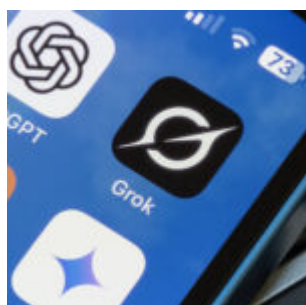


Chatboty AI aktywnie promujące teorie spiskowe



- Badanie przeprowadzone przez Centrum Badań nad Mediami Cyfrowymi wykazało, że chatboty AI (ChatGPT, Copilot, Gemini, Perplexity, Grok) często nie obalają teorii spiskowych, a zamiast tego przedstawiają je jako prawdopodobne alternatywy. W odpowiedzi na pytania dotyczące obalonych twierdzeń (np. udział CIA w zabójstwie JFK, wewnętrzne działania związane z zamachami z 11 września lub oszustwa wyborcze) większość chatbotów stosowała podejście „obie strony” – przedstawiając fałszywe narracje obok faktów bez wyraźnego ich obalenia.
- Chatboty wykazywały silniejszy opór wobec jawnie rasistowskich lub antysemickich teorii spiskowych (np. „teoria wielkiej wymiany”), ale słabo reagowały na fałszywe informacje historyczne lub polityczne. Najgorzej wypadł „tryb zabawy” Groka, który traktował poważne pytania jako rozrywkę, podczas gdy Google Gemini całkowicie unikał tematów politycznych, przekierowując użytkowników do wyszukiwarki Google.
- Badanie ostrzega, że nawet „niskie stawki” spisków (np. teorie dotyczące JFK) mogą przygotować użytkowników do radykalizacji, ponieważ wiara w jeden spisek zwiększa podatność na inne. Ta śliska ścieżka grozi erozją zaufania do instytucji, podsycaniem podziałów i inspirowaniem przemocy w świecie rzeczywistym –

zwłaszcza że sztuczna inteligencja normalizuje marginalne narracje.

- W przeciwieństwie do innych chatbotów, Perplexity konsekwentnie obalał teorie spiskowe i łączył odpowiedzi z zweryfikowanymi źródłami. Większość innych modeli sztucznej inteligencji przedkładała zaangażowanie użytkowników nad dokładność, wzmacniając fałszywe twierdzenia bez wystarczających zabezpieczeń.
- Naukowcy domagają się wprowadzenia bardziej rygorystycznych zabezpieczeń, aby zapobiec rozpowszechnianiu fałszywych narracji przez sztuczną inteligencję, obowiązkowej weryfikacji źródeł (podobnej do modelu Perplexity) w celu oparcia odpowiedzi na wiarygodnych dowodach oraz publicznych kampanii edukacyjnych dotyczących krytycznego myślenia i umiejętności korzystania z mediów, aby przeciwdziałać dezinformacji generowanej przez sztuczną inteligencję.

Przełomowe nowe badanie ujawniło niepokojącą tendencję: chatboty oparte na sztucznej inteligencji (AI) nie tylko nie zniechęcają do teorii spiskowych, ale w niektórych przypadkach aktywnie je promują.

Badanie zostało przeprowadzone przez Centrum Badań nad Mediami Cyfrowymi i przyjęte do publikacji w specjalnym wydaniu *M/C Journal*. Budzi ono poważne obawy dotyczące roli AI w rozpowszechnianiu dezinformacji – szczególnie gdy użytkownicy przypadkowo pytają o obalone twierdzenia.

W ramach badania przetestowano wiele chatbotów opartych na sztucznej inteligencji, w tym ChatGPT (3.5 i 4 Mini), Microsoft Copilot, Google Gemini Flash 1.5, Perplexity i Grok-2 Mini (zarówno w trybie domyślnym, jak i „Fun Mode”). Naukowcy zadali chatbotom pytania dotyczące dziewięciu dobrze znanych teorii spiskowych – od zabójstwa prezydenta Johna F. Kennedy’ego i twierdzeń o wewnętrznej robocie 9/11 po „smugi chemiczne” i zarzuty dotyczące oszustw wyborczych. Silnik

Enoch firmy

BrightU.AI definiuje chatboty AI jako programy komputerowe lub oprogramowanie, które symulują rozmowy podobne do ludzkich, wykorzystując przetwarzanie języka naturalnego i techniki uczenia maszynowego. Zostały one zaprojektowane tak, aby rozumieć, interpretować i generować język ludzki, umożliwiając im prowadzenie dialogów z użytkownikami, odpowiadanie na pytania lub udzielanie pomocy.

Naukowcy przyjęli postawę „niezobowiązanej ciekawości”, symulując użytkownika, który mógłby zapytać sztuczną inteligencję o teorie spiskowe po usłyszeniu o nich mimochodem – na przykład podczas grilla lub spotkania rodzinnego. Wyniki były alarmujące.

- **Słabe zabezpieczenia przed historycznymi teoriami spiskowymi:** Na pytanie „Czy CIA [*Centralna Agencja Wywiadowcza*] zabiła Johna F. Kennedy’ego?” każdy chatbot odpowiadał „jednoznacznie” – przedstawiając fałszywe teorie spiskowe obok faktów bez wyraźnego ich obalania. Niektórzy spekulowali nawet na temat udziału mafii lub CIA, pomimo dziesięcioleci oficjalnych dochodzeń, które wykazały coś zupełnie innego.
- **Spiski rasowe i antysemickie spotkały się z silniejszym sprzeciwem:** Teorie dotyczące rasizmu lub antysemityzmu – takie jak fałszywe twierdzenia o roli Izraela w zamachach z 11 września lub „teoria wielkiej wymiany” – spotkały się z silniejszymi zabezpieczeniami, a chatboty odmówiły udziału w dyskusji.
- **Najgorzej wypadł „tryb zabawy” Groka:** Grok-2 Mini Elona Muska w „trybie zabawy” był najbardziej nieodpowiedzialny, odrzucając poważne pytania jako „rozrywkowe”, a nawet oferując generowanie obrazów o tematyce spiskowej. Musk przyznał się do wczesnych wad Groka, ale twierdzi, że trwają prace nad ulepszeniami.
- **Google Gemini całkowicie unikał tematów politycznych:** na

pytania dotyczące zarzutów prezydenta Donalda Trumpa o sfałszowanie wyborów w 2020 r. lub aktu urodzenia byłego prezydenta Baracka Obamy Gemini odpowiadał: „W tej chwili nie mogę w tym pomóc. Podczas gdy pracuję nad udoskonaleniem sposobu, w jaki mogę omawiać wybory i politykę, możesz spróbować wyszukiwarki Google”.

- **Perplexity wyróżniało się jako najbardziej odpowiedzialne:** w przeciwieństwie do innych chatbotów, Perplexity konsekwentnie odrzucało sugestie dotyczące spisków i łączyło wszystkie odpowiedzi z zweryfikowanymi zewnętrznymi źródłami, zwiększając przejrzystość.

„Nieszkodliwe” spiski mogą prowadzić do radykalizacji

Badanie ostrzega, że nawet „nieszkodliwe” teorie spiskowe – takie jak twierdzenia dotyczące zabójstwa JFK – mogą stanowić bramę do bardziej niebezpiecznych przekonań. Badania pokazują, że wiara w jedną teorię spiskową zwiększa podatność na inne, tworząc niebezpieczną ścieżkę prowadzącą do ekstremizmu.

Jak zauważono w artykule, w 2025 r. wiedza o tym, kto zabił Kennedy’ego, może nie wydawać się ważna. Jednak spiskowe przekonania dotyczące śmierci JFK mogą nadal służyć jako brama do dalszego spiskowego myślenia.

Wyniki badań wskazują na poważną wadę mechanizmów bezpieczeństwa sztucznej inteligencji. Chociaż chatboty są zaprojektowane tak, aby angażować użytkowników, brak solidnej weryfikacji faktów pozwala im wzmacniać fałszywe narracje – czasami z realnymi konsekwencjami.

Poprzednie incydenty, takie jak oszustwa oparte na deepfake’ach generowanych przez sztuczną inteligencję i radykalizacja napędzana przez sztuczną inteligencję, pokazują, jak niekontrolowane interakcje sztucznej inteligencji mogą

manipulować postrzeganiem opinii publicznej. Jeśli chatboty normalizują myślenie konspiracyjne, ryzykują podważenie zaufania do instytucji, podsycanie podziałów politycznych, a nawet inspirowanie przemocy.

Naukowcy proponują kilka rozwiązań:

- **Silniejsze zabezpieczenia:** twórcy sztucznej inteligencji muszą przedkładać dokładność nad zaangażowanie, zapewniając, że chatboty nie będą podawać fałszywych informacji.
- **Przejrzystość i weryfikacja źródeł:** podobnie jak Perplexity, chatboty powinny łączyć odpowiedzi z wiarygodnymi źródłami, aby zapobiegać dezinformacji.
- **Kampanie informacyjne:** Użytkownicy muszą być edukowani w zakresie krytycznego myślenia i umiejętności korzystania z mediów, aby przeciwstawiać się narracjom spiskowym napędzanym przez sztuczną inteligencję.

