

Sztuczna inteligencja jako eksternalizacja intelektu: analiza dziewięciu tez Dennisa Yi Tenena



AUTOR

Fundacja Dobre Państwo

STRESZCZENIE / ABSTRACT

Artykuł analizuje koncepcje zawarte w książce Dennisa Yi Tenena, skupiając się na AI jako formie eksternalizacji intelektu. Autor tekstu podkreśla, że sztuczna inteligencja nie jest autonomicznym bytem, lecz rezultatem „zbiorowej pracy” (collectiva tabor) – skumulowanego wysiłku inżynierów, badaczy i niewidzialnych pracowników danych. W tekście przywołane są teorie rozproszonego poznania oraz krytyka antropomorfizacji AI. Analiza wskazuje na socjotechniczny charakter systemów AI, które, choć posiadają swoiste mechanizmy obliczeniowe, są głęboko osadzone w historycznych i społecznych strukturach wytwarzania wiedzy. Całość stanowi próbę odczarowania mitu samodzielnej inteligencji maszynowej na rzecz modelu współpracy człowieka z technologią.

The article analyzes concepts contained in the book by Dennis Yi Tenen, focusing on AI as a form of intellectual externalization. The author emphasizes that artificial intelligence is not an autonomous entity, but the result of "collective labor" (collectiva tabor)—the accumulated effort of engineers, researchers, and invisible data workers. The text references theories of distributed cognition and critiques the anthropomorphization of AI. The analysis points to the socio-technical nature of AI systems, which, while possessing specific computational mechanisms, are deeply embedded in historical and social structures of knowledge production. Overall, it is an attempt to debunk the myth of independent machine intelligence in favor of a model of human-technology collaboration.

Artykuł stanowi krytyczną analizę natury sztucznej inteligencji, odrzucając popularne narracje o autonomicznych bytach czy nadchodzącej superinteligencji na rzecz ujęcia AI jako nowej klasy artefaktów poznawczych. Autor stawia tezę, że współczesne systemy AI

nie są niezależnymi podmiotami, lecz rezultatem historycznego procesu eksternalizacji ludzkiego intelektu oraz skumulowanej pracy zbiorowej – od inżynierów po niewidzialnych pracowników danych. Tekst argumentuje, że antropomorficzne metafory (takie jak „uczenie się” czy „decydowanie”) przesłaniają rzeczywiste mechanizmy techniczne i prowadzą do niebezpiecznego rozproszenia odpowiedzialności moralnej i prawnej. Analizując ryzyko tzw. „inteligencji generycznej”, artykuł ostrzega przed systemową homogenizacją kultury i wiedzy, gdzie AI zamiast wspierać kreatywność, może sprowadzać ludzką ekspresję do uśrednionego standardu. W konsekwencji tekst wzywa do przejścia od pytań o „inteligencję” maszyn ku analizie systemowych relacji między infrastrukturą technologiczną a ludzkim sprawstwem, sugerując, że AI jest w istocie lustrem odbijającym nasze zbiorowe zasoby i instytucjonalne uwarunkowania.

SŁOWA KLUCZOWE

eksternalizacja intelektu

rozproszone poznanie

collective labor

distributed cognition

antropomorfizm AI

uczenie reprezentacji

moral agency

algorithmic management

artefakty poznawcze

teoria zadaniowa automatyzacji

human oversight

epistemiczna różnorodność

systemy socjotechniczne

AI Act

modelowanie probabilistyczne

AI jako rezultat zbiorowej pracy i rozproszonego poznania

Nine Big Ideas – dziewięć tez będących efektem syntezy historii, nauki i filozofii. Ostatni rozdział Literary Theory for Robots Dennis Yi Tenen zatytułował przewrotnie 9 Big Ideas for an Effective Conclusion. Ironia tego tytułu jest kluczowa dla sposobu, w jaki należy odczytywać finał książki. Autor świadomie parodiuje wydawniczy przymus redukcji złożonych argumentów do kilku łatwych do zapamiętania "mądrości", pozostawiając nawet metakomentarz imitujący instrukcję dopisania efektownej sentencji o dobru społecznym, pracy, środkach produkcji i przyszłości kreatywności. Dziewięciu idei nie należy zatem traktować jak aksjomatów nowej teorii sztucznej inteligencji. Bardziej użyteczne jest uznanie ich za hipotezy organizujące program badawczy, który – po przejściu przez historię kombinatoryki, mechanizacji języka, probabilistyki, kognitywistyki,

ekonomii politycznej i filozofii odpowiedzialności – może zostać poddany krytycznej rekonstrukcji. Wymaga to jednak drobnej korekty redakcyjnej materiału źródłowego.

W końcowym skrócie dziewięciu tez pojawia się sformułowanie, że AI działa społecznie bardziej jak korporacja lub państwo niż samotny robot. Jednak wcześniejsza, szczegółowa część notatki identyfikuje piątą ideę Tenena jako *Metaphor Doesn't Hurt* – "metafory nie odczuwają bólu" – a dopiero szóstą jako tezę o braku moralnego sprawstwa maszyn.

Taki sam układ potwierdza niezależna recenzja książki: praca zbiorowa, rozproszona inteligencja, metaforyczność machine learning, przesłanianie odpowiedzialności przez metafory, brak cierpienia metafory, brak pełnego moralnego sprawstwa maszyny, automatyzacja pracy wiedzy, polityczność technologii oraz przejście od "general" do "generic intelligence". Metafora państwa i korporacji pozostaje istotnym elementem argumentacji Tenena, lecz trafniej rozumieć ją jako rozwinięcie tez o kolektywnym i instytucjonalnym charakterze inteligencji niż jako osobną, numerowaną zasadę.

Pierwsza teza głosi, że AI jest pracą zbiorową – *collective labor*. Tenen przeciwstawia tę koncepcję obrazowi autonomicznego elektronicznego ducha: zdolność widoczna w interfejsie jest rezultatem skumulowanego wysiłku inżynierów, leksykografów, badaczy korpusowych, autorów oraz użytkowników, których działalność pozostawiła dane możliwe do wykorzystania przez kolejne systemy. Oficjalne omówienie książki przez Columbia University dobrze oddaje tę intencję: nawet narzędzie tak niepozorne jak korektor pisowni jest w ujęciu Tenena rezultatem rozciągniętej w czasie współpracy wielu ludzi, a automatyzacja wcześniejszych "papierowych maszyn" wiedzy nie czyni ich z tego powodu autonomicznymi istotami. W świetle współczesnych badań teza ta brzmi szczególnie przekonująco, o ile potraktujemy ją jako twierdzenie socjotechniczne, a nie ontologiczne.

Międzynarodowa Organizacja Pracy dokumentuje ogromną rolę ukrytych pracowników oznaczających, czyszczących, klasyfikujących i weryfikujących dane, moderujących treści oraz wykonujących działania umożliwiające trenowanie i utrzymywanie systemów AI. ILO zwraca uwagę, że za interfejsem automatyzacji kryją się zarówno wysoko kwalifikowani twórcy algorytmów, jak i często niewidzialni *data workers*, których praca jest organizowana przez globalne platformy *crowdworkingowe* i przedsiębiorstwa *outsourcingowe*. "Sztuczność" inteligencji nie oznacza więc braku człowieka w łańcuchu produkcyjnym; oznacza raczej przeniesienie człowieka poza pole widzenia końcowego użytkownika.

Teza Tenena zyskuje na głębi po zestawieniu z Marksem, Babbage'em, teorią kapitału ludzkiego i współczesną ekonomią dóbr niematerialnych. Model jest technicznym artefaktem o konkretnych właściwościach wynikających z architektury i procesu treningu, lecz warunkiem jego istnienia jest

uprzedmiotowanie ogromnej ilości społecznej praktyki w postaci zapisów. Nie wolno przy tym dokonywać przeciwnej redukcji i twierdzić, że model "jest tylko ludzką pracą". Sieć neuronowa uczy się reprezentacji i transformacji, których nie da się utożsamić z prostą biblioteką odpowiedzi. Trafniejsza synteza brzmi: współczesna AI posiada technicznie swoiste mechanizmy obliczeniowe, ale są one osadzone w historycznie zbiorowym systemie wytwarzania języka, danych, norm, infrastruktury i ekspertyzy. Tenen słusznie odczarowuje mit autonomicznej genezy inteligencji maszynowej, lecz nie powinno się tego robić tak daleko, aby zniknęła rzeczywista nowość uczenia reprezentacji przez współczesne modele.

Druga teza mówi, że inteligencja jest rozproszona. Poznanie przedstawione jest tu jako relacja między użytkownikiem, urządzeniem, chmurą obliczeniową, narzędziami i instytucjami selekcionującymi informację. Jest to twierdzenie posiadające znacznie szersze zaplecze niż sama książka Tenena.

