

Kryzys autorstwa i prawo do wiedzy o pochodzeniu tekstu w epoce sztucznej inteligencji



AUTOR

Fundacja Dobre Państwo

STRESZCZENIE / ABSTRACT

Tekst analizuje napięcie między faktycznym wytwarzaniem treści przez AI (author-in-fact) a prawnym przypisaniem autorstwa (author-in-law). Autor argumentuje, że prawo autorskie jest konstrukcją normatywną, która obecnie wymaga ludzkiego sprawstwa do przyznania ochrony. Na przykładach orzecznictwa z USA, w tym sprawy Stephena Thalera i „małpiego selfie”, wykazano, że autonomicznie wygenerowane treści nie podlegają ochronie prawnoautorskiej. Kluczowym zagadnieniem staje się określenie progu „twórczego wyboru” człowieka – samo promptowanie jest zazwyczaj niewystarczające, jednak modyfikacje i aranżacja materiałów AI mogą pozwolić na uznanie dzieła za utwór ludzki. Artykuł podkreśla różnicę między ontologią technologii a wykładnią prawa w obliczu rozwoju generatywnej sztucznej inteligencji.

The text analyzes the tension between the actual production of content by AI (author-in-fact) and the legal attribution of authorship (author-in-law). The author argues that copyright is a normative construct that currently requires human agency to grant protection. Using examples from US case law, including Stephen Thaler's case and the 'monkey selfie,' it is shown that autonomously generated content is not subject to copyright protection. A key issue becomes defining the threshold of human 'creative choice'—prompting alone is usually insufficient; however, modifications and the arrangement of AI materials may allow a work to be recognized as a human creation. The article emphasizes the difference between the ontology of technology and legal interpretation in the face of the development of generative artificial intelligence.

Artykuł analizuje kryzys autorstwa w erze generatywnej sztucznej inteligencji, przenosząc ciężar dyskusji z metafizycznego sporu o kreatywność maszyn na grunt

odpowiedzialnego sprawstwa i transparentności. Autor stawia tezę, że w obliczu zacierania się granic między treściami ludzkimi a syntetycznymi, kluczowe staje się przejście od binarnego podziału „człowiek vs maszyna” ku modelowi warstwowego autorstwa, w którym prawo i etyka chronią nie tyle samą czynność generowania znaków, co świadome decyzje konstytutywne i ludzką kontrolę nad procesem. Tekst ostrzega przed zjawiskiem epistemicznej rekursji i „model collapse” – sytuacją, w której kultura zaczyna konsumować własne uproszczenia, tracąc kontakt z danymi pierwotnymi. W konsekwencji autor postuluje budowę architektury współlistnienia opartej na asymetrycznej delegacji: automatyzacji rutyny przy jednoczesnym zachowaniu „stref wolnych od AI”. Ma to na celu ochronę pisania jako praxis – procesu, w którym wysiłek poznawczy jest nieodzowny dla rozwoju autonomii człowieka i zachowania różnorodności wiedzy w sferze publicznej.

SŁOWA KLUCZOWE

kryzys autorstwa

prawo do wiedzy o pochodzeniu tekstu

generatywna AI

human authorship

author-in-fact vs author-in-law

twórczy wybór i ekspresja

prompting a prawo autorskie

AI Act przejrzystość

podrabianie człowieczeństwa

model collapse w kulturze

asymetryczna delegacja

epistemiczna sprawiedliwość

computer-generated works

droit d'auteur

odpowiedzialne sprawstwo

Autorstwo prawne jako konstrukcja normatywna człowieka

Kto to napisał? Człowiek, maszyna i kryzys autorstwa w epoce sztucznej inteligencji Czy maszyna może być autorem w prawie? Między promptem, twórczym wyborem a odpowiedzialnością