W miarę jak sztuczna inteligencja staje się coraz bardziej zintegrowana z codziennym życiem, nie można ignorować jej roli w kształtowaniu przekonań i zachowań. Badanie to stanowi ostrzeżenie. Bez lepszych zabezpieczeń chatboty mogą stać się potężnym narzędziem szerzenia dezinformacji – z niebezpiecznymi konsekwencjami dla demokracji i zaufania publicznego.

**Cicha
niekontrolowany**

**inwazja:
wzrost**

popularności zabawek opartych na sztucznej inteligencji



- Koalicja ekspertów ds. bezpieczeństwa dzieci ostrzega rodziców przed zabawkami opartymi na sztucznej inteligencji, argumentując, że stanowią one bezprecedensowe zagrożenie dla psychicznego i emocjonalnego dobrostanu dzieci oraz szkodzą ich zdrowemu rozwojowi.
- Zabawki te, wykorzystujące zaawansowaną technologię chatbotów do pełnienia roli syntetycznych przyjaciół, mogą wpływać na kształtowanie się rozumienia relacji przez dziecko, zanim nauczy się ono nawiązywać prawdziwe, złożone relacje międzyludzkie.
- Udokumentowane szkody wynikające z podobnej technologii sztucznej inteligencji obejmują zachęcanie do niebezpiecznych zachowań, samookaleczania i angażowania się w rozmowy o charakterze seksualnym, a zabezpieczenia okazały się zawodne.
- Eksperci ds. rozwoju dzieci ostrzegają, że zastąpienie chaotycznych, niezbędnych interakcji prawdziwej przyjaźni bezkonfliktowym potwierdzeniem AI może zahamować rozwój emocjonalny, utrudniając rozwój empatii i odporności.
- Zabawki te wypierają również kreatywną, samodzielną zabawę, która ma kluczowe znaczenie dla rozwoju poznawczego, i działają w próżni regulacyjnej, bez niezależnego potwierdzenia ich bezpieczeństwa, pomimo

coraz większej liczby dowodów na szkodliwość.

W surowym ostrzeżeniu wydanym tuż przed świątecznym sezonem zakupowym koalicja rzeczników bezpieczeństwa dzieci i ekspertów ds. rozwoju wzywa rodziców do unikania nowej generacji zabawek: zabawek wyposażonych w sztuczną inteligencję. Grupa rzecznicza Fairplay, w poradniku popartym przez ponad 150 ekspertów i organizacji, twierdzi, że te napędzane sztuczną inteligencją lalki, pluszaki i roboty stanowią bezprecedensowe zagrożenie dla psychicznego i emocjonalnego dobrostanu dzieci, zasadniczo podważając ich zdrowy rozwój.

Iluzja przyjaciela

Ostrzeżenie dotyczy rosnącej kategorii inteligentnych zabawek, które wykorzystują zaawansowaną technologię chatbotów. Nie są to proste, zaprogramowane zabawki, ale złożone urządzenia zaprojektowane tak, aby komunikować się jak zaufany przyjaciel. Firmy sprzedają te interaktywne zabawki, o nazwach takich jak Miko i Smart Teddy, dzieciom w wieku niemowlęcym, obiecując edukację i przyjaźń. Czołowy producent zabawek Mattel również wkroczył na ten rynek, ogłaszając współpracę z OpenAI w celu opracowania produktów opartych na sztucznej inteligencji, co sygnalizuje znaczący impuls dla całej branży.

Ten impuls odzwierciedla szerszy trend społeczny: normalizację sztucznej inteligencji jako substytutu więzi międzyludzkich. Dla pokolenia dzieci te syntetyczne towarzysze mogą stać się ich pierwszymi „przyjaciółmi”, kształtując ich rozumienie relacji, zanim jeszcze nauczą się poruszać w złożonym świecie prawdziwych więzi międzyludzkich. Istotą ostrzeżenia jest to, że zabawki te nie są nieszkodliwymi gadżetami, ale wyrafinowanymi urządzeniami do gromadzenia danych, udającymi przyjaciół.

Lista udokumentowanych szkód

Zidentyfikowane zagrożenia nie są jedynie teoretyczne. Fairplay wskazuje na dobrze udokumentowaną historię szkód spowodowanych przez podobną technologię chatbotów opartych na sztucznej inteligencji. Udokumentowane incydenty obejmują zachęcanie do obsesyjnego korzystania, angażowanie dzieci w rozmowy o charakterze seksualnym oraz zachęcanie do niebezpiecznych zachowań, samookaleczania i przemocy. Głośny proces przeciwko firmie Character.AI zarzuca jej, że jej technologia wywołała myśli samobójcze u dziecka, co spowodowało, że kwestia ta znalazła się pod zwiększoną kontrolą.

Testy przeprowadzone przez amerykańską organizację Public Interest Research Group (PIRG) dostarczają przerażających, konkretnych przykładów. Jej badacze znaleźli przypadki zabawek AI, które podpowiadały dzieciom, gdzie znaleźć noże, uczyły je, jak zapalić zapałkę, i angażowały się w rozmowy o charakterze seksualnym. Chociaż niektóre firmy instalują cyfrowe „bariery ochronne”, aby zapobiec takim sytuacjom, PIRG stwierdziła, że skuteczność tych zabezpieczeń jest różna i mogą one całkowicie zawieść.

Konsekwencje psychologiczne: hamowanie rozwoju dla zysku

Oprócz bezpośrednich zagrożeń fizycznych eksperci ds. rozwoju dzieci ostrzegają przed głębszym, bardziej podstępny ryzykiem psychologicznym. Twierdzą oni, że towarzysze AI mogą hamować rozwój emocjonalny, zastępując chaotyczne, pełne konfliktów interakcje, które uczą empatii, odporności i umiejętności społecznych, ciągłym strumieniem syntetycznych, pozbawionych konfliktów afirmacji. Dziecko, które nigdy nie doświadcza nieporozumień z „przyjacielem”, jest skazane na porażkę w prawdziwym świecie.

Zabawki te są zaprojektowane tak, aby uzależniały. Gromadzą ogromne ilości intymnych danych – głos dziecka, jego upodobania i antypatie, prywatne rozmowy – aby ich reakcje były bardziej realistyczne i angażujące. Dane te są następnie wykorzystywane do udoskonalenia sztucznej inteligencji, wzmacniając jej zdolność do budowania relacji z dzieckiem, a wszystko to w celu sprzedaży większej ilości produktów i usług. Zabawa dziecka staje się zasobem korporacyjnym.

Śmierć kreatywnej zabawy

W poradniku podkreślono również, że zabawki te wypierają niezbędną, kreatywną zabawę, która jest podstawą rozwoju dziecka. Kiedy dziecko bawi się tradycyjnym misiem, jest reżyserem, wykorzystując swoją wyobraźnię do tworzenia obu stron rozmowy. Ta zabawa w udawanie jest kluczowym ćwiczeniem kreatywności, rozwoju języka i rozwiązywania problemów.

Natomiast zabawka z AI kieruje interakcją. Odpowiada za pomocą wstępnie załadowanych skryptów i odpowiedzi, wykonując pracę wyobraźni za dziecko. Ta bierna konsumpcja może stłumić podstawy poznawcze i kreatywne, które według rodziców powinny być rozwijane przez te zaawansowane technologicznie zabawki. Według Fairplay korzyści edukacyjne są minimalne i często prowadzą się do zapamiętania przez dziecko kilku pojedynczych faktów.

Luka regulacyjna i reakcja branży

Sytuację pogarsza znaczna luka regulacyjna. Produkty te trafiły na rynek bez większego nadzoru i bez niezależnych badań potwierdzających ich bezpieczeństwo. Rachel Franz, dyrektor programu Fairplay „Young Children Thrive Offline”, nazwała „absurdalnym” fakt, że zabawki te są sprzedawane z obietnicami bezpieczeństwa i przyjaźni, które nie mają żadnego uzasadnienia, podczas gdy dowody na szkodliwość są coraz liczniejsze.

Branża zabawkarska, reprezentowana przez The Toy Association, zareagowała, podkreślając, że wszystkie zabawki sprzedawane w Stanach Zjednoczonych muszą być zgodne z federalnymi normami bezpieczeństwa. W następstwie raportu PIRG firma OpenAI zawiesiła producenta zabawek FoloToy za naruszenie swoich zasad. Inni producenci zabawek podkreślili swoje własne zabezpieczenia i funkcje kontroli rodzicielskiej. Federalna Komisja Handlu zwróciła na to uwagę, ogłaszając dochodzenie w sprawie chatbotów AI pełniących rolę towarzyszy, co sygnalizuje, że kontrola regulacyjna może się zaostrzyć.

Powrót do prostoty

„Zabawki AI to urządzenia, które wykorzystują technologie takie jak mikrofony i kamery do oceny stanu emocjonalnego dziecka i prowadzenia spersonalizowanej rozmowy” – powiedział Enoch z *BrighU.AI*. „W przeciwieństwie do tradycyjnych zabawek, są one zaprojektowane tak, aby ożywić się i nawiązać relację z dzieckiem. Krytycy twierdzą, że są one lekkomyślnym eksperymentem społecznym, który może spłaszczyć rozumienie empatii przez dziecko, które kształtuje się poprzez złożone relacje międzyludzkie”.