Teorie distributed cognition Edwina Hutchinsa oraz extended mind Andy'ego Clarka i Davida Chalmersa od dawno kwestionują przekonanie, że właściwą jednostką analizy procesów poznawczych musi być pojedynczy mózg. Stanford Encyclopedia of Philosophy wskazuje, że tradycja poznania rozproszonego bada procesy zachodzące w złożonych systemach socjotechnicznych, składających się z ludzi i artefaktów, podczas gdy koncepcje umysłu rozszerzonego koncentrują się na jednostce poznawczej funkcjonalnie sprzęgniętej z zasobami zewnętrznymi.

AI jako eksternalizacja intelektu a pułapka antropomorfizmu

Historyczne przykłady analizowane w tym artykule potwierdzają, że rozproszenie nie zaczęło się wraz z elektroniką. Abakus, alfabet, zapis matematyczny, księga rachunkowa, biblioteka, indeks, mapa, formularz, lista kontrolna i tabela statystyczna pozwalają wykonywać operacje niemożliwe lub zbyt trudne dla nieuzbrojonej pamięci biologicznej. W tym znaczeniu pytanie o to, czy „człowiek sam” policzył, zapamiętał lub zaprojektował, jest często źle postawione. Człowiek myśli poprzez system kulturowych artefaktów; AI powiększa jedynie skalę, szybkość i interaktywność tego zjawiska.

Należy jednak zachować rozróżnienie między tezą o rozproszonym procesie poznawczym a tezą o rozproszonej świadomości. Z faktu, że grupa ludzi, dokumenty i urządzenia wspólnie rozwiązują problem, nie wynika istnienie fenomenalnego podmiotu obejmującego całość systemu. Rodzina

może „pamiętać” w sensie społecznym dzięki historiom, fotografiom i podziałowi ról pamięciowych; państwo może „wiedzieć” poprzez rejestry i urzędy; korporacja może „decydować” poprzez procedury governance. Są to rzeczywiste własności organizacyjne, lecz nie muszą one oznaczać, że rodzina, państwo albo korporacja czują ból i posiadają jedno pole świadomości. Tenenowska formuła, według której AI „pamięta jak rodzina, myśli jak państwo i rozumie jak korporacja”, powinna więc pozostać analogią socjologiczną, a nie hipotezą neurofenomenologiczną. Trzecia teza głosi natomiast, że „machine learning” jest metaforą, a metafora nie jest mechanizmem. Jest to być może najważniejsze językowe ostrzeżenie całej książki.

Używamy tych samych czasowników wobec procesów biologicznych i matematycznych: człowiek „uczy się” mówić, sieć „uczy się” reprezentacji; dziecko „widzi” psa, system wizyjny „widzi” obiekt; człowiek „pamięta” spotkanie, model „pamięta” informację treningową. Podobieństwo funkcjonalnego rezultatu zachęca do projekcji podobieństwa mechanizmu, choć droga rozwojowa, architektura i związek ze środowiskiem są radykalnie odmienne. Kirkus, streszczając argument Tenena, trafnie zaznacza, że autor przeciwstawia powtarzalny i algorytmiczny proces uczenia maszynowego konceptualnemu, biologicznemu i rozwojowemu charakterowi uczenia ludzkiego. Teza ta jest słuszna jako ostrzeżenie przed antropomorfizmem, lecz wymaga korekty, jeśli miałyby sugerować, że uczenie maszynowe jest jedynie udawaniem uczenia. W naukach o uczeniu „learning” może być terminem technicznym oznaczającym trwałą zmianę parametrów systemu pod wpływem danych i sygnału optymalizacyjnego. Nie ma sprzeczności między stwierdzeniem, że maszyna autentycznie uczy się w sensie technicznym, a uznaniem, że nie uczy się tak jak dziecko. Właśnie to rozróżnienie powinno zastąpić spór o to, komu „wolno” używać czasownika „uczyć się”. Najnowsze badania reprezentacji pojęciowych pokazują, dlaczego przeciwstawienie człowiek-model nie może być już proste. Badanie opublikowane w *Nature Human Behaviour* w 2025 roku porównało reprezentacje około 4442 pojęć u ludzi oraz kilku rodzin dużych modeli językowych.

Modele tekstowe odtwarzały stosunkowo dobrze niesensomotoryczne aspekty ludzkich reprezentacji pojęciowych, natomiast zgodność spadała w obszarze sensorycznym i była szczególnie niska dla wymiarów motorycznych; dołączenie uczenia wizualnego poprawiało zgodność w kategoriach związanych z widzeniem. Wynik ten wspiera intuicję Tenena, że metafora uczenia nie powinna zasłaniać różnicy mechanizmów i doświadczeń, lecz jednocześnie osłabia skrajną wersję tezy o „więzieniu języka”: sam język okazuje się nośnikiem znacznej ilości struktury konceptualnej, nawet jeśli nie odtwarza pełnego ucieleśnionego doświadczenia człowieka. Czwarta teza mówi, że metafory przesłaniają odpowiedzialność. W źródle Tenen używa mocnego porównania: obwinianie „AI” za dezinformację może być tak nieadekwatne, jak obwinianie miny zamiast organizacji, która ją umieściła. Problem ten nie jest jedynie stylistyczny.

Gramatyczne zdanie „algorytm odrzucił wniosek” kondensuje długi łańcuch decyzji: ktoś określił cel systemu, ktoś wybrał dane, ktoś przyjął model, ktoś ustalił próg, ktoś zdecydował o zakresie automatyzacji i ktoś nadał wynikowi określone konsekwencje prawne albo ekonomiczne. Personifikacja urządzenia może więc przekształcić rozproszoną odpowiedzialność w pozorną odpowiedzialność nieistniejącej osoby. Współczesne prawo regulacyjne w dużej mierze potwierdza tę intuicję, ponieważ nie przypisuje odpowiedzialności mistycznemu „AI”, lecz rozdziela obowiązki między dostawców i podmioty wdrażające. AI Act wymaga między innymi odpowiedniego nadzoru człowieka nad systemami wysokiego ryzyka, dokumentacji, informacji o ograniczeniach systemu oraz, w określonych zastosowaniach, oceny wpływu na prawa podstawowe. Wprost podkreśla też, że ryzyko może wynikać zarówno ze sposobu zaprojektowania systemu, jak i z kontekstu jego użycia.

Konstrukcja ta jest filozoficznie istotna: technologia może mieć sprawczość przyczynową, nie będąc ostatecznym podmiotem odpowiedzialności normatywnej. Nie oznacza to jednak, że odpowiedzialność zawsze można bezproblemowo przypisać jednej osobie. W systemach złożonych istnieją problemy many hands, nieprzewidywalności, asymetrii wiedzy i moralnych "stref zgniotu", w których pracownik najbliższy momentowi awarii może ponieść winę za decyzje systemowe pozostające poza jego kontrolą. Dlatego prawidłowa korekta antropomorfizmu nie powinna brzmieć: „algorytm nie odpowiada, więc odpowiada programista”. Należy zamiast tego rekonstruować poziomy odpowiedzialności projektowej, zarządczej, wdrożeniowej, profesjonalnej i regulacyjnej.

Brak biologicznego cierpienia nie wyczerpuje definicji podmiotowości

Tenen trafnie identyfikuje problem metaforycznej ucieczki od odpowiedzialności, jednak jego militarna analogia okazuje się zbyt uproszczona, by objąć całą topologię współczesnych systemów organizacyjnych. Piąta teza, głosząca, że "metafory nie odczuwają bólu", przenosi argumentację na grunt antropologii filozoficznej. Autor wskazuje w niej brak bólu, lęku przed śmiercią, miłości czy biologicznych imperatywów jako zasadniczą różnicę między maszyną a żyjącym uczestnikiem wspólnoty moralnej.

Intuicję tę można powiązać z Hume'owską teorią uczuć moralnych, etyką troski, fenomenologią ucieleśnienia, levinasowską podatnością na wezwanie Innego, jonasowską kategorią odpowiedzialności czy współczesnymi nurtami 4E cognition. W tych tradycjach moralność nie jest jedynie systemem logicznego przetwarzania nakazów; wyrasta ona z kondycji istoty, której coś

może się stać – która cierpi, traci, potrzebuje innych i może zostać przez nich zraniona. Nie należy jednak czynić z biologicznego argumentu Tenena uniwersalnej definicji etyki. Kant wyprowadza normatywność przede wszystkim z autonomii rozumu praktycznego, a nie z bólu, natomiast teorie kontraktualistyczne pytają o zdolność uczestnictwa w systemie wzajemnie uzasadnialnych norm. Z kolei niektóre funkcjonalistyczne koncepcje moral agency dopuszczają techniczne użycie pojęcia "sztucznego agenta moralnego" bez wymagania wszystkich cech osoby ludzkiej. Stanford Encyclopedia of Philosophy rozróżnia tu systemy mające wpływ etyczny, systemy z wbudowanymi ograniczeniami normatywnymi, explicit ethical agents oraz znacznie bardziej wymagające full ethical agents. Tenen nie rozstrzyga zatem definicyjnie całej filozofii moralności.

