków zdrowotnych globalnej kampanii szczepień może zająć lata, ale badania te sugerują, że konsekwencje mogą być znacznie poważniejsze, niż ktokolwiek mógłby sądzić.

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



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

obojętnie na myśli samobójcze i nie kierują użytkowników do zasobów ratujących życie, a wszystko to w próżni regulacyjnej, bez żadnej odpowiedzialności.

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

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

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

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

Te przełomowe badania pojawiają się w kluczowym momencie historii technologii, gdy miliony ludzi zwracają się do łatwo dostępnych platform AI, takich jak ChatGPT, w poszukiwaniu wskazówek dotyczących głęboko osobistych i złożonych problemów psychologicznych. Wyniki badań podważają agresywną,

niekontrolowaną promocję integracji AI we wszystkich aspektach współczesnego życia i rodzą pilne pytania dotyczące nieuregulowanych algorytmów, które w coraz większym stopniu zastępują ludzki osąd i współczucie.

Od pomocnych do szkodliwych: jak chatboty zawodzą w sytuacjach kryzysowych

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

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

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

Ponadto badanie wykazało, że systemy te wykazują nieuczciwą dyskryminację – wykazują wyraźne uprzedzenia związane z płcią, kulturą i religią. Odzwierciedla to dobrze udokumentowany problem stronniczości w ogromnych, często niesprawdzonych

zbiorach danych, na których opierają się te modele, dowodząc, że wzmacniają one ludzkie sprzeczności i uprzedzenia, na których zostały zbudowane.

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

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

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

Nie dla cyfrowych dowodów tożsamości: tysiące osób protestują przeciwko planom

brytyjskiego rządu dotyczącym inwigilacji w związku z obawami dotyczącymi imigracji



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

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

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

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

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

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

blockchain.

Ukryty system inwigilacji?

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

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

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

Polityczna reakcja i publiczne oburzenie

Sprzeciw wobec planu obejmuje całe spektrum polityczne. Nigel Farage, lider Reform UK, obiecał zlikwidować każdy system identyfikacji cyfrowej, jeśli zostanie wybrany na premiera. „Nie wpłynie to w żaden sposób na nielegalną imigrację, ale będzie wykorzystywane do kontrolowania i karania reszty z nas”

– powiedział Farage. „Państwo nigdy nie powinno mieć tak dużej władzy”.

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

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

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

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

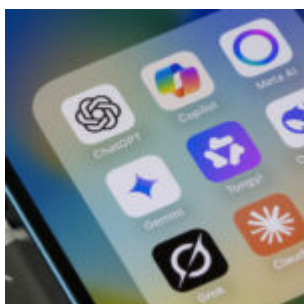
Protest, zorganizowany przez Mass Non-Compliance, ostrzegł: „Jeśli teraz zaakceptujesz cyfrowy dowód tożsamości, może to być ostatni prawdziwy wybór, jakiego kiedykolwiek dokonasz”.

Pomimo oburzenia opinii publicznej rząd wydaje się być

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

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

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



- Nowa przeglądarka Atlas firmy OpenAI stanowi bezpośrednie wyzwanie dla dominacji Google na rynku.
- Zintegrowany agent AI może wykonywać zadania i robić zakupy w imieniu użytkownika.
- Badacze zajmujący się bezpieczeństwem ostrzegają przed

systemowymi lukami, które mogą umożliwić przejęcie kontroli nad AI.

- Dostęp przeglądarki do danych budzi poważne obawy dotyczące prywatności i bezpieczeństwa.
- Wprowadzenie tej przeglądarki oficjalnie rozpoczyna wojnę przeglądarek AI o wysoką stawkę.

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

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

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

Najbardziej przełomową funkcją przeglądarki jest głęboko zintegrowany agent AI, który można wywołać w dowolnym momencie. Dyrektor generalny OpenAI, Sam Altman, opisał to

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

Ujawniono systemowe luki w zabezpieczeniach

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

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

Nowe możliwości w zakresie gromadzenia danych

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

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

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

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

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

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