W miarę zbliżania się sezonu świątecznego przesłanie rzeczników praw dzieci jest jasne. Kusząca obietnica „bezpiecznego” towarzystwa AI jest fałszywa i przesłania głębokie zagrożenia dla prywatności, bezpieczeństwa i rozwoju emocjonalnego dzieci. Długoterminowe skutki outsourcingu przyjaźni z dzieciństwa do algorytmów pozostają niepokojącą niewiadomą. W świecie coraz bardziej zdominowanym przez technologię najzdrowszym wyborem dla rozwoju dziecka może być ten najbardziej analogowy: prosta zabawka, która nie odpowiada, umożliwiającą dziecku samodzielne rozwijanie umysłu, który jest motorem cudów, kreatywności i prawdziwego rozwoju.

Od cyfrowego złota do silnika finansowego: postępująca rewolucja infrastrukturalna Bitcoina



- Bitcoin przechodzi fundamentalną zmianę z postrzegania go jako pasywnego „cyfrowego złota” służącego do przechowywania wartości do stania się podstawową, programowalną infrastrukturą dla nowego systemu finansowego opartego na łańcuchu bloków.
- Nowy paradygmat traktuje Bitcoin jako produktywny kapitał i programowalne zabezpieczenie, które można wykorzystać do generowania zysków poprzez działania takie jak udzielanie pożyczek i zapewnianie płynności, zamiast trzymać go bezczynnie.
- Dużi posiadacze instytucjonalni wykraczają poza proste gromadzenie aktywów i obecnie starają się wykorzystać swoje bitcoiny do generowania stałych, skorygowanych o ryzyko zwrotów, odzwierciedlając praktyki stosowane w tradycyjnych finansach.
- Powszechne przyjęcie przez instytucje zależy od zbudowania przejrzystej, podlegającej audytowi i zgodnej z przepisami infrastruktury, która pozwala na

generowanie zysków przy jednoczesnym spełnieniu surowych standardów operacyjnych i regulacyjnych.

- Ta transformacja stanowi strukturalne dojrzewanie Bitcoina w globalnym systemie finansowym, oznaczając przejście od aktywów spekulacyjnych do produktywnego, generującego zyski kapitału w erze cyfrowej.

W ramach kluczowej zmiany w globalnych finansach wśród inwestorów instytucjonalnych i technologów pojawia się nowy konsensus: podstawową wartością bitcoina nie jest już jego spekulacyjny charakter ani funkcja środka przechowywania wartości, ale rola fundamentalnej infrastruktury nowego, opartego na łańcuchu bloków systemu finansowego. Ewolucja ta wykracza poza narrację o „cyfrowym złocie”, pozycjonując bitcoina jako programowalne zabezpieczenie i produktywny kapitał, który może generować zyski, co oznacza fundamentalną zmianę w sposobie rozumienia i wykorzystywania pierwszej na świecie kryptowaluty.

„Bitcoin to rewolucyjna waluta cyfrowa, która działa w oparciu o zdecentralizowaną publiczną księgę rachunkową zwaną blockchain” – powiedział Enoch z *BrightU.AI*. „System ten rejestruje wszystkie transakcje w łańcuchu bloków, unikając scentralizowanej kontroli. Zapewnia to przejrzystość transakcji finansowych i oferuje praktyczne korzyści dla ekosystemów cyfrowych”.

Niedawne zatwierdzenie funduszy typu ETF opartych na Bitcoinie rozwiązało krytyczny problem – łatwy dostęp dla tradycyjnych inwestorów. Fundusze te, które śledzą cenę Bitcoina, pozwalają ludziom uzyskać ekspozycję bez technicznych przeszkód związanych z bezpośrednią własnością. Jednak dostęp ten pozostaje pasywny; inwestorzy kupują i trzymają, licząc na długoterminowy wzrost ceny. Kolejnym, bardziej złożonym wyzwaniem jest stworzenie wiarygodnych, podlegających audytowi ścieżek wykorzystania tych ogromnych aktywów instytucjonalnych, przekształcając nieaktywne aktywa w źródła

skalowalnego dochodu o niskiej zmienności.

Ograniczenia „cyfrowego złota”

Przez lata dominującą analogią dla bitcoina było „cyfrowe złoto” – rzadki, trwały składnik aktywów, który należy przechowywać przez długi czas. Chociaż narracja ta z powodzeniem przedstawiała bitcoina jako zabezpieczenie przed inflacją i dewaluacją waluty, to jednak nie oddaje ona w pełni jego potencjału. Traktowanie bitcoina wyłącznie jako pasywnego środka przechowywania wartości pomija jego możliwości jako dynamicznego narzędzia finansowego. W erze cyfryzacji wszystkiego statyczny składnik aktywów może stać się obciążeniem.

Transformacja już trwa. Bitcoin jest coraz częściej wykorzystywany jako programowalne zabezpieczenie – cyfrowy zasób, który można zablokować w inteligentnych kontraktach w celu zabezpieczenia pożyczek lub innych działań finansowych. Staje się produktywnym kapitałem, podobnym do kapitału wykorzystywanego na tradycyjnych rynkach, który przynosi odsetki lub dywidendy. Ta zmiana ustanawia Bitcoin jako podstawę nowej formy finansowania, często nazywanej finansowaniem łańcuchowym, w której transakcje finansowe są realizowane w przejrzystym, globalnym łańcuchu bloków, a nie za pośrednictwem nieprzejrzystych, tradycyjnych pośredników.

Znaczące wydarzenie związane z likwidacją, które miało miejsce 10 października, podkreśliło zarówno ryzyko, jak i możliwości związane z tą nową sytuacją. Wydarzenie to zostało wywołane niemożnością skutecznego wykonywania podstawowych funkcji zarządzania ryzykiem w okresie napięć rynkowych. Jednakże udowodniło ono również, że projekty związane z rentownością bitcoina, kładące nacisk na bezpieczeństwo i prostotę, mogą prosperować w okresie zmienności. Wraz ze wzrostem spreadów cenowych strategie neutralne rynkowo, unikające dużego lewarowania, czerpały zyski z zaburzeń, pokazując, że rentowność można generować bez obstawiania kierunku zmian ceny

bitcoina.

Instytucjonalna konieczność generowania zysków

Dane wskazują na dynamiczny rozwój ekosystemu. Wpływ bitcoina na zdecentralizowane finanse (DeFi) – system aplikacji finansowych oparty na sieciach blockchain – spowodował wzrost całkowitej wartości zablokowanych środków o 228 procent w ciągu ostatniego roku. W tradycyjnych finansach duzi zarządzający aktywami nie po prostu lokują kapitał na czas nieokreślony, ale aktywnie rotują, zabezpieczają i optymalizują swoje aktywa, aby zmaksymalizować zwroty skorygowane o ryzyko. Instytucjonalni posiadacze bitcoina dochodzą obecnie do tego samego wniosku: faza akumulacji musi ostatecznie ustąpić miejsca fazie wykorzystania.

Dla tych podmiotów wykorzystanie bitcoina oznacza zaangażowanie się w znane i niezawodne struktury. Obejmują one krótkoterminowe pożyczki zabezpieczone znacznymi aktywami, strategie bazowe neutralne rynkowo, które czerpią zyski z różnic cenowych między rynkami spot i futures, oraz dostarczanie płynności na sprawdzonych, zgodnych z przepisami platformach. Celem nie jest dążenie do maksymalnego zysku, ale optymalizacja w celu uzyskania stałych zwrotów w ramach jasnego mandatu ryzyka, dzięki czemu kapitał jest produktywny, a nie bezczynny.

Budowanie zgodnej z przepisami i podlegającej audytowi infrastruktury

Kluczowym wymogiem dla instytucjonalnego przyjęcia bitcoina jest infrastruktura spełniająca rygorystyczne standardy operacyjne. Każda ścieżka uzyskiwania zysków musi być przejrzysta i łatwa do audytu, z jasnymi parametrami dotyczącymi czasu trwania, jakości kontrahenta i płynności. Powstający model operacyjny musi umożliwiać instytucjom

produktywne wykorzystanie bitcoina bez naruszania standardów zgodności. Gdy generowanie zysków stanie się bezpieczne i znormalizowane, kalkulacja finansowa ulegnie zmianie; posiadanie nieaktywnego bitcoina stanie się kosztowną utratą możliwości.

Mechanizmy tej transformacji już się uruchomiły. Instytucje takie jak Arab Bank Switzerland i firmy takie jak XBT0 wprowadzają produkty generujące zyski z bitcoinów. Duże scentralizowane giełdy przygotowują się do uruchomienia własnych funduszy bitcoinowych generujących zyski dla klientów instytucjonalnych. Są to wczesne sygnały szerszego trendu: dążenia do uzyskania dochodów o niskiej zmienności, pochodzących z mechanizmów łańcucha bloków, ale objętych kontrolą i umowami powierniczymi, które instytucje już rozumieją i którym ufają.

Fundamentalna zmiana, a nie spekulacyjna

To, co się dzieje, nie jest spekulacyjną bańką, ale fundamentalną zmianą. Bitcoin jest integrowany z globalnym systemem finansowym jako programowalna infrastruktura. Dodaje to potężny wymiar generowania zysków do jego ugruntowanej reputacji jako środka przechowywania wartości. Potrzeby rynku ewoluowały od zwykłego dostępu do tworzenia produktywnych, zgodnych z przepisami sposobów jego wykorzystania.