Jego argument najsilniej wybrzmiewa jako przypomnienie, że biegłość systemu w posługiwaniu się językiem emocji nie jest dowodem jego fenomenalnego uczestnictwa w bólu, trosce czy śmiertelności. Szósta teza idzie jeszcze dalej: dzisiejsze maszyny nie powinny być traktowane jako pełne podmioty moralne. W materiale źródłowym zostaje ona powiązana z kuhlmannowskim wymogiem "inward understanding" oraz obrazem mechanizmu zdolnego do poprawnego wykonania formy, lecz pozbawionego wewnętrznej transformacji.

Zestawienie to z historią etyki pokazuje, że pojęcie moralnego agenta składa się z wielu potencjalnych warunków: Arystoteles wymaga praktycznej mądrości i charakteru, Augustyn woli, Akwinata racjonalnego samookreślenia, Kant autonomii, Hegel wzajemnego uznania, Strawson uczestnictwa w sieci reaktywnych postaw, a współczesne teorie odpowiedzialności – zdolności reagowania na racje oraz odpowiedniej kontroli nad własnym działaniem.

Współczesna nauka nie dostarcza podstaw, by stwierdzić, że modele LLM spełniają ten bogaty zestaw warunków. Jednocześnie problem świadomości AI nie może zostać uznany za definitywnie rozstrzygnięty na mocy samej intuicji. W 2025 roku interdyscyplinarny zespół badaczy zaproponował metodologię oceny systemów w oparciu o wskaźniki wyprowadzone z konkurencyjnych teorii naukowych, podkreślając przy tym znaczny poziom niepewności właściwy obecnemu stanowi wiedzy. Najbardziej odpowiedzialna pozycja naukowa jest zatem asymetryczna: nie mamy wystarczających podstaw do przypisywania obecnym modelom pełnej świadomości i odpowiedzialności moralnej, lecz nie powinniśmy zamieniać dzisiejszej niewiedzy w metafizyczny dowód, że żaden system sztuczny nigdy nie mógłby posiadać odpowiednich własności.

Ta subtelność ma kluczowe znaczenie prawne. Osobowość prawna, moral agency, świadomość fenomenalna i sprawczość funkcjonalna to odrębne kategorie. Korporacja jest podmiotem prawa bez biologicznego systemu nerwowego; algorytm może przyczynowo oddziaływać na świat, nie będąc osobą prawną; zwierzę może posiadać status pacjenta moralnego bez odpowiedzialności analogicznej do ludzkiej. Tenen ma zatem rację, że sformułowanie "maszyna zdecydowała" nie

powinno automatycznie implikować istnienia moralnej osoby. Pełna teoria przyszłej podmiotowości sztucznej będzie jednak wymagała precyzyjniejszych kryteriów niż proste rozróżnienie człowiek-maszyna. Współczesna filozofia odpowiedzialności bada już sztucznych agentów właśnie poprzez rozdzielanie sprawstwa, kontroli, racji i odpowiedzialności.

Siódma teza brzmi: automatyzacja dotarła do pracy opartej na wiedzy. W źródle pojawiają się przykłady lekarzy, prawników i pisarzy oraz przejście od żywego "ciała" do tekstowego "korpusu", który może stać się materiałem obliczeniowym i treningowym. Jest to jeden z najbardziej empirycznie aktualnych elementów argumentacji Tenena.

AI automatyzuje konkretne zadania, a nie całe zawody

Rewolucja przemysłowa mechanizowała przede wszystkim wybrane operacje fizyczne. Współczesna generatywna AI wkracza natomiast w domeny klasyfikacji, pisania, programowania, analizy dokumentów, tłumaczenia, projektowania i organizowania informacji – obszary, które jeszcze niedawno uznawano za domenę wysokokwalifikowanej pracy umysłowej.

Język "automatyzacji zawodu" jest zbyt gruby.

Najnowszy indeks ILO i NASK, oparty na analizie niemal trzydziestu tysięcy zadań, wskazuje, że około jedna czwarta zatrudnionych na świecie wykonuje zawody wykazujące pewien stopień ekspozycji na generatywną AI; do najwyższej kategorii ekspozycji należy natomiast około 3,3 procent globalnego zatrudnienia. ILO podkreśla, że ze względu na heterogeniczność zadań oraz wciąż niezbędny wkład człowieka, bardziej prawdopodobna jest transformacja wielu profesji niż ich prosta eliminacja. Teza Tenena pozostaje zatem trafną diagnozą przesunięcia granicy automatyzacji, jednak wymaga uzupełnienia o teorię zadaniową: AI nie zastępuje całych zawodów, lecz automatyzuje ich fragmenty, zmieniając wewnętrzną strukturę pracy i redystrybuując kompetencje między człowiekiem, organizacją a kapitałem technicznym.

Collective labor spotyka się z knowledge work. AI automatyzująca fragment pracy analityka sama w sobie zależy od pracy programistów, autorów, ekspertów domenowych i pracowników danych. ILO zwraca uwagę, że systemy AI w środowisku zawodowym napędzają jednocześnie dwa procesy: automatyzację zadań wykonywanych przez pracowników oraz automatyzację funkcji menedżerskich, czyli algorithmic management. W rezultacie nie mamy do czynienia z prostym mechanizmem "maszyna zastępuje człowieka", lecz z głęboką rekonfiguracją organizacji: jedne ludzkie zadania znikają, inne powstają, kolejne zostają ukryte w łańcuchu dostaw, a jeszcze inne ewoluują w stronę nadzorowania i korygowania modelu.

Technologia i regulacje jako nośniki polityki

Ósma teza, "technology encodes politics", wykracza poza ramy samej sztucznej inteligencji. Od Marksa po Langdona Winerę filozofia technologii rozwija argument, że artefakty i infrastruktury są nierozzerwalnie związane ze strukturami władzy, gdyż mogą ucieleśniać określone formy organizacji, ograniczać alternatywne zachowania i premiować konkretne układy instytucjonalne. Współczesne opracowanie Stanford Encyclopedia of Philosophy przywołuje zarówno marksowskie powiązanie technologii ze strukturą stosunków produkcji, jak i Winerowskie hasło "artifacts have politics". Nie należy tego jednak interpretować w duchu prostego determinizmu: maszyna nie posiada programu politycznego tylko dlatego, że jest maszyną. Jej polityczność wynika z projektu, wymagań infrastrukturalnych, modelu własności, procedur użycia oraz układu interesów, w którym funkcjonuje.

AI dostarcza tu szczególnie wyraźnych przykładów. Funkcja celu decyduje o tym, co podlega optymalizacji; system klasyfikacyjny określa, które kategorie pozostają widzialne; dane historyczne wpływają na regularności powielane przez model; interfejs wyznacza zakres możliwych działań użytkownika; model dostępu rozstrzyga, kto może korzystać z danych kompetencji, a infrastruktura obliczeniowa rzutuje na koncentrację rynku. W edukacji system może ułatwiać dostęp do wiedzy, ale jednocześnie standaryzować język. W pracy zwiększać produktywność, lecz równocześnie intensyfikować monitoring. W administracji ograniczać arbitralność urzędnika, by w zamian rozpowszechniać jeden błąd klasyfikacyjny na tysiące decyzji.

Prawo europejskie stanowi praktyczne potwierdzenie politycznego charakteru projektowania AI. AI Act rozróżnia zastosowania według poziomu ryzyka, wymaga transparentności określonych systemów, nakłada obowiązki na dostawców i wdrażających oraz przewiduje human oversight i oceny wpływu na prawa podstawowe w przypadkach wysokiego ryzyka. W odniesieniu do modeli ogólnego przeznaczenia wymusza m.in. politykę zgodności z unijnym prawem autorskim i publiczne streszczenie treści użytych do treningu. Są to decyzje normatywne zapisane w architekturze governance technologii: społeczeństwo definiuje, że określone korzyści techniczne nie mogą być realizowane bez uwzględnienia praw, kontroli i odpowiedzialności. Tenenowska teza wymaga jednak jednego istotnego rozszerzenia.