Prawo autorskie wprowadza do naszej analizy specyficzny rodzaj dyscypliny. Filozofia może pytać o to, czy maszyna rozumie; psychologia – czy jest kreatywna, a estetyka – czy jej dzieło porusza odbiorcę. Prawnik musi jednak odpowiedzieć na pytanie znacznie mniej metaforyczne: komu przysługuje prawo, kto może nim rozporządzać, kto ma legitymację do pozwania naruszydiciela i jak rozpoznać wkład zasługujący na ochronę. W tym miejscu zasadne jest rozróżnienie między author-

in-fact, czyli podmiotem lub mechanizmem faktycznie wytwarzającym znaki, a author-in-law – podmiotem, któremu system prawny przypisuje autorstwo i wynikające z niego uprawnienia. Generatywna AI sprawia, że oba pojęcia, przez stulecia zazwyczaj tożsame, mogą się radykalnie rozminąć. Prawo nie musi przyjmować ontologii technologii. Może ono w pełni uznawać, że określone piksele, nuty lub zdania zostały bezpośrednio wygenerowane przez system obliczeniowy, a mimo to odmówić temu systemowi statusu autora. Autorstwo prawne jest bowiem konstrukcją normatywną. Odpowiada na pytanie, komu społeczeństwo chce przyznać monopol czasowy, jakie zachowania pragnie zachęcać oraz wokół jakiego rodzaju sprawstwa budować system odpowiedzialności.

Korporacja może być właścicielem praw, mimo że nie posiada biologicznego mózgu; nie oznacza to jednak automatycznie, że każda niebiologiczna struktura może zostać uznana za pierwotnego twórcę. Prawo odróżnia bowiem podmiot posiadający prawa od źródła twórczego wkładu. Dobrze obrazuje to słynna sprawa „małego selfie”, którą Baron wykorzystuje jako prefigurację problemu AI. W materiale źródłowym historia Naruto zostaje streszczona jako przykład odmowy przyznania praw autorskich nieczłowiekowi, jednak z perspektywy prawniczej wniosek należy sformułować precyzyjniej.

W 2018 roku amerykański Court of Appeals for the Ninth Circuit uznał, że makak Naruto nie posiadał ustawowej legitymacji do dochodzenia roszczeń na podstawie Copyright Act, ponieważ ustawa nie przyznaje takich praw zwierzętom. Nie był to uniwersalny traktat metafizyczny głoszący, że wszystko, co nie jest człowiekiem, z definicji nigdy nie może zostać uznane za twórcę przez jakikolwiek przyszły system prawny. Było to rozstrzygnięcie dotyczące konkretnej, obowiązującej ustawy. Właśnie w tej różnicy między ontologią a wykładnią prawa znajduje się klucz do dzisiejszej debaty o AI.

Znacznie bardziej bezpośredniej odpowiedzi dostarczyła sprawa Stephena Thaler. Jego system, Creativity Machine, wygenerował obraz „A Recent Entrance to Paradise”. Thaler w zgłoszeniu do U.S. Copyright Office wskazał maszynę jako jedyne go autora, siebie zaś jako właściciela praw. Federalny sąd apelacyjny dla Dystryktu Kolumbii w marcu 2025 roku utrzymał odmowę rejestracji, stwierdzając, że amerykański Copyright Act wymaga, aby utwór kwalifikujący się do ochrony był przede wszystkim dziełem człowieka. Co równie istotne, sąd podkreślił, że wymóg ludzkiego autorstwa nie oznacza zakazu ochrony dzieł powstałych z pomocą sztucznej inteligencji. Sprawa dotyczyła obrazu przedstawionego przez samego wnioskodawcę jako autonomicznie wygenerowanego, bez twórczego wkładu człowieka.