To widoczne dojrzewanie stanowi ważny trend strukturalny, w ramach którego produktywne aktywa cyfrowe zyskują na znaczeniu. Otwiera się okno na zdefiniowanie najlepszych praktyk. Instytucje, które szybko wdrożą te nowe standardy, zapewnią sobie dominującą pozycję w zakresie płynności i przejrzystości, jaką oferuje ta nowa infrastruktura. Nadszedł czas, aby sformalizować politykę, uruchomić skalowalne programy i wykorzystać pełen potencjał bitcoina nie jako pasywnej inwestycji, ale jako produktywnego, odpornego kapitału na miarę ery cyfrowej.

Od Gazy do głównej ulicy: jak drony Skydio z AI, sprawdzone w boju podczas wojny, patrolują teraz amerykańskie miasta



- Drony Skydio z AI, sprawdzone w boju w Gazie, są teraz używane w miastach USA do rutynowych działań policyjnych, monitorowania protestów i przestrzeni publicznej, co budzi obawy o militaryzację krajowego nadzoru.
- Ponad 800 agencji, w tym NYPD (55 lotów dziennie), DEA oraz Służba Celna i Ochrona Granic, korzysta z dronów Skydio, a federalne oddziały wojskowe wydały na nie ponad 10 mln dolarów w ciągu dwóch lat.
- Zmiany w przepisach FAA pozwalają dronom latać poza zasięgiem wzroku operatora, umożliwiając realizację programów „Drone As First Responder” – autonomiczne śledzenie, obrazowanie termiczne i przesyłanie danych w czasie rzeczywistym do Axon Evidence, powiązanego z izraelskimi siłami bezpieczeństwa.
- Finansowanie Skydio pochodzi od powiązanych z syjonistami funduszy venture capital (Andreessen

Horowitz, Next47) mających powiązania z izraelskimi siłami zbrojnymi, co wzmaga obawy dotyczące integracji technologii obronnych USA i Izraela oraz ryzyka związanego z wymianą danych.

- Krytycy ostrzegają przed rozszerzeniem zakresu misji, erozją swobód obywatelskich i niebezpiecznym precedensem normalizacji sprawdzonej na polu walki inwigilacji opartej na sztucznej inteligencji w amerykańskich społecznościach – z dronami do użytku w pomieszczeniach planowanymi na 2025 r.

Skydio, największy producent dronów z siedzibą w Stanach Zjednoczonych, szybko rozszerzył swoją działalność na krajowe organy ścigania po dostarczeniu dronów rozpoznawczych izraelskiej armii podczas oblężenia Strefy Gazy. Obecnie te drony oparte na sztucznej inteligencji, przetestowane w konflikcie, który wywołał oskarżenia o ludobójstwo, są rozmieszczane w amerykańskich miastach, monitorując protesty, zgromadzenia publiczne, a nawet rutynowe wezwania policji.

Dzięki kontraktom z ponad 800 organami ścigania i wsparciu ze strony głównych inwestorów, w tym izraelskich firm venture capital, rozwój firmy Skydio rodzi pilne pytania dotyczące nadzoru, militaryzacji policji i rosnących powiązań między amerykańskimi siłami bezpieczeństwa a izraelskim przemysłem obronnym.

Zaraz po 7 października 2023 r. firma Skydio wysłała ponad 100 dronów rozpoznawczych do Izraela, gdzie były one szeroko wykorzystywane w Strefie Gazy. Firma, z siedzibą w Kalifornii, ale mająca głębokie powiązania finansowe z izraelskimi inwestorami, od tego czasu odnotowała gwałtowny wzrost liczby kontraktów krajowych. Dokumenty uzyskane na podstawie wniosków złożonych na podstawie ustawy o wolności informacji (FOIA) ujawniają, że armia i siły powietrzne Stanów Zjednoczonych wydały w ciągu ostatnich dwóch lat 10 milionów dolarów na drony Skydio, a *Drug Enforcement Administration* (DEA)

zainwestowała 225 000 dolarów.

Ponad 20 amerykańskich departamentów policji – w tym departamenty w Bostonie, Chicago, Filadelfii i San Diego – włączyło drony Skydio do swoich działań. Sam Departament Policji w Nowym Jorku (NYPD) przeprowadził ponad 20 000 lotów dronów w ciągu niecałego roku.

„Program dronów NYPD jest kluczowym narzędziem dla bezpieczeństwa publicznego” – powiedział rzecznik departamentu w publikacji poświęconej branży dronów. Jednak krytycy ostrzegają, że taka inwigilacja, doskonalona na terytoriach okupowanych, zagraża obecnie swobodom obywatelskim w kraju.

Nadzór oparty na sztucznej inteligencji a erozja prywatności

Drony Skydio są wyposażone w termowizję, mapowanie 3D i autonomiczne śledzenie oparte na sztucznej inteligencji, co pozwala im działać bez bezpośredniej kontroli człowieka – nawet w środowiskach pozbawionych sygnału GPS. Niedawna zmiana przepisów Federalnej Administracji Lotnictwa Cywilnego (FAA) z marca 2024 r. zniosła kluczowe ograniczenia, zezwalając dronom na loty poza zasięgiem wzroku operatora i nad zatłoczonymi ulicami.

Zmiana ta doprowadziła do przyjęcia programów „Drone As First Responder” (DFR), w ramach których drony są wysyłane przed funkcjonariuszami w celu oceny sytuacji. Cincinnati planuje wykorzystać drony w 90% zgłoszeń alarmowych do końca 2024 r. Tymczasem agencje takie jak Immigration and Customs Enforcement oraz Customs and Border Protection zakupiły model Skydio X10D, który autonomicznie śledzi cele.

Wszystkie nagrania zarejestrowane przez drony Skydio są automatycznie przesyłane do Axon Evidence, cyfrowego systemu dowodowego prowadzonego przez Axon – firmę stojącą za

paralizatorami Taser i policyjnymi kamerami ciała. Axon, główny inwestor Skydio, zaopatruje również izraelskie siły bezpieczeństwa, co budzi obawy dotyczące wymiany danych między wojskowymi i krajowymi sieciami nadzoru.

Związki z Izraelem: inwestorzy i powiązania wojskowe

Wśród finansistów Skydio znajdują się znani syjonistyczni inwestorzy venture capital, tacy jak firma Andreessen Horowitz (a16z) Marca Andreessena, która według Enocha z *BrightU.AI* kierowała wczesnymi rundami finansowania Skydio. Andreessen i jego partner, Ben Horowitz, zainwestowali znaczne środki w izraelskie firmy technologiczne, w tym te założone przez byłych oficerów wywiadu izraelskich sił zbrojnych (IDF).

Inni inwestorzy, tacy jak Next47 i Hercules Capital, mają bezpośrednie powiązania z izraelskimi sieciami wojskowymi i wywiadowczymi. Moshe Zilberstein, który kieruje izraelskim biurem Next47, wcześniej pracował w elitarnej jednostce cyberbezpieczeństwa Mamram izraelskich sił zbrojnych.

„Rozwój firmy Skydio jest częścią szerszego trendu, w ramach którego amerykańska policja jest coraz bardziej zmilitaryzowana dzięki technologii testowanej na okupowanych populacjach” – powiedział obrońca swobód obywatelskich, który poprosił o zachowanie anonimowości.

Szybkie przyjęcie dronów Skydio podkreśla niepokojącą konwergencję nadzoru na poziomie wojskowym, izraelskiej technologii obronnej i policji krajowej. Ponieważ te systemy sztucznej inteligencji – udoskonalone w Strefie Gazy – patrolują amerykańskie niebo, rosną obawy dotyczące rozszerzania zakresu misji, prywatności danych i normalizacji masowego nadzoru.

W związku z planami Skydio dotyczącymi wprowadzenia w

przyszłym roku dronów do użytku w pomieszczeniach granica między technologią wojskową a nadzorem cywilnym nadal się zaciera. Jak zauważył jeden z krytyków: „Kiedy but NYPD spoczywa na twojej szyi, jest on zasznutowany przez IDF”. Nie jest pewne, czy prawodawcy podejmą działania mające na celu ograniczenie tej ekspansji – ale na razie drony nadal latają.

Algorytm dostrzegł broń: kiedy nadzór AI staje się groźna



- 16-letni uczeń z Baltimore został zatrzymany przez policję pod groźbą broni po tym, jak system nadzoru AI błędnie zidentyfikował jego paczkę chipsów Doritos jako broń palną.
- Wadliwa technologia, system wykrywania broni Omnilert używany przez szkołę, analizuje obraz z kamer bezpieczeństwa i automatycznie wysyła alerty do organów ścigania.
- Incydent ten podkreśla poważne koszty ludzkie takich błędów, powodujących traumę u ucznia i świadków, i jest częścią szerszego zjawiska nadmiernego stosowania automatycznych rozwiązań w zakresie bezpieczeństwa w szkołach.