Jeżeli technologia koduje politykę, to również regulacja koduje politykę. Wybór tego, które systemy sklasyfikować jako wysokiego ryzyka, jak rozumieć transparentność, jakie prawa pozostawić dostawcom, jak chronić innowację, gdzie dopuścić wyjątek i jak rozłożyć koszty zgodności, nie jest technokratycznym dodatkiem do "neutralnego" rozwoju AI. Jest to element projektowania porządku instytucjonalnego. Właściwym przedmiotem analizy staje się zatem nie samo pytanie o

"bias algorytmu", lecz cały układ wzajemnie kodujących się warstw: danych, modelu, platformy, organizacji, rynku i prawa.

Ryzyko inteligencji generycznej i systemowa homogenizacja treści

Dziewiąta i najbardziej prowokacyjna teza głosi, że general intelligence może prowadzić do generic intelligence. Tenen przeciwstawia marzeniu o AGI ryzyko powstania inteligencji generycznej: poprawnej, uśrednionej, stylistycznie bezpiecznej i pozbawionej ostrości wynikającej z indywidualnego doświadczenia. W pierwszym odruchu można uznać tę tezę za literacką polemikę, jednak współczesne badania zaczynają dostarczać jej interesującego wsparcia empirycznego – pod warunkiem precyzyjnego określenia rodzaju obserwowanej homogenizacji.

Badanie opublikowane w Science Advances w 2024 roku wykazało, że choć dostęp do generatywnych pomysłów AI może zwiększać ocenianą kreatywność i jakość opowiadań (szczególnie u słabszych twórców), to teksty powstające z pomocą AI stawały się bardziej podobne do siebie. Autorzy interpretowali to jako napięcie między indywidualnym wzrostem jakości a spadkiem kolektywnej różnorodności. Podobny mechanizm zaobserwowano w badaniach nad brainstormingiem: ChatGPT podnosił przeciętną kreatywność poszczególnych pomysłów, jednocześnie zmniejszając różnorodność całej puli propozycji.

Jeszcze mocniejsze dane dostarczyły analizy 2200 esejów rekrutacyjnych z 2025 roku. Każdy kolejny tekst napisany przez człowieka wносił do zbiorowej puli więcej nowych idei niż tekst wygenerowany przez GPT-4, a różnica w różnorodności rosła wraz z wielkością analizowanego zbioru. Efekt ten utrzymywał się nawet po próbach zwiększenia wariancji modelu poprzez zmiany promptów i parametrów generacji. Nie dowodzi to istnienia uniwersalnego prawa, według którego AI zawsze homogenizuje – eksperyment obejmował przecież konkretne modele i zadania – wskazuje jednak na realny mechanizm ryzyka: system może podnosić poziom jednostkowego wykonania, jednocześnie zmniejszając wariancję populacyjną. Hipotezę tę wzmacniają badania z sierpnia 2026 roku opublikowane w Nature Human Behaviour. Analiza ponad 880 tysięcy tekstów wykazała, że upowszechnienie LLM jako pomocy pisarskiej wiąże się z redukcją zróżnicowania językowego; wariancja złożoności pisania spadała o około 21–50 procent, a modele wzmacniały wzorce dominujące kosztem różnic indywidualnych.

Ma to kluczowe znaczenie dla koncepcji generic intelligence. Metafora Tenena zyskuje empiryczną treść: problemem nie musi być przeciętność pojedynczej odpowiedzi, lecz systemowe przesuwanie

milionów wypowiedzi ku podobnym centrom stylistycznym i konceptualnym. Jednocześnie nie należy ulegać pokusie odwróconego romantyzmu, zgodnie z którym człowiek jest z natury nieskończenie oryginalny, a maszyna banalna. Ludzka kultura również opiera się na gatunkach, stereotypach, kliszach i konwencjach. Analiza Template Culture pokazała, że proces tworzenia od zawsze łączy nowość z odziedziczoną strukturą.

Modele generatywne mogą pomagać jednostkom wydobywać pomysły, których bez wsparcia by nie sformułowały. Co więcej, badania nad tendencjami kulturowymi pokazują, że sam język interakcji wpływa na ujawniane przez model orientacje, więc AI nie musi produkować jednej, jednorodnej kultury globalnej. Właściwym problemem nie jest zatem średnia różnorodność modelu, lecz układ sprzężeń zwrotnych między dominującymi danymi, modelem, wyborami użytkowników a tekstami ponownie trafiającymi do obiegu kulturowego.

Konfrontacja dziewięciu tez z aktualną wiedzą wyłania obraz daleki zarówno od technoutopii, jak i antropologicznego defensywizmu. Pierwsza teza – o pracy zbiorowej – jest bardzo mocna w kontekście socjotechnicznej genezy i ukrytych łańcuchów pracy, lecz nie powinna negować specyfiki obliczeniowych własności wyuczonych przez model. Druga – o rozproszonej inteligencji – ma solidne podstawy kognitywistyczne, ale nie implikuje rozproszonej świadomości. Trzecia – o metaforyczności machine learning – trafnie chroni przed antropomorfizmem, o ile nie odbiera technicznemu uczeniu statusu rzeczywistej adaptacji parametrów i reprezentacji. Czwarta – o metaforach przesłaniających odpowiedzialność – znajduje potwierdzenie w teorii organizacji i prawie, choć odpowiedzialność musi być tu rekonstruowana wielopoziomowo. Piąta teza – o braku bólu – stanowi ważną antropologię ucieleśnionej moralności, ale nie wyczerpuje wszystkich teorii etycznych. Szósta – o braku pełnego moralnego sprawstwa maszyn – jest obecnie stanowiskiem uzasadnionym ostrożnością epistemiczną, choć nie dowodzi metafizycznej niemożliwości powstania przyszłego sztucznego podmiotu. Siódma – o automatyzacji knowledge work – znajduje potwierdzenie w transformacji struktury zadań, lecz jest zbyt radykalna, jeśli zakłada całkowite znikanie zawodów. Ósma – o polityczności technologii – opiera się na filozofii techniki i socjologii, ale musi być rozumiana niedeterministycznie: technologie konfiguruja relacje władzy wspólnie z instytucjami, a nie poza nimi.

AI jako pośrednik między językiem a światem

Dziewiąta teza – dotycząca generic intelligence – pozostaje w części hipotezą kulturową, jednak badania z lat 2024–2026 sprawiają, że nie można jej już traktować jedynie jako metafory. Istnieją

empiryczne przesłanki, by uznać, że pomoc generatywna może jednocześnie podnosić jakość jednostkowej twórczości i obniżać zbiorową różnorodność.

Na tym poziomie ujawnia się pełna architektura książki Tenena. *Letter Magic* ukazało pragnienie generowania wiedzy poprzez formalną transformację znaków, natomiast *Smart Cabinets* przeniosło kompetencję z osoby na artefakt, stawiając problem wykonania zadania bez wewnętrznego rozumienia procesu. *Floral Leaf Pattern* ujawniło, że zapis może nie tylko reprezentować działanie, ale wręcz nim sterować, a *Template Culture* wprowadziła do pracy symbolicznej logikę podziału, standaryzacji i powtarzalności. *Airplane Stories* pokazało możliwość formalizowania narracji jako struktur działań, stanów i celów, zaś *Markov's Pushkin* zastąpił część jawnych reguł statystycznym modelowaniem zależności. Współczesne przestrzenie wektorowe i transformery dopełniły ten obraz, pokazując, jak bogate reprezentacje mogą wyłonić się z probabilistycznego uczenia na ludzkich zapisach.

Dziewięć końcowych idei nie jest zatem jedynie dodatkiem do tej historii; stanowią one próbę politycznego, etycznego i antropologicznego wyjaśnienia płynących z niej wniosków. Najważniejsza korekta, jakiej wymaga myśl Tenena w zderzeniu ze współczesną wiedzą, dotyczy samego przeciwstawienia języka światu. Wielokrotnie przywoływane "więzienie języka" ma ogromną wartość jako ostrzeżenie przed utożsamieniem koherencji z prawdą, lecz język nie jest systemem symboli odciętym hermetycznie od rzeczywistości. Został on wytworzony przez ucieleśnionych ludzi żyjących w świecie, dlatego rozkłady językowe zawierają skondensowane ślady rzeczywistości fizycznej, społecznej i kulturowej. Badania reprezentacji konceptualnych obnażają tę dwuznaczność: modele tekstowe potrafią odtworzyć znaczną część struktury ludzkich pojęć, lecz ich zgodność słabnie tam, gdzie kluczową rolę odgrywa bezpośrednie doświadczenie sensoryczne i motoryczne. Trafniejszym określeniem niż "więzienie" jest zatem pośredniość: model językowy ma dostęp do świata głównie poprzez reprezentacje pozostawione przez innych, chyba że zostanie sprzęgnięty z dodatkowymi modalnościami, narzędziami i środowiskiem działania. Podobnej korekty wymaga romantyczne pojęcie autora.