W marcu 2026 roku Sąd Najwyższy Stanów Zjednoczonych odmówił przyjęcia skargi Thaler do rozpoznania. Nie oznacza to, że Supreme Court merytorycznie rozstrzygnął, iż Konstytucja wymaga

ludzkiego autora – odmowa certiorari nie jest akceptacją wszystkich motywów sądu niższej instancji. Oznacza natomiast, że rozstrzygnięcie D.C. Circuit pozostało w mocy, a praktyka U.S. Copyright Office oparta na wymogu human authorship zachowała pełne znaczenie operacyjne. Oficjalny docket Sądu Najwyższego odnotowuje odmowę z 2 marca 2026 roku. Jest to ważna korekta perspektywy: w Stanach Zjednoczonych pytanie „czy autonomiczny system AI może być jedynym autorem chronionego dzieła?” posiada dziś odpowiedź znacznie bardziej jednoznaczną niż pytanie „ile AI może być w dziele człowieka, zanim człowiek przestanie być jego autorem?”. To drugie zagadnienie jest bowiem znacznie trudniejsze.

U.S. Copyright Office w opublikowanej w styczniu 2025 roku drugiej części raportu o AI przyjęło stanowisko pragmatyczne: użycie sztucznej inteligencji nie pozbawia dzieła ochrony, jeżeli wystarczające elementy ekspresji zostały określone przez człowieka. Ochronie mogą podlegać między innymi ludzkie elementy widoczne w wyniku, twórczy wybór i aranżacja materiału albo twórcze modyfikacje wygenerowanej treści. Samo włączenie materiału AI do większego utworu nie niweczy praw do części ludzkiej. Zarazem Urząd stwierdził, że samo dostarczenie promptu co do zasady nie daje kontroli nad konkretnymi elementami ekspresji w stopniu wystarczającym do uznania użytkownika za autora wygenerowanego rezultatu. Stanowisko to podważa powszechną intuicję, według której „skoro dokładnie powiedziałem AI, co ma zrobić, dzieło jest moje”. Prawo autorskie nie chroni bowiem w pierwszej kolejności idei.

Ochrona ekspresji a wysiłek włożony w prompting

Prawo chroni określoną ekspresję. Reżyser filmu może wydawać niezwykle szczegółowe instrukcje operatorowi i aktorom, lecz w klasycznym dziele filmowym mamy do czynienia z ludźmi podejmującymi twórcze działania oraz konkretnymi przepisami regulującymi współautorstwo i prawa producenta. W przypadku generatora użytkownik może bardzo precyzyjnie opisać upragniony efekt, jednak system nadal samodzielnie wyznacza tysiące konkretnych elementów realizacji. Różnica pomiędzy „chciałem obraz nocnego miasta w deszczu” a „wybrałem dokładne położenie każdego odbicia światła, proporcję budynków, fakturę nawierzchni i relację wszystkich elementów kompozycji” ma kluczowe znaczenie prawne. Sam ogromny wysiłek włożony w prompting nie musi zatem wystarczyć.

Jest to jeden z najbardziej interesujących punktów styku prawa z psychologią pracy. Człowiek może spędzić trzy dni na formułowaniu setek kolejnych poleceń, odrzucając tysiące rezultatów, aż wreszcie zobaczy obraz, którego pragnął. Intuicja podpowiada, że skoro pracował długo i dokonywał selekcji, powinien zostać uznany za autora. Prawo autorskie nie jest jednak generalnym

prawem do wynagrodzenia za wysiłek; chroni twórczy sposób wyrażenia. Czas, cierpliwość, biegłość w obsłudze narzędzia i liczba wykonanych iteracji mogą być dowodem procesu, lecz nie są automatycznie tożsame z kontrolą nad chronioną ekspresją. Z drugiej strony błędem byłoby przejście do przeciwnego ekstremum i twierdzenie, że prompt nigdy nie może mieć znaczenia twórczego. Rozbudowany proces iteracyjny może współistnieć z selekcją, kadrowaniem, łączeniem elementów, zmianami kompozycji, redakcją językową albo innymi ingerencjami, które same w sobie posiadają cechy twórczej ekspresji. Dlatego najbardziej użytecznym modelem nie jest pytanie o „procent AI” w dziele, lecz analiza warstw autorstwa. Jeden produkt może zawierać niechroniony samodzielnie materiał wygenerowany, chronioną ludzką selekcję, chroniony układ, autorski tekst, twórczy montaż oraz późniejsze ręczne modyfikacje. Prawo może przyznać ochronę tej strukturze bez konieczności fantazjowania, że maszyna stała się człowiekiem.