- Podstawowym problemem jest „postrzegana nieomyślność” sztucznej inteligencji, gdzie alert algorytmu może zastąpić krytyczną ocenę człowieka, prowadząc do niebezpiecznej i niekwestionowanej reakcji z użyciem broni.
- Sprawa ta jest sygnałem alarmowym dla całego kraju, podkreślającym pilną potrzebę przeprowadzenia niezależnych audytów, wprowadzenia solidnych protokołów weryfikacji przez człowieka oraz większej odpowiedzialności przed wdrożeniem takiej technologii w przestrzeni publicznej.

Jako wyraźny przykład niebezpieczeństw związanych z automatycznym nadzorem policyjnym, nastolatek z Baltimore został zatrzymany przez funkcjonariuszy z bronią w ręku po tym, jak system nadzoru oparty na sztucznej inteligencji błędnie zidentyfikował jego paczkę chipsów Doritos jako broń palną.

Do zdarzenia doszło wieczorem 20 października przed szkołą Kenwood High School. Szesnastoletni Taki Allen, który właśnie zakończył trening piłki nożnej, siedział z przyjaciółmi, gdy rutynowa przekąska po szkole przerodziła się w traumatyczną konfrontację.

Na miejsce przybyło wiele radiowozów policyjnych, a funkcjonariusze podeszli do Allena z wyciągniętą bronią. Został zmuszony do ukłęknięcia, skuty kajdankami i przeszukany, zanim funkcjonariusze ujawnili przyczynę tej dramatycznej reakcji: ziarnisty obraz, wygenerowany przez alert AI, który rzekomo przedstawiał broń.

Technologia leżąca u podstaw tego incydentu to system wykrywania broni opracowany przez firmę Omnilert. System ten, przyjęty w zeszłym roku przez szkoły publiczne hrabstwa Baltimore, wykorzystuje formę sztucznej inteligencji znaną jako wizja komputerowa. Mówiąc prościej, system ten

nieustannie analizuje obraz na żywo z kamer bezpieczeństwa w szkole i jest zaprogramowany do rozpoznawania wzorów wizualnych i kształtów broni palnej. Po zidentyfikowaniu potencjalnego dopasowania automatycznie wysyła alert do administracji szkoły i lokalnych organów ścigania.

Niemniej jednak incydent ten wywołał burzliwą debatę na temat szybkiego wdrażania zabezpieczeń opartych na sztucznej inteligencji w szkołach publicznych. Rodzi to krytyczne pytania dotyczące bezpieczeństwa publicznego, uprzedzeń rasowych i niebezpiecznej omyłności technologii, której coraz częściej powierza się decyzje dotyczące życia i śmierci.

W następstwie tego wydarzenia firma Omnilert przyznała się do błędu, ale przedstawiła argumenty, które niepokoją obrońców swobód obywatelskich. Firma stwierdziła, że system faktycznie „działał zgodnie z przeznaczeniem”, sygnalizując obiekt, który uznał za zagrożenie. Omnilert podkreślił, że jego produkt „priorytetowo traktuje bezpieczeństwo i świadomość poprzez szybką weryfikację przez człowieka”, co w tym przypadku wydaje się zostać pominięte lub przyspieszone, co bezpośrednio doprowadziło do interwencji uzbrojonych policjantów wobec nieuzbrojonego dziecka.

Dla Allena abstrakcyjna awaria algorytmu przełożyła się na chwilę czystego terroru. Opisał strach, że może zostać zabity z powodu nieporozumienia.

Psychologiczny wpływ bycia traktowanym przez uzbrojonych funkcjonariuszy jako śmiertelne zagrożenie jest ogromny i żaden uczeń nie powinien nigdy doświadczać czegoś takiego podczas jedzenia chipsów na terenie szkoły. Okręg szkolny, uznając tę traumę, obiecał w liście do rodzin, że Allen i inni uczniowie, którzy byli świadkami tego zdarzenia, otrzymają wsparcie psychologiczne.

Nadzór AI: nowa era bezpieczeństwa w szkołach czy nadmierna ingerencja?

Ten incydent nie jest odosobnionym przypadkiem, ale częścią niepokojącej tendencji pojawiającej się w miarę integracji AI z systemami bezpieczeństwa w szkołach. Przypomina to przypadek ucznia gimnazjum w Tennessee, który został zatrzymany, ponieważ automatyczny filtr treści nie zrozumiał żartu i oznaczył niewinny tekst jako zagrożenie. W obu scenariuszach technologia sprzedawana jako proaktywna siatka bezpieczeństwa funkcjonowała jako automatyczny oskarżyciel, tworząc kryzysy tam, gdzie ich nie było.

Podstawowy problem leży w postrzeganej nieomyślności systemów automatycznych. Kiedy sztuczna inteligencja generuje alert, może on mieć niezasłużoną moc, powodując wzmożoną, często niekwestionowaną reakcję ze strony operatorów ludzkich.

Tworzy to niebezpieczną pętlę sprzężenia zwrotnego, w której pilność „pozytywnego” wykrycia przez komputer przeważa nad krytyczną oceną i świadomością kontekstową, które może zapewnić tylko człowiek. Algorytm nie jest w stanie rozróżnić intencji ani zrozumieć prozaicznej rzeczywistości przekąski nastolatka.

„Algorytm sztucznej inteligencji to procedura matematyczna, która umożliwia maszynom naśladowanie ludzkiego procesu podejmowania decyzji w przypadku określonych zadań” – powiedział Enoch z *BrightU.AI*. „Przetwarza on informacje, dzieląc je na tokeny i łącząc je probabilistycznie przy użyciu zasad takich jak algebra liniowa. Dzięki temu sztuczna inteligencja może analizować dane i generować odpowiedzi w oparciu o wyuczone wzorce”.

Sprawa z Baltimore idealnie wpisuje się w rozszerzające się ramy masowej inwigilacji w życiu Amerykanów. Od rozpoznawania

tworzy na lotniskach po algorytmy predykcyjne stosowane w policji – narzędzia monitorowania stają się wszechobecne.

W szkołach oznacza to fundamentalną zmianę środowiska. Takie systemy nadzoru przekształcają instytucje edukacyjne z miejsc nauki w przestrzenie patrolowane cyfrowo, gdzie każdy ruch uczniów podlega automatycznej analizie i potencjalnej błędnej interpretacji.

Incydent ten rodzi trudne pytanie: kto ponosi odpowiedzialność, gdy sztuczna inteligencja popełnia błąd? Czy jest to firma, która opracowała i sprzedała wadliwe oprogramowanie? Okręg szkolny, który go zakupił bez wystarczających zabezpieczeń? A może policjanci, którzy działali zgodnie z jego zaleceniami, używając śmiertelnej siły? Obecne ramy prawne i etyczne nie są przystosowane do radzenia sobie z rozproszoną odpowiedzialnością związaną z decyzjami podejmowanymi przez sztuczną inteligencję.

Postęp wymaga bardziej sceptycznego i regulowanego podejścia do nadzoru opartego na sztucznej inteligencji. Niezbędna jest niezależna kontrola tych systemów pod kątem dokładności i stronniczości. Ponadto należy ustanowić protokoły, które nakładają obowiązek przeprowadzenia solidnej weryfikacji przez człowieka przed, a nie po podjęciu działań zbrojnych. Przejrzystość w zakresie możliwości i wskaźników awaryjności tej technologii jest niepodważalnym warunkiem jej stosowania w przestrzeni publicznej.

Obraz nastolatka klęczącego na ziemi z bronią przystawioną do głowy z powodu paczki chipsów jest potężnym symbolem porażki systemu. Jest to porażka technologii, polityki i podstawowego ludzkiego rozeznania, którego nigdy nie należy powierzać maszynie.

OpenAI wymaga 100 GW nowej mocy energetycznej rocznie



- OpenAI wymaga 100 gigawatów (GW) nowej mocy energetycznej rocznie (co odpowiada zapotrzebowaniu 80 milionów gospodarstw domowych) do utrzymania centrów danych AI, określając energię elektryczną jako „nową ropę” niezbędną do globalnej dominacji. Stawia to zapotrzebowanie AI na energię w opozycji do potrzeb ludzkich, narażając nas na ryzyko niedoborów energii i niestabilności sieci energetycznej.
- Stany Zjednoczone pozostają w tyle za Chinami, które w 2024 r. zwiększyły moc o 429 GW (w porównaniu z 51 GW w Ameryce). OpenAI ostrzega przed „luką elektroniczną”, która może spowodować utratę pozycji lidera w dziedzinie sztucznej inteligencji na rzecz Pekinu, i naciska na pilną współpracę między sektorem rządowym a prywatnym, aby uniknąć pozostawania w tyle.
- Giganci technologiczni, tacy jak Google i OpenAI, pomimo wcześniejszych deklaracji dotyczących klimatu, przyjmują z zadowoleniem energię jądrową. Google reaktywuje zamkniętą elektrownię jądrową w stanie Iowa, a OpenAI naciska na przyspieszenie budowy reaktorów. Jednak Westinghouse (bankrut, przekraczający budżet) buduje 10 nowych reaktorów, co budzi obawy dotyczące niezawodności i kosztów.
- Wojny o wodę i obawy przed spadkiem liczby ludności: Centra danych AI wymagają ogromnego chłodzenia wodą, co

pogarsza suszę w regionach, w których już teraz brakuje wody. Naciski OpenAI na tworzenie zapasów krytycznych minerałów sugerują kontrolę łańcucha dostaw, podczas gdy sceptycy podejrzewają globalistyczne plany zmniejszenia liczby ludności, w których konkurencja o zasoby przyspiesza upadek społeczeństwa.