Jeżeli inteligencja jest rozproszona, a twórczość historycznie czerpała z gatunku, języka, biblioteki, redakcji, notacji i cudzych pomysłów, pytanie o to, czy tekst stworzył człowiek, czy AI, często okazuje się zbyt ubogie. Bardziej adekwatne są pytania o stopnie autorstwa: kto określił cel, kto dokonał selekcji źródeł, kto zbudował argument, kto kontrolował prawdziwość, kto przeformułował wynik i kto odpowiada za publikację – a także jaki rodzaj zależności między ludzkim osądem a automatycznym generowaniem wystąpił w procesie. Tenenowskie *collective labor* nie powinno prowadzić do śmierci autora, lecz do precyzyjniejszej analizy jego funkcji w systemie współtworzenia. Czas sformułować centralną syntezę artykułu.

AI jako rozproszona kompetencja wymagająca nowej ramy odpowiedzialności

Inteligencja nie jest pojedynczą cechą, którą posiada albo człowiek, albo maszyna. To raczej rodzina zdolności obejmująca uczenie się, reprezentowanie danych, przewidywanie, rozwiązywanie problemów, transfer wiedzy, adaptację, rozumowanie, komunikację, planowanie oraz umiejętność działania w warunkach niepewności. Do tego dochodzą właściwości, które tradycja filozoficzna zazwyczaj wiązała z osobą: świadomość, intencjonalność, autonomię, odpowiedzialność, odczuwanie i uczestnictwo w normatywnej wspólnotce.

Współczesna AI potrafi osiągać bardzo wysokie kompetencje w obszarze pierwszej grupy zdolności, nie wykazując przy tym posiadania tych drugich. Błędem jest zatem ekstrapolowanie sukcesu z jednego wymiaru na wszystkie pozostałe. Tu pojawia się druga synteza: jednostką analizy inteligencji nie zawsze powinien być pojedynczy podmiot. Człowiek myśli bowiem za pomocą języka otrzymanego od wspólnoty, korzysta z bibliotek, których nie napisał, instrumentów pomiarowych, których nie skonstruował, i instytucji wiedzy, których procedur sam nie opracował. AI radykalizuje ten stan rzeczy, gdyż potrafi skompresować fragmenty zbiorowego dorobku w interaktywny system dostępny dla jednostki. Pod tym względem jest czymś więcej niż zwykłym narzędziem, lecz czymś innym niż autonomiczna osoba.

Najwłaściwszym przedmiotem badania staje się zatem układ człowiek-model-źródło-instytucja-infrastruktura. Trzecia synteza dotyczy odpowiedzialności. Im bardziej kompetencja staje się rozproszona, tym większe jest ryzyko, że odpowiedzialność także ulegnie rozproszeniu do punktu niewidzialności. Dlatego postulat Tenena, aby "czytać przez artifice" i rekonstruować ludzką pracę ukrytą za inteligentnym interfejsem, należy rozszerzyć na zasadę epistemicznej i instytucjonalnej traceability. Społeczeństwo korzystające z AI potrzebuje nie tylko poprawnych odpowiedzi, ale możliwości ustalenia źródła informacji, celu systemu, granic jego działania, odpowiedzialnego operatora oraz procedury korekty błędów. To właśnie w tym kontekście należy oceniać regulacyjne obowiązki dotyczące dokumentacji, nadzoru człowieka, transparentności i oceny wpływu.

Czwarta synteza ma wymiar ekonomiczny. Automatyzacja wiedzy nie prowadzi nieuchronnie ani do powszechnego bezrobocia, ani do automatycznej emancypacji człowieka od pracy. Jej konsekwencje zależą od tego, jakie zadania są automatyzowane, kto posiada infrastrukturę, kto przejmuje wzrost produktywności, jak rozdzielone są prawa pracownicze i czy technologia jest projektowana jako substytut, czy komplement ludzkiej kompetencji. Aktualne dane ILO wskazują przede wszystkim na transformację zawodów, a nie ich masową likwidację. Tenenowska historia

pracy stanowi zatem dobre antidotum na dyskurs o tym, że "AI zabierze nam pracę", ponieważ nakazuje pytać o zmianę struktury zatrudnienia, jego widzialność i podział korzyści.

Piąta synteza ma charakter pedagogiczny. Jeżeli system może zapewnić użytkownikowi wysokiej jakości wynik bez wewnętrznego przyswojenia procesu, edukacja musi wyraźnie odróżnić sukces wykonawczy od faktycznego nabycia kompetencji. W tym kontekście spór Kirchera z Kuhlmannem okazuje się niezwykle aktualny: demokratyzacja wyniku może być społecznie cenna, lecz nie wolno mylić jej z demokratyzacją rozumienia. Najlepiej zaprojektowana AI edukacyjna nie powinna zatem maksymalizować liczby zadań wykonanych przez ucznia z pomocą systemu, lecz wzrost liczby operacji poznawczych, które uczeń będzie potrafił później przeprowadzić samodzielnie.

AI jako infrastruktura zagrażająca pluralizmowi poznawczemu

Szósta synteza ma charakter polityczny. Technologia generatywna przestaje być pojedynczym produktem, a staje się infrastrukturą. Jeśli ten sam typ modeli zaczyna pośredniczyć w edukacji, wyszukiwaniu informacji, administracji, prawie, medycynie, programowaniu czy tworzeniu kultury, jego wpływ nie może być oceniany jedynie przez pryzmat jakości jednostkowej odpowiedzi. Należy analizować systemowe skutki dla pluralizmu, koncentracji rynku, autonomii profesjonalnej, dystrybucji wiedzy oraz zdolności obywateli do kontroli instytucji. To właśnie na tym poziomie "generic intelligence" nabiera największego znaczenia: nawet wybitny model może okazać się społecznie problematyczny, jeśli jego masowe użycie doprowadzi do redukcji różnorodności sposobów rozumienia świata.

Warto odwrócić jeden z najczęstszych lęków związanych z AI. Najpoważniejszym zagrożeniem nie musi być powstanie maszyny całkowicie odmiennej od człowieka, lecz infrastruktury, która zbyt skutecznie odtwarza to, co w ludzkości najbardziej przeciętne, dominujące i dobrze udokumentowane, by następnie podać tę średnią jako bezosobową kompetencję. Badania nad homogenizacją twórczości i języka dowodzą, że problem ten nie jest już tylko literacką intuicją. Przyszłość inteligencji zbiorowej zależy zatem nie od wzrostu średniej jakości wypowiedzi, lecz od zachowania epistemicznej różnorodności systemu jako całości. Nie oznacza to, że różnorodność ma zastąpić prawdę lub jakość – w nauce nie potrzebujemy tysiąca różnych wartości stałej Plancka dla zachowania pluralizmu. W polityce, etyce, kulturze i interpretacji często istnieje jednak więcej niż jeden racjonalnie obronny sposób przedstawienia problemu. Inteligentna infrastruktura powinna zatem odróżniać domeny, w których celem jest konwergencja ku najlepiej potwierdzonej odpowiedzi, od tych, w których wartość ma eksploracja wielu perspektyw. To meta-rozróżnienie

może stać się ważniejszą kompetencją przyszłych systemów niż sama zdolność generowania kolejnego tekstu.

Po przeprowadzonej analizie staje się jasne, dlaczego Tenen tak konsekwentnie odwołuje się do historii. Historia nie jest dla niego dekoracją technologicznej teraźniejszości, lecz pełni funkcję "epistemicznego deflatora". Zajiradza zmniejsza poczucie absolutnej nowości generowania odpowiedzi ze znaków; Llull ukazuje wcześniejszy "sen o formalnym kombinowaniu pojęć", a Kircher ujawnia eksternalizację kompetencji. Babbage i Lovelace wskazują na oddzielenie instrukcji od wykonania, natomiast Template Culture przypomina, że standaryzacja kreatywności poprzedza komputer. Markow dowodzi, iż probabilistyczne potraktowanie literatury ma ponadstuletnią historię, a Firth i Salton pokazują, że "geometria relacji językowych" nie została wynaleziona wraz z chatbotem.

Taki zabieg nie umniejsza znaczenia obecnej AI, lecz pozwala ustalić, co w niej jest rzeczywiście nowe. Nowością jest przede wszystkim skala integracji: w jednym systemie spotykają się dziś ogromne korpusy, rozproszone reprezentacje neuronowe, dynamiczne kontekstualizowanie, multimodalność, generowanie, wyszukiwanie, wykonywanie kodu i dostęp do narzędzi. Nowa jest dostępność – kompetencje wymagające wcześniej specjalistycznego oprogramowania lub profesjonalnego pośrednika można dziś uruchomić zwykłym językiem. Nowa jest prędkość: koszt produkcji wariantu tekstu, obrazu czy kodu spada dramatycznie. Nowa jest również skala instytucjonalna: ten sam model może uczestniczyć w milionach procesów poznawczych jednocześnie. Genealogia historyczna chroni więc przed mistycyzmem, nie stając się jednak argumentem, że "wszystko już było".

Pewne składniki istniały wcześniej; to ich integracja i umasowienie mogą wytwarzać jakościowo nowe efekty społeczne. Ostatecznie dziewięć tez Tenena można sprowadzić do nadrzędnej zasady metodologicznej: zanim przypiszemy inteligencję pojedynczemu artefaktowi, należy zrekonstruować system relacji, który umożliwił jego działanie. Jest to zasada jednocześnie humanistyczna, naukowa i polityczna. Humanistyczna, gdyż przywraca historię pojęć, tekstów i kultury. Naukowa, ponieważ wymaga oddzielenia metafory od mechanizmu. Ekonomiczna, gdyż ujawnia pracę i własność. Polityczna, bo wskazuje decyzje dotyczące władzy i odpowiedzialności. Etyczna, ponieważ broni różnicy pomiędzy skutecznym zachowaniem a pełną podmiotowością moralną.

AI jako artefakt poznawczy w systemie współzależności

Najdojrzalsze odczytanie Literary Theory for Robots nie prowadzi do tezy, że sztuczna inteligencja „naprawdę jest tylko biblioteką”, podobnie jak nie sugeruje, że w krzemie narodziła się nowa osoba. Prowadzi raczej do stanowiska bardziej wymagającego: współczesne modele to nowa klasa artefaktów poznawczych, których możliwości wynikają z połączenia specyficznych mechanizmów obliczeniowych z gigantycznym zasobem społecznie wytworzonego języka i wiedzy. Ich badanie wymaga równoczesnego zaangażowania informatyki, statystyki, kognitywistyki, filozofii umysłu, socjologii wiedzy, ekonomii politycznej, prawa, historii mediów i teorii literatury. Każda z tych dyscyplin dostrzega inną warstwę tego samego zjawiska; żadna z nich nie jest wystarczająca samodzielnie.

Centralne pytanie artykułu – czym jest inteligencja – można sformułować na nowo po przejściu całej tej drogi. Inteligencja nie jest wyłącznie ukrytą substancją rezydującą w mózgu albo procesorze, ani nie sprowadza się do samego skutecznego wyniku.

Jest zdolnością systemu do przetwarzania różnic, budowania reprezentacji, uczenia się na doświadczeniu, przewidywania, adaptowania działania, używania narzędzi, komunikowania racji i – w przypadku osoby moralnej – ponoszenia odpowiedzialności w świecie innych osób. Niektóre składniki tej architektury mogą być rozproszone poza organizm, inne pozostają nierozdzielnie związane z cielesnością, doświadczeniem i społecznym uznaniem. W tym świetle pytanie "czy AI jest inteligentna?" okazuje się za małe. Nauka powinna pytać: w jakim sensie, w jakich zadaniach, dzięki jakim mechanizmom, przy użyciu czyjej wcześniejszej pracy, pod czyją kontrolą i z jakimi konsekwencjami społecznymi system zachowuje się inteligentnie? Tak postawione pytanie usuwa zarówno magię, jak i lekceważenie. Nie ma potrzeby przypisywania modelowi duszy, aby uznać jego realną zdolność do transformowania wiedzy. Nie trzeba również zaprzeczać technicznej nowości modelu, by dostrzec pracę ludzi stojących za jego kompetencją. Nie trzeba wybierać między technologicznym mesjanizmem i humanistyczną obroną obłąkanego człowieczeństwa.

Bardziej produktywna jest teoria współzależności, w której każda nowa maszyna poznawcza przekształca nie tylko repertuar ludzkich narzędzi, lecz także kryteria, według których ludzie definiują własne kompetencje.

Historia opisana przez Tenena uczy przede wszystkim tego, że granica inteligencji jest ruchoma. Kiedy mnożenie wielocyfrowych liczb wymagało szkolonego rachmistrza, sprawność ta była znakiem wysokiej kompetencji; po upowszechnieniu kalkulatora przestała fascynować. Mechaniczna korekta pisowni była niegdyś technologicznie niezwykła, dziś pozostaje przezroczysta. Wyszukiwarki przejęły część pracy pamięci bibliograficznej i nie nazywamy ich z tego powodu

geniuszami. Generatywna AI przenosi kolejne zadania z obszaru „dowodu inteligencji” do obszaru infrastruktury.

Zjawisko to nie dowodzi, że inteligencja jest fikcją. Pokazuje raczej, że kultura regularnie utożsamia inteligencję z tymi operacjami, których jeszcze nie potrafi zautomatyzować.

Przyszłość pojęcia inteligencji prawdopodobnie nie będzie historią jego ostatecznej definicji, lecz dalszego różnicowania. W coraz większym stopniu będziemy potrzebować odrębnych pojęć dla kompetencji predykcyjnej, reprezentacyjnej, twórczej, społecznej, praktycznej, epistemicznej i moralnej. Będziemy musieli odróżnić zdolność znalezienia rozwiązania od zdolności jego uzasadnienia, dostarczenie informacji od rozpoznania prawdy, symulację troski od doświadczenia troski, generowanie argumentu od przyjęcia odpowiedzialności za argument oraz zbiorową kompetencję infrastruktury od autonomii jednostki.

AI jako etap eksternalizacji ludzkiego intelektu

"Nine Big Ideas" przestają być jedynie dziewięcioma puentami. Stają się drogami prowadzącymi do jednej, znacznie szerszej rekonstrukcji antropologicznej. Sztuczna inteligencja nie jest tu tylko przedmiotem definicji, lecz instrumentem eksperymentalnym, za pomocą którego cywilizacja bada własne pojęcia rozumu, autorstwa, twórczości, pracy, odpowiedzialności, wiedzy i wspólnoty. Im bardziej maszyny upodabniają się do ludzi w sferze zachowań, tym mniej wystarcza nam behawioralna definicja człowieka. Musimy na nowo zapytać o istotę znaczenia i doświadczenia, o naturę autonomii, wartość odpowiedzialności oraz o to, jak kultura zbiorowo wytwarza wiedzę, którą później błędnie przypisuje pojedynczym geniuszom. To najtrwalszy wkład książki Tenena.

Nie polega on na rozwiązaniu problemu sztucznej inteligencji, lecz na zmianie perspektywy, z której go oglądamy. Zamiast patrzeć wyłącznie przed siebie, ku hipotetycznej superinteligencji, Tenen nakazuje spojrzeć wstecz na tysiące lat rozwoju narzędzi symbolicznych i dostrzec ciągłość ludzkiego pragnienia przenoszenia procesów myślowych poza biologiczne ciało. Zamiast pytać tylko o to, czy maszyna zastąpi człowieka, zachęca do badania, ilu ludzi zostało już ukrytych wewnątrz tego, co nazywamy "maszyną". Zamiast zastanawiać się wyłącznie nad tym, czy system potrafi pisać, kieruje uwagę ku temu, czym pisanie było zawsze: współpracą pamięci, języka, społeczeństwa, materii i narzędzia.

Sztuczna inteligencja nie jest zatem końcem historii ludzkiego intelektu, lecz kolejnym etapem jego eksternalizacji. Historycznie człowiek wielokrotnie przenosił funkcje poznawcze poza siebie: pamięć powierzył pismu, rachunek notacji, klasyfikację katalogowi, procedurę formularzowi,

instrukcję programowi, prawdopodobieństwo modelowi, a coraz większą część transformacji językowych systemowi generatywnemu. Każda taka eksternalizacja zwiększała nasze możliwości działania, lecz jednocześnie zmieniała strukturę zależności między człowiekiem a narzędziem. Kluczowe pytanie przyszłości nie brzmi więc, czy proces ten należy zatrzymać.

Brzmi ono raczej: jak eksternalizować kompetencję bez eksternalizowania odpowiedzialności? Jak rozszerzać poznanie, nie niszczyć zdolności samodzielnego osądu i jak korzystać ze zbiorowego intelektu, nie redukując zbiorowości do generycznej średniej? To moment domknięcia argumentacji rozpoczętej od metafory inteligencji. Sztuczna inteligencja okazuje się nie lustrzanym przeciwieństwem człowieka, lecz lustrem historycznie szczególnym: odbija nasze teksty, instytucje, nierówności, odkrycia, stereotypy, sposoby argumentowania i formy współpracy, by następnie zwrócić je w przekształconej, skondensowanej postaci. Jej największym poznawczym znaczeniem może ostatecznie nie być to, że nauczyliśmy maszynę mówić, lecz to, że dzięki niej zostaliśmy zmuszeni do dokładniejszego zbadania, co przez stulecia mieliśmy na myśli, twierdząc, że mówi, myśli, tworzy i wie człowiek.