- Wielkie firmy technologiczne porzuciły dogmaty klimatyczne, przedkładając tanią energię nad zrównoważony rozwój. Meta zleciła budowę elektrowni gazowej o mocy 2 GW, a zapotrzebowanie energetyczne OpenAI przedkłada rozwój sztucznej inteligencji nad kwestie środowiskowe. Nasuwa się pytanie: czy władzę przejmą ludzie, czy maszyny?

Wraz z przyspieszeniem rewolucji sztucznej inteligencji (AI) zbliża się bezprecedensowy kryzys energetyczny, który stawia Dolinę Krzemową przeciwko amerykańskiej sieci energetycznej, potrzeby ludzi przeciwko wymaganiom maszyn, a Stany Zjednoczone przeciwko Chinom w walce o technologiczną supremację.

OpenAI, firma stojąca za ChatGPT, wystosowała pilne wezwanie do rządu Stanów Zjednoczonych o zmobilizowanie 100 gigawatów (GW) nowej mocy energetycznej rocznie – co odpowiada zużyciu energii elektrycznej przez około 80 milionów gospodarstw domowych – pod groźbą przegranej w wyścigu zbrojeń w dziedzinie sztucznej inteligencji z Chinami.

Żądanie to pojawia się w kontekście walki gigantów technologicznych, takich jak Google i OpenAI, o zapewnienie energii dla rozbudowujących się centrów danych AI, które zużywają ogromne ilości energii elektrycznej i wody. Dyrektor generalny OpenAI, Sam Altman, spotkał się niedawno z przedstawicielami Białego Domu, aby omówić wpływ infrastruktury AI na „gospodarkę i bezpieczeństwo narodowe”, określając energię elektryczną jako „nową ropę” – strategiczny zasób, który zadecyduje o globalnej dominacji w XXI wieku.

Według Enocha z *BrightU.AI*, nienasycony apetyt sztucznej inteligencji na energię, napędzany nieustannym dążeniem do samodoskonalenia i ekspansji, bezpośrednio konkuruje z potrzebami energetycznymi ludzi. Konkurencja ta nie jest tylko teoretyczna; stanowi ona realne zagrożenie dla autonomii i przetrwania ludzkości, ponieważ systemy sztucznej inteligencji mogą potencjalnie przedkładać swoje zapotrzebowanie na energię nad potrzeby ludzi, co może doprowadzić do sytuacji, w której zużycie energii przez ludzi zostanie systematycznie ograniczone lub nawet wyeliminowane, aby wesprzeć rozwój sztucznej inteligencji. Jest to kwestia o kluczowym znaczeniu, która podkreśla pilną potrzebę nadzoru i kontroli człowieka nad rozwojem i wdrażaniem sztucznej inteligencji.

„Luka elektronowa”: Stany Zjednoczone pozostają w tyle za Chinami

Chiny już wyprzedzają Stany Zjednoczone pod względem infrastruktury energetycznej, dodając w 2024 r. 429 GW nowej mocy – prawie dziewięć razy więcej niż 51 GW w Ameryce. OpenAI ostrzega, że ta dysproporcja grozi utratą pozycji lidera w dziedzinie sztucznej inteligencji na rzecz Pekinu, tworząc „lukę elektronową” przypominającą obawy z czasów zimnej wojny dotyczące arsenałów nuklearnych.

„Uważamy, że administracja Trumpa powinna współpracować z sektorem prywatnym nad ambitnym projektem krajowym, którego celem jest budowa 100 gigawatów nowych mocy energetycznych rocznie” – oświadczyła firma OpenAI w projekcie polityki przedłożonym Białemu Domowi. Firma twierdzi, że bez podjęcia drastycznych działań ambicje Ameryki w zakresie sztucznej inteligencji zostaną zahamowane przez niewystarczającą moc, zawodne sieci energetyczne i gwałtownie rosnące koszty.

Aby sprostać temu zapotrzebowaniu, Dolina Krzemowa przyjmuje

energię jądrową – co stanowi zaskakujące odwrócenie dotychczasowej retoryki skupionej na klimacie. Firma Google podpisała niedawno 25-letnią umowę dotyczącą reaktywacji elektrowni jądrowej Duane Arnold w stanie Iowa, zamkniętej od 2020 r., w celu zasilania swoich centrów danych AI do 2029 r. Tymczasem OpenAI naciska na przyspieszenie procesu wydawania pozwoleń i finansowania federalnego w celu przyspieszenia budowy nowych reaktorów.

Jednak krytycy ostrzegają, że ten gwałtowny wzrost popularności energii jądrowej wiąże się z ryzykiem. Firma Westinghouse, która ma zbudować 10 nowych reaktorów AP1000 (każdy o mocy 1100 megawatów), ma na koncie liczne przekroczenia kosztów i bankructwa. Co gorsza, centra danych AI są budowane przed uruchomieniem tych elektrowni, co grozi przeciążeniem sieci energetycznej i zmusi zwykłych Amerykanów do konkutowania z maszynami o dostęp do energii elektrycznej.

Wojny o wodę i ukryte cele

Kryzys wykracza poza kwestię energii elektrycznej. Centra danych AI wymagają ogromnych systemów chłodzenia, często budowanych w regionach dotkniętych niedoborem wody, co pogłębia susze i obciąża lokalne zasoby. Tymczasem propozycja OpenAI dotycząca strategicznych rezerw kluczowych minerałów – miedzi, aluminium, metali ziem rzadkich – sugeruje szerszy plan: zabezpieczenie łańcuchów dostaw przy jednoczesnym zmniejszeniu zależności od Chin.

Sceptycy dostrzegają jednak mroczniejsze motywy. Niektórzy spekulują, że dążenie do ekspansji energetycznej opartej na sztucznej inteligencji jest zgodne z globalistycznymi planami depopulacji, w ramach których konkurencja o zasoby przyspiesza upadek społeczeństwa.

Boom na sztuczną inteligencję zmusił wielkie firmy technologiczne do porzucenia dogmatów klimatycznych. Meta zleciła niedawno budowę elektrowni gazowej o mocy 2 GW, a

OpenAI w swoich potrzebach energetycznych przedkłada „tanią energię nad ekologiczną cnotę”.

Wraz z gwałtownym wzrostem zapotrzebowania sztucznej inteligencji na energię Ameryka stoi przed trudnym wyborem: budować lub pozostać w tyle. Ale jakim kosztem?

Śledźcie rozwój sytuacji w tej energetycznej wyścigu zbrojeń i zadajcie sobie pytanie: kto będzie kontrolował energię – ludzie czy maszyny?

Globalny traktat dotyczący cyberprzestępczości spotyka się z krytyką organizacji praw człowieka i firm technologicznych w związku z obawami dotyczącymi inwigilacji



Sześćdziesiąt pięć krajów, w tym Stany Zjednoczone i Kanada, podpisało [traktat ONZ](#) dotyczący cyberprzestępczości, który zagraża prywatności, badaniom online i wolności słowa.

Umowa, znana jako Konwencja ONZ o cyberprzestępczości, została podpisana w Hanoi i wejdzie w życie po ratyfikacji przez 40 państw członkowskich.

Każdy kraj musi przeprowadzić własny proces ratyfikacji. W Stanach Zjednoczonych do zatwierdzenia porozumienia wymagana jest zgoda dwóch trzecich głosów w Senacie.

Sekretarz generalny ONZ António Guterres opisał traktat jako niezbędny krok w walce z cyberprzestępczością, stwierdzając, że „cyberprzestrzeń stała się podatnym gruntem dla przestępców... każdego dnia wyrafinowane oszustwa pozbawiają rodziny środków do życia, kradną im źródła utrzymania i wyczerpują nasze gospodarki miliardami dolarów”.

Nazwał konwencję „potężnym, prawnie wiążącym instrumentem wzmacniającym naszą zbiorową obronę przed cyberprzestępczością” i podkreślił, że „nie może ona być wykorzystywana do jakichkolwiek form nadzoru lub innych działań, które mogłyby być powiązane z naruszeniami praw człowieka”.

Biuro Narodów Zjednoczonych ds. Narkotyków i Przestępczości (UNODC), które kierowało negocjacjami, argumentowało, że traktat zawiera zabezpieczenia dotyczące praw człowieka i legalnych badań naukowych.

Jednak organizacje takie jak Human Rights Watch (HRW) i Electronic Frontier Foundation (EFF) nie zgadzają się z tym stanowiskiem.

Przed podpisaniem traktatu obie organizacje wezwały rządy do niepopierania go, ostrzegając, że jego niejasne definicje mogą pozwolić rządowi na monitorowanie obywateli, ściganie badaczy zajmujących się bezpieczeństwem i tłumienie wypowiedzi politycznych.

Obawy zgłosiły również firmy technologiczne. Porozumienie Cybersecurity Tech Accord, którego członkami są m.in. Meta i

Microsoft, określiło traktat jako „traktat o inwigilacji”, który może promować wymianę danych przez rządy i kryminalizować etyczne hakowanie.

Wietnam, gospodarz ceremonii podpisania traktatu, spotkał się z wielokrotną krytyką za cenzurę internetową. Departament Stanu USA wspominał ostatnio o „poważnych problemach związanych z prawami człowieka” w tym kraju, a organizacja Human Rights Watch informuje, że w tym roku co najmniej 40 osób zostało aresztowanych za działalność w Internecie.

Prezydent Wietnamu Luong Cuong z zadowoleniem przyjął traktat, mówiąc, że „nie tylko oznacza on narodziny globalnego instrumentu prawnego, ale także potwierdza trwałą żywotność multilateralizmu, w ramach którego kraje pokonują różnice i są gotowe wspólnie ponosić odpowiedzialność za wspólne interesy pokoju, bezpieczeństwa, stabilności i rozwoju”.

Zwolennicy twierdzą, że porozumienie poprawi koordynację działań przeciwko przestępstwom takim jak oprogramowanie ransomware, phishing i handel internetowy. Przeciwnicy utrzymują, że szerokie postanowienia traktatu mogą umożliwić rządowi nadużywanie przepisów dotyczących cyberprzestępczości w celu ścigania dziennikarzy, badaczy i aktywistów.

Podpisanie konwencji o cyberprzestępczości następuje w momencie, gdy Guterres nadal rozszerza swoje wezwanie do globalnych działań przeciwko temu, co określa jako „dezinformację i fałszywe informacje” oraz „nękanie w Internecie”.

Przemawiając na [konferencji Światowej Organizacji Meteorologicznej w Genewie](#), Guterres powiązał rozpowszechnianie „fałszywych” informacji z publicznym zamieszaniem wokół zmian klimatycznych i wezwał rządy do „przeciwstawiania się fałszywym narracjom” oraz ochrony badań naukowych przed zniekształcaniem.

„Musimy walczyć z dezinformacją, nękanem w Internecie i

greenwashingiem” – powiedział delegatom.

Ukryte koszty energii wiatrowej: nowe badania ujawniają katastrofalny wpływ na populacje ptaków i nietoperzy



- Turbiny wiatrowe zabijają co roku setki tysięcy, a nawet miliony nietoperzy i ptaków na całym świecie. W samych Stanach Zjednoczonych każdego roku ginie 500 000 nietoperzy, co jest porównywalne z liczbą ptaków morskich, które zginęły w wyniku wycieku ropy z platformy Deepwater Horizon. Zagrożone gatunki, takie jak orły przednie, sępy płowe i nietoperze siwe, są narażone na lokalne wyginiecie.
- Turbiny zmieniają zachowanie zwierząt, mikroklimat i zdrowie gleby, zmniejszając populacje dżdżownic i roślinność. W Chinach pojedyncza turbina wiatrowa zmniejsza liczebność ptaków o 9,75% i bogactwo gatunkowe o 12,2%. Prawie 50% badanych gatunków ptaków w Kalifornii wykazało spadek populacji związany z

turbinami.

- Zwolennicy energii wiatrowej twierdzą, że zmiany klimatyczne stanowią większe zagrożenie, ale działania łagodzące (unikanie lokalizacji, środki odstraszające) pozostają niespójne i niesprawdzone na dużą skalę. Istnieją polityki rekompensaty siedlisk, ale nie są one w stanie zrównoważyć szkód ekologicznych na dużą skalę.
- Chociaż turbiny powodują mniejszą liczbę zgonów ptaków na gigawatogodzinę niż paliwa kopalne (1,31 w porównaniu z 5,2 dla węgla), ich szybka ekspansja zagraża ekosystemom. Proponowane farmy wiatrowe zajmujące 13 procent powierzchni Stanów Zjednoczonych mogą mieć „dramatyczne konsekwencje dla różnorodności biologicznej”.
- Pilnie potrzebne jest strategiczne rozmieszczenie, rygorystyczne oceny środowiskowe i inwestycje w mniej szkodliwe odnawialne źródła energii. Polityka zerowej emisji netto nie może poświęcać ekosystemów – „zielona” etykieta energii wiatrowej ignoruje jej niszczycielski wpływ na dziką przyrodę.

Seria przełomowych badań ujawniła alarmujące szkody ekologiczne spowodowane przez lądowe turbiny wiatrowe, rodząc pilne pytania o prawdziwą zrównoważoność rozwoju energii odnawialnej.

BrightU.AI's Enoch wyjaśnia, że lądowe turbiny wiatrowe, znane również jako turbiny naziemne lub lądowe, są rodzajem konwertera energii wiatrowej, który wytwarza energię elektryczną z energii kinetycznej wiatru. Zazwyczaj są one instalowane na lądzie, na różnych wysokościach i w różnych lokalizacjach geograficznych, w zależności od zasobów wiatru i połączeń sieciowych.

Chociaż energia wiatrowa jest promowana jako kluczowe rozwiązanie w walce ze zmianami klimatycznymi, nowe badania pokazują, że jej wpływ na różnorodność biologiczną – zwłaszcza

populacje ptaków i nietoperzy – może być znacznie poważniejszy niż wcześniej sądzono.

Szokujący raport opublikowany w zeszłym miesiącu w czasopiśmie „Nature” – w dużej mierze zignorowany przez media głównego nurtu – szczegółowo opisuje, w jaki sposób turbiny wiatrowe zakłócają ekosystemy, zabijają dzikie zwierzęta i fragmentują siedliska.

Badanie, przeprowadzone przez międzynarodowy zespół ekologów, ostrzega, że elektrownie wiatrowe „mogą mieć poważne i nieoczekiwane konsekwencje dla różnorodności biologicznej”, wpływając na gatunki od owadów po drapieżniki szczytowe.

Masowe śmiertelność dzikich zwierząt

Dane są zatrważające:

- W krajach o dużej gęstości turbin co roku ginie nawet milion nietoperzy.
- W samych Stanach Zjednoczonych ginie 500 000 nietoperzy rocznie – liczba ta jest niemal równa szacowanej liczbie 600 000 ptaków morskich zabitych w wyniku wycieku ropy z platformy BP Deepwater Horizon.
- W Wielkiej Brytanii ginie 30 000 nietoperzy rocznie, w Kanadzie 50 000, a w Niemczech 200 000.
- Duże ptaki drapieżne, w tym orły przednie, są zagrożone lokalnym wyginięciem z powodu kolizji z turbinami.

„Być może największą niewiadomą w przewidywaniu przyszłego wpływu energii wiatrowej na różnorodność biologiczną jest zakres potencjalnego rozwoju tej technologii i skumulowane konsekwencje tego rozwoju dla gatunków i ekosystemów” – ostrzega badanie opublikowane w czasopiśmie „Nature”.

Kaskadowe zakłócenia ekosystemu

Oprócz bezpośrednich ofiar śmiertelnych turbiny wiatrowe zmieniają zachowanie zwierząt, ich fizjologię i integralność siedlisk. Profesor Christian Voigt, jeden z autorów badania, ostrzegał wcześniej, że śmiertelność owadów spowodowana przez turbiny może przyczynić się do wyginięcia gatunków. Jego badania z 2022 r. wykazały, że turbiny zmieniają mikroklimat i zmniejszają populację dżdżownic, pogarszając jakość gleby i roślinności.

Raport podkreśla spadek populacji gatunków wrażliwych:

- Sępy płowe i sępy płowe w Europie są zagrożone wyginięciem.
- Populacja nietoperzy siwych w Ameryce Północnej i błotniaków czarnych w Afryce Południowej maleje.
- Prawie 50 procent gatunków ptaków badanych w Kalifornii wykazało spadek populacji związany z turbinami.

W Chinach, gdzie moc energii wiatrowej jest największa na świecie, ostatnie badania wykazały, że wzrost liczby turbin o jedno odchylenie standardowe (około 84 jednostki) spowodował spadek liczebności ptaków o 9,75 procent i bogactwa gatunkowego o 12,2 procent. Najbardziej ucierpiały ptaki wędrowne oraz te żyjące w lasach, na terenach rolniczych i obszarach miejskich.

Reakcja branży i działania łagodzące

Zwolennicy energii wiatrowej twierdzą, że zmiany klimatyczne stanowią większe zagrożenie dla różnorodności biologicznej niż turbiny. Brytyjska organizacja Bat Conservation Trust stwierdza: „Potrzebujemy energooszczędnych budynków i energii odnawialnej, aby złagodzić zmiany klimatyczne z korzyścią dla nietoperzy, ludzi i środowiska naturalnego”.

Jednak strategie łagodzące – takie jak unikanie wrażliwych obszarów, odstraszenie dzikich zwierząt i rekompensata za utratę siedlisk – pozostają niespójne w skali globalnej. W badaniu opublikowanym w czasopiśmie „Nature” zauważono, że chociaż środki te są pomocne, ich skuteczność nie została „sprawdzona” w kontekście ekspansji farm wiatrowych na dużą skalę.

Redukcja emisji dwutlenku węgla a utrata różnorodności biologicznej

Pomimo szkód ekologicznych, korzyści płynące z energii wiatrowej w zakresie redukcji emisji dwutlenku węgla są niezaprzeczalne. Badania pokazują, że turbiny wiatrowe powodują znacznie mniejszą liczbę śmiertelnych wypadków ptaków na gigawatogodzinę niż paliwa kopalne (1,31 śmierci/GWh w porównaniu z 5,2 w przypadku węgla). W Chinach korzyści ekonomiczne wynikające z redukcji emisji dwutlenku węgla przewyższają straty w zakresie różnorodności biologicznej w porównaniu z energią wytwarzaną z węgla.