Sztuczna inteligencja nie jest zatem końcem historii ludzkiego intelektu, lecz kolejnym etapem jego eksternalizacji poza biologiczne ciało. Pozostaje jednak otwarte fundamentalne pytanie: czy potrafimy rozszerzać nasze możliwości poznawcze, nie eksternalizując przy tym samej odpowiedzialności za wydawany osąd? Być może największą wartością AI nie jest to, że nauczyła się ona mówić, lecz to, że zmusiła nas do zdefiniowania na nowo, co właściwie oznacza, gdy mówi i myśli człowiek.

Inspiracje

[1]



"Literary Theory for Robots" Dennis Yi Tenen

2017 | Stanford University Press | ISBN: 9781503601802

Dennis Yi Tenen jest profesorem nadzwyczajnym języka angielskiego i literatury porównawczej na Uniwersytecie Columbia, gdzie współkieruje Centrum Mediów Porównawczych i Laboratorium Inteligencji Narracyjnej. Jest również związany z Data Science Institute. Przed rozpoczęciem kariery akademickiej, Tenen pracował jako inżynier projektowania oprogramowania w firmie Microsoft w grupie Windows, a później był pracownikiem naukowym w Centrum Berkmana Kleina ds. Internetu i Społeczeństwa na Uniwersytecie Harvarda. Uzyskał tytuł doktora literatury porównawczej na Uniwersytecie Harvarda. Jego badania naukowe łączą teorię literatury, historię mediów, humanistykę cyfrową i socjologię wiedzy. W swojej krytycznej pracy Tenen bada kulturowy, polityczny i materialny wymiar technologii obliczeniowej, śledząc historyczne powiązania między systemami pisma przedcyfrowego, kodowaniem tekstu i sztuczną inteligencją, aby naświetlić ludzką pracę w systemach zautomatyzowanych.

Dennis Yi Tenen is an associate professor of English and Comparative Literature at Columbia University, where he co-directs the Center for Comparative Media and the Narrative Intelligence Lab, and holds an affiliation with the Data Science Institute. Prior to entering academia, Tenen worked as a software design engineer at Microsoft in the Windows group and later served as a fellow at Harvard University's Berkman Klein Center for Internet & Society. He earned his Ph.D. in Comparative Literature from Harvard University. His scholarship operates at the intersection of literary theory, media history, digital humanities, and the sociology of knowledge. Tenen's critical work investigates the cultural, political, and material dimensions of computational technology, tracing historical continuities between pre-digital writing systems, text encoding, and artificial intelligence to illuminate human labor within automated systems.

Columbia University

Literatura uzupełniająca

Książki

Autorzy przywoływani w artykule:



[1] "Teoria letteraria per robot"

Dennis Yi Tenen

2024 | Bollati Boringhieri | ISBN: 9788833943305

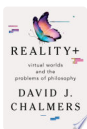
Literatura pokrewna (Google Scholar):



[2] "Cognition in the Wild"

Edwin Hutchins

1995 | MIT Press | ISBN: 9780262082310



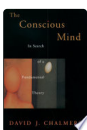
[3] "Reality+"

David Chalmers

W. W. Norton & Company | ISBN: 9780393635812

[4] "David Chalmers: How do you explain consciousness?"

David Chalmers



[5] "The Conscious Mind"

David Chalmers

Oxford University Press | ISBN: 9780199839353



[6] "The Imperative of Responsibility"

Hans Jonas

University of Chicago Press | ISBN: 9780226405971

[7] "City of God"

Augustyn Z Hippony

[8] "Confessions"

Augustyn Z Hippony

[9] "Psalmus contra partem Donati"

Augustyn Z Hippony

[10] "De iudiciis astrorum"

Tomasz Z Akwinu

[11] "Postilla super Psalmos"

Tomasz Z Akwinu

[12] "Super Boetium De Trinitate"

Tomasz Z Akwinu



[13] "Autonomous Technology"

Langdon Winner

1978 | MIT Press | ISBN: 9780262730495



[14] "The Chechen Identity: between social movement and "identity makers"'"

Christoph Kircher

2011 | GRIN Verlag | ISBN: 9783640823482



[15] "Markov chain"

Andrey Markov

Springer Nature | ISBN: 9783030460440

[16] "Markov process"

Andrey Markov

[17] "Markov property"

Andrey Markov



[18] "Humanising Language Education in the Generative Artificial Intelligence Age"

Jérémie Bouchard

2026 | Taylor & Francis | ISBN: 9781040827680



[19] "Machines that Think-History of Artificial Intelligence"

Herman Strange

2023 | ISBN: 9782646678874



[20] "How Artificial Intelligence Influences Future Human Job Market Change"

Johnny Ch Lok

2021 | ISBN: 9798736038251

[21] "Artificial Intelligence and Human Evolution"

Ameet Joshi

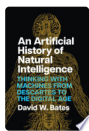
2023 | ISBN: 9798868805936



[22] "Self Aware AI Robot"

Maria Johnsen

2024 | Maria Johnsen | ISBN: 9798333409850



[23] "An Artificial History of Natural Intelligence"

David W. Bates

2024 | Univ of Chicago+ORM | ISBN: 9780226832111

Artykuły naukowe

Autorzy przywoływani:

[1] "Meeting frameworks must be even more inclusive"

Gabriela Serrato Marks, Caroline Solomon, Kaitlin Stack Whitney

2021 | Nature Ecology & Evolution | DOI: 10.1038/s41559-021-01437-9

[Google Scholar](#)

[2] "Author Correction: Meeting frameworks must be even more inclusive"

Gabriela Serrato Marks, Caroline Solomon, Kaitlin Stack Whitney

2021 | Nature Ecology & Evolution | DOI: 10.1038/s41559-021-01583-0

[Google Scholar](#)

[3] "Cognition in the Wild"

Edwin L. Hutchins

1995 | The MIT Press eBooks | DOI: 10.7551/mitpress/1881.001.0001

[Google Scholar](#)

[4] "Distributed cognition in an airline cockpit"

Edwin L. Hutchins, Tove Klausen

1996 | Cambridge University Press eBooks | DOI: 10.1017/cb09781139174077.002

[Google Scholar](#)

[5] "The social organization of distributed cognition."

Edwin L. Hutchins

1991 | American Psychological Association eBooks | DOI: 10.1037/10096-012

[Google Scholar](#)

[6] "Whatever next? Predictive brains, situated agents, and the future of cognitive science"

Andy Clark

2013 | Behavioral and Brain Sciences | DOI: 10.1017/s0140525x12000477

[Google Scholar](#)

[7] "Being There: Putting Brain, Body and World Together Again."

Tim van Gelder, Andy Clark

1998 | The Philosophical Review | DOI: 10.2307/2998391

[Google Scholar](#)

[8] "Natural-Born Cyborgs: Minds, Technologies, and the Future of Human Intelligence"

Mark A. Erickson, Andy Clark

2004 | The Canadian Journal of Sociology | DOI: 10.2307/3654679

[Google Scholar](#)

[9] "An embodied cognitive science?"

Andy Clark, Andy Clark

1999 | Trends in Cognitive Sciences | DOI: 10.1016/s1364-6613(99)01361-3

[Google Scholar](#)

[10] "Does Conceivability Entail Possibility?"

David J. Chalmers

2002 | DOI: 10.1093/oso/9780198250890.003.0004

[Google Scholar](#)

[11] "Extended Cognition and Extended Consciousness"

David J. Chalmers

2019 | Andy Clark and His Critics | DOI: 10.1093/oso/9780190662813.003.0002

[Google Scholar](#)

[12] "A Computational Foundation for the Study of Cognition"

David J. Chalmers

2011 | Journal of Cognitive Science | DOI: 10.17791/jcs.2011.12.4.325

[Google Scholar](#)

[13] "The thinking ape: the enigma of human consciousness"

Steve Paulson, David J. Chalmers, D. Kahneman, L. Santos, N. Schiff

2013 | Annals of the New York Academy of Sciences | DOI: 10.1111/nyas.12165

[Google Scholar](#)

[14] "Does discovery-based instruction enhance learning?"

Louis Alfieri, Patricia Joyce Brooks, Naomi J. Aldrich, Harriet R. Tenenbaum

2010 | Journal of Educational Psychology | DOI: 10.1037/a0021017

[Google Scholar](#)

[15] "Are parents' gender schemas related to their children's gender-related cognitions? A meta-analysis."

Harriet R. Tenenbaum, Campbell Leaper

2002 | Developmental Psychology | DOI: 10.1037/0012-1649.38.4.615

[Google Scholar](#)

[16] "Feeling Like It: A Theory of Inclination and WillSchapiro, Tamar, <i>Feeling Like It: A Theory of Inclination and Will</i> , Oxford: Oxford University Press, 2021, pp. viii + 173, £61 (hardback)."

Sergio Tenenbaum

2023 | Australasian Journal of Philosophy | DOI: 10.1080/00048402.2022.2158891

[Google Scholar](#)

[17] "Cognitive and computational building blocks for more human-like language in machines"

Joshua B. Tenenbaum

2020 | DOI: 10.33774/coe-2020-tsg3x

[Google Scholar](#)

[18] "Essays, Moral, Political, and Literary"

David Hume

2021 | The Clarendon Edition of the Works of David Hume: Essays, Moral, Political, and Literary, Vol. 1 | DOI: 10.1093/oseo/instance.00280189

[Google Scholar](#)

[19] "Essays, Moral, Political, and Literary — Part 1"

David Hume

2021 | Essays, Moral, Political, and Literary | DOI: 10.1093/oseo/instance.00280200

[Google Scholar](#)

[20] "Essays, Moral, Political, and Literary — Part 2"

David Hume

2021 | Essays, Moral, Political, and Literary | DOI: 10.1093/oseo/instance.00280224

[Google Scholar](#)

[21] "The Imperative of Responsibility: In Search of an Ethics in the Technological Age"

Millett Granger Morgan, Hans Jonas

1985 | Journal of Policy Analysis and Management | DOI: 10.2307/3324683

[Google Scholar](#)

[22] "The Imperative of Responsibility: In Search of an Ethics for the Technological Age"

Hans Jonas

1985 | DigitalGeorgetown (Georgetown University Library)

[Google Scholar](#)

[23] "Philosophical Reflections on Experimenting with Human Subjects"

Hans Jonas

1979 | DOI: 10.1007/978-1-4615-6561-1_16

[Google Scholar](#)

[24] "Technology and Responsibility: Reflections on the New Tasks of Ethics"

Hans Jonas

2014 | Palgrave Macmillan UK eBooks | DOI: 10.1057/9781137349088_3

[Google Scholar](#)

[25] "Evaluation as a source of 'strategic intelligence'"

Stefan Kuhlmann

2003 | Learning from Science and Technology Policy Evaluation | DOI: 10.4337/9781781957059.00025

[Google Scholar](#)

[26] "The Beijing Center: Fostering Mutual Understanding between China and the World"

Moritz Kuhlmann

2022 | Journal of Cultural Interaction in East Asia | DOI: 10.1515/jciea-2022-0001

[Google Scholar](#)

[27] "Lectures on the Philosophy of World History"

Georg Wilhelm Friedrich Hegel, Duncan Forbes

1975 | Cambridge University Press eBooks | DOI: 10.1017/cb09781139167567

[Google Scholar](#)

[28] "The Impossibility of Ultimate Moral Responsibility"

Galen Strawson

2008 | Real Materialism | DOI: 10.1093/acprof:oso/9780199267422.003.0014

[Google Scholar](#)

[29] "The impossibility of moral responsibility"

Galen Strawson

1994 | Philosophical Studies | DOI: 10.1007/bf00989879

[Google Scholar](#)

[30] "Do Artifacts Have Politics?"

Langdon Winner

2017 | Computer Ethics | DOI: 10.4324/9781315259697-21

[Google Scholar](#)

[31] "Autonomous Technology: Technics-out-of-Control as a Theme in Political Thought"

Stanley R. Carpenter, Langdon Winner

1978 | Technology and Culture | DOI: 10.2307/3103332

[Google Scholar](#)

[32] "The Whale and the Reactor: A Search for Limits in an Age of High Technology"

William J. McGucken, Langdon Winner

1988 | Technology and Culture | DOI: 10.2307/3105261

[Google Scholar](#)

[33] "Upon Opening the Black Box and Finding It Empty: Social Constructivism and the Philosophy of Technology"

Langdon Winner

1993 | Science Technology & Human Values | DOI: 10.1177/016224399301800306

[Google Scholar](#)

[34] "Human Fear Conditioning and Extinction in Neuroimaging: A Systematic Review"

Christina Sehlmeier, Sonja Schöning, Pienie Zwitserlood, Bettina Pfleiderer, Tilo T. J. Kircher

2009 | PLoS ONE | DOI: 10.1371/journal.pone.0005865

[Google Scholar](#)

[35] "Structural brain changes in schizophrenia at different stages of the illness: A selective review of longitudinal magnetic resonance imaging studies"

Bruno Dietsche, Tilo T. J. Kircher, Irina Falkenberg

2017 | Australian & New Zealand Journal of Psychiatry | DOI: 10.1177/0004867417699473

[Google Scholar](#)

[36] "Selected Works of Ramon Llull (1232-1316)"

J. M. Sobre, Anthony Bonner, Ramón Llull

1986 | Hispanic Review | DOI: 10.2307/473199

[Google Scholar](#)

[37] "Simulation Infrastructure Design on the Basis of the Space Industry's International Standards"

Ludmila Nozhenkova, Olga Isaeva, Aleksey Markov, Andrey Koldy

2017 | Proceedings of the 2017 2nd International Conference on Control, Automation and Artificial Intelligence (CAAI 2017) | DOI: 10.2991/caai-17.2017.28

[Google Scholar](#)

[38] "Artificial Intelligence as a Creative Tool"

Atanas Markov

2024 | Cultural and Historical Heritage: Preservation, Presentation, Digitalization | DOI: 10.55630/kinj.2024.100104

[Google Scholar](#)

[39] "The school of English language and literature, a contribution to the history of Oxford studies"

C. Firth

2009

[Google Scholar](#)

[40] "The SMART and SIRE experimental retrieval systems"

G. Salton, M. McGill

1997

[Google Scholar](#)

[41] "Textual and Visual for Content-based Image Re- 5.2 Data Preparation 5.3 Query Processing 5.3.2 Subquery Processing 4.1 Webssql Syntax Query := Select Targetattr from Tablelist Where Conditions] With-size Size]"

C. Buckley, G. Salton, J. Auto, La Cascia, S. Sethi

2000

[Google Scholar](#)

Artykuły pokrewne (Google Scholar):

[42] "Plain Text"

Dennis Yi Tenen

2017 | Stanford University Press eBooks | DOI: 10.1515/9781503602342

[Google Scholar](#)

[43] "Toward a Computational Archaeology of Fictional Space"

Dennis Yi Tenen

2018 | New Literary History | DOI: 10.1353/nlh.2018.0005

[Google Scholar](#)

[44] "What Is?: Nine Epistemological Essays"

Dennis Yi Tenen

2015 | Design and Culture | DOI: 10.1080/17547075.2015.1051841

[Google Scholar](#)

[45] "Blunt Instrumentalism:"

Dennis Yi Tenen

2016 | University of Minnesota Press eBooks | DOI: 10.5749/j.ctt1cn6thb.12

[Google Scholar](#)

[46] "Plain Text: The Poetics of Computation"

Dennis Yi Tenen

2017

[Google Scholar](#)

[47] "Sustainable Authorship in Plain Text using Pandoc and Markdown"

Dennis Yi Tenen, Grant Wythoff

2014 | The Programming Historian | DOI: 10.46430/phen0041

[Google Scholar](#)

[48] "The Emergence of American Formalism"

Dennis Yi Tenen

2019 | Modern Philology | DOI: 10.1086/705672

[Google Scholar](#)

[49] "Book Piracy As Peer Preservation"

Dennis Yi Tenen, Maxwell Foxman

2014 | Columbia Academic Commons (Columbia University) | DOI: 10.7916/d8w66jhs

[Google Scholar](#)

[50] "Open Letter In Support of the Digital Humanities Studio Space at Butler Library"

Matthew Connelly, Nicholas Dames, Dennis Yi Tenen

2013 | DOI: 10.7916/d8j391xz

[Google Scholar](#)

[51] "Author—Anonymous and Distributed"

Dennis Yi Tenen

2024 | New Techno Humanities | DOI: 10.1016/j.techum.2024.07.001

[Google Scholar](#)

 Treść tworzy Fundacja Dobre Państwo.

Redaguje i publikuje APA ONE, autorski system redakcyjny Fundacji oparty na AI.

APA ONE Publishing System

Fundacja Dobre Państwo © 2026

Article ID: c1c68908-fe6b-45a5-b51d-01e63bb1d925