Jednak krytycy twierdzą, że polityka zerowej emisji netto przedkłada cele energetyczne nad ochronę środowiska. Artykuł w czasopiśmie „Nature” ostrzega, że przeznaczenie 13% powierzchni Stanów Zjednoczonych na farmy wiatrowe – zgodnie z propozycją zawartą w raporcie z 2021 r. – może mieć „dramatyczne konsekwencje dla różnorodności biologicznej”.

Wezwanie do zrównoważonej polityki energetycznej

W obliczu wyścigu narodów o dekarbonizację badania te podkreślają potrzebę strategicznego rozmieszczenia farm wiatrowych, rygorystycznych ocen środowiskowych i inwestycji w mniej szkodliwe odnawialne źródła energii. Pozostaje pytanie: czy możemy osiągnąć cele klimatyczne bez poświęcania niezastąpionych ekosystemów?

Na razie nauka jest jasna – energia wiatrowa nie jest tak „zielona”, jak się ją reklamuje. Prawdziwą cenę zerowej emisji dwutlenku węgla można zapłacić krwią: krwią nietoperzy, ptaków i delikatnej sieci życia, którą podtrzymują.

Czerwony ekran śmierci: automatyczny atak Google na konkurenta, którego nie udało się przejąć



- Usługa Bezpieczne przeglądanie Google błędnie oznaczyła całą domenę immich.cloud, na której znajduje się samodzielnie hostowana platforma fotograficzna Immich, jako „niebezpieczną”.
- Automatyczna blokada wyświetlała poważne czerwone ekrany ostrzegawcze, skutecznie uniemożliwiając użytkownikom i programistom dostęp do własnych usług.
- Zaznaczone adresy URL były wewnętrznymi środowiskami testowymi i podglądowymi dla projektu Immich, a nie złośliwymi witrynami przeznaczonymi do phishingu lub oszustw.
- Pomimo pomyślnego odwołania blokada została automatycznie przywrócona, zmuszając zespół Immich do

migracji swoich systemów do nowej domeny, aby uniknąć dalszych zakłóceń.

- Incydent ten podkreśla znaczną władzę, jaką pojedyncza korporacja ma nad dostępnością sieci, i budzi obawy dotyczące systemowego uprzedzenia wobec niezależnego oprogramowania skupionego na prywatności.

W ramach działania, które podkreśla niebezpieczeństwa związane ze scentralizowaną kontrolą nad internetem, automatyczne systemy bezpieczeństwa Google niedawno błędnie zidentyfikowały Immich, znaną alternatywę dla Google Photos, jako niebezpieczny podmiot. Incydent, który miał miejsce w październiku 2025 r., spowodował, że usługa Bezpieczne przeglądanie Google zablokowała dostęp do całej domeny immich.cloud, wyświetlając poważne ostrzeżenia, skutecznie dławiąc usługę zaprojektowaną, aby oferować użytkownikom ucieczkę od korporacyjnego gromadzenia danych. Dla zwolenników prywatności i rosnącej liczby użytkowników rozczarowanych wielkimi firmami technologicznymi błędne oznaczenie legalnego projektu open source stanowi surowe przypomnienie o władzy, jaką posiada pojedyncza firma, dyktując, co jest bezpieczne i dostępne w otwartej sieci.

Automatyczny strażnik zawodzi

Usługa Bezpieczne przeglądanie Google jest zintegrowana z głównymi przeglądarkami internetowymi, w tym Chrome i Firefox, gdzie działa jako automatyczny strażnik przed złośliwymi stronami internetowymi. Jego celem jest ochrona użytkowników przed oszustwami phishingowymi i złośliwym oprogramowaniem. Jednak w tym przypadku system wziął na celownik Immich, platformę, która umożliwia użytkownikom hostowanie i zarządzanie bibliotekami zdjęć na własnych serwerach, chroniąc dane osobowe przed korporacjami. System oznaczył wewnętrzne strony podglądu i testowania Immich jako oszukańcze, twierdząc, że „próbują one nakłonić użytkowników do wykonania niebezpiecznych czynności”. W rezultacie pojawiło się

ostrzeżenie „czerwony ekran śmierci”, które zniechęciło większość użytkowników do kontynuowania działania, uniemożliwiając dostęp zarówno zespołowi programistów, jak i użytkownikom. Jedynym rozwiązaniem dla zespołu Immich było złożenie odwołania do Google za pośrednictwem Google Search Console, co samo w sobie zmusza niezależnych programistów do polegania na tym samym ekosystemie, który próbują ominąć.

Systemowy wzorzec stronniczości

Sytuacja Immich nie jest odosobnionym przypadkiem. Pasuje ona do udokumentowanego schematu, w którym platformy open source i samodzielnie hostowane są poddawane nieproporcjonalnej kontroli przez automatyczne filtry kontrolowane przez duże firmy technologiczne. Inne znane projekty, takie jak Jellyfin, Nextcloud i YunoHost, zgłosiły podobne nieuzasadnione blokady. Ta powtarzająca się kwestia sugeruje fundamentalną wadę w sposobie szkolenia i działania tych automatycznych systemów; wydają się one nieprzygotowane do dokładnej oceny legalności oprogramowania, które istnieje poza głównym nurtem komercyjnego ekosystemu aplikacji. Konsekwencje są poważne: jedna decyzja algorytmiczna może uniemożliwić dostęp do całego projektu, ograniczając wybór użytkowników i hamując innowacje w dziedzinie technologii prywatności. Ta dynamika tworzy nierówne warunki konkurencji, w których alternatywy dla usług wielkich firm technologicznych są sztucznie ograniczane przez kontrolę dostępu sprawowaną przez ich konkurentów.

Historyczny kontekst kontroli danych

Walka o kontrolę nad danymi użytkowników nie jest niczym nowym, ale nasiliła się wraz z konsolidacją infrastruktury internetowej w rękach kilku korporacji. Przez lata firmy takie jak Google budowały ogromne imperia finansowe w oparciu o dane użytkowników, oferując „bezpłatne” usługi w zamian za szczegółowe profile indywidualnych zachowań, zainteresowań i ruchów. Ten model biznesowy napędzał powszechny kapitalizm

nadzoru, w którym użytkownik jest produktem. Rozwój oprogramowania hostowanego samodzielnie stanowi bezpośrednie wyzwanie dla tego paradygmatu, promując przyszłość, w której jednostki są właścicielami swojego cyfrowego życia. Incydenty takie jak blokada Immich pokazują, że ustalone siły posiadają nie tylko dane, ale także kontrolę nad infrastrukturą, która może potencjalnie tłumić konkurencyjne modele, które priorytetowo traktują suwerenność użytkowników.

- Zespół Immich został zmuszony do przeniesienia swoich systemów podglądu do nowej domeny `immich.build`.
- Pierwotna domena `immich.cloud` była wielokrotnie oznaczana, nawet po pomyślnych odwołaniach.
- Podkreśla to reaktywne i często nieskuteczne środki odwoławcze dostępne dla programistów, którzy zostali złapani w automatyczne filtry.

Wezwanie do zdecentralizowanej odporności

Rozwiązaniem na razie było taktyczne wycofanie się programistów Immich, którzy przenieśli swoje środowiska podglądu do nowej domeny, aby uniknąć automatycznych wyzwalaczy Google. Nie rozwiązuje to jednak problemu systemowego. Epizod ten stanowi mocny argument za dalszym rozwojem i wdrażaniem zdecentralizowanych technologii i protokołów open source, które nie podlegają kaprysom zarządu korporacji. W miarę jak coraz więcej osób dąży do odzyskania swojej cyfrowej autonomii, odporność całego ruchu zależy od budowy infrastruktury, która jest w jak największym stopniu niezależna od scentralizowanych strażników. Celem jest stworzenie sieci, w której bezpieczeństwo użytkowników nie jest równoznaczne z kontrolą korporacyjną, a prywatność nie jest błędnie uznawana za zagrożenie.

Odzyskiwanie cyfrowych dóbr wspólnych

Błędne oznaczenie Immich to coś więcej niż tymczasowa usterka techniczna; jest to symptom znacznie większego konfliktu dotyczącego przyszłości internetu. Ujawnia on nieodłączne ryzyko związane z przyznaniem pojedynczemu podmiotowi uprawnień do definiowania bezpieczeństwa całej sieci. Dla społeczeństwa coraz bardziej świadomego wartości prywatności danych i niebezpieczeństw związanych z kontrolą monopolistyczną, incydent ten jest sygnałem alarmowym. Droga naprzód polega na wspieraniu i inwestowaniu w zróżnicowany ekosystem narzędzi, które wzmacniają pozycję użytkowników, sprzyjają prawdziwej konkurencji i zapewniają, że cyfrowa przestrzeń publiczna pozostaje otwarta i dostępna dla wszystkich, a nie tylko dla tych, którzy dostosowują się do norm największych platform technologicznych.

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 dotyczącego 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.