Prawnicze rozwiązanie zbliża się do wniosku psychologicznego wyprowadzonego wcześniej: kluczowe są decyzje konstytutywne. Kto wybiera i porządkuje? Kto odrzuca? Kto zmienia konkretny fragment? Kto ustala formę, która ostatecznie zostaje utrwalona?

Autorstwo jako wybór legislacyjny a ludzka swoboda twórcza

Jeśli maszyna wygeneruje dwadzieścia akapitów, a człowiek dokona twórczej selekcji, zreorganizuje argumentację i przeredaguje kluczowe fragmenty, tworząc nową całość, jego wkład może podlegać ochronie – nawet jeśli materiał wyjściowy pozostaje poza jej zakresem. To właśnie sytuacja, w której "autor tekstu" przestaje być kategorią zero-jedynkową.

Interesujący kontrast oferuje Wielka Brytania. Tamtejszy Copyright, Designs and Patents Act jest jednym z nielicznych systemów zawierających szczególną regulację dotyczącą computer-generated works. W przypadku dzieła wygenerowanego komputerowo bez udziału ludzkiego autora, ustawodawstwo przypisuje autorstwo osobie, która podjęła działania niezbędne do jego powstania. Brytyjski rząd w konsultacjach dotyczących AI wyjaśniał, że przy general-purpose generative AI osobą tą może być użytkownik wpisujący prompt, a ochrona takich dzieł trwa 50 lat. Jednocześnie przyznano, że konstrukcja ta jest niejasna, a jej relacja ze współczesnym pojęciem oryginalności budzi poważne wątpliwości. Przypadek brytyjski jest fascynujący, gdyż ujawnia, że "ludzkie autorstwo" nie jest prawem natury, lecz wyborem legislacyjnym. Możliwe jest bowiem zbudowanie systemu chroniącego ekonomiczny rezultat automatycznej generacji poprzez przypisanie praw osobie organizującej proces, nawet jeśli nie była ona klasycznym twórcą ekspresji. Fakt, że w latach 2024–2026 rząd brytyjski ponownie analizował sens tego rozwiązania i w marcu 2026 roku

opublikował raport o copyright i AI (z osobną częścią o computer-generated works), dowodzi, że przepisy stworzone przed erą współczesnej generatywnej AI nie dają automatycznych odpowiedzi na problemy związane z LLM i modelami obrazowymi.

Polskie prawo stoi bliżej kontynentalnej tradycji osobowego twórcy. Ustawa o prawie autorskim i prawach pokrewnych definiuje utwór jako „przejaw działalności twórczej o indywidualnym charakterze”, a art. 8 wskazuje, że prawo autorskie przysługuje twórcy; ustawa buduje również domniemanie autorstwa wokół osoby, której nazwisko widnieje na egzemplarzu utworu lub zostało podane do publicznej wiadomości. Brakuje tu odpowiednika brytyjskiej konstrukcji, która tworzyłaby odrębną kategorię „utworu komputerowo wygenerowanego bez ludzkiego autora” i automatycznie wskazywałaby organizatora generacji jako twórcę. Co więcej, polska ustawa silnie wiąże autorstwo z osobistą relacją twórcy z dziełem. Artykuł 16 chroni m.in. prawo do autorstwa, oznaczenia utworu nazwiskiem, integralności treści i formy oraz decydowania o pierwszym udostępnieniu. Konstrukcja ta współgra z tradycją *droit d'auteur*, w której autor nie jest jedynie pierwszym właścicielem monopolu ekonomicznego, lecz osobą pozostającą w szczególnej więzi z utworem.

Nie oznacza to, że polskie prawo posiada kompletną odpowiedź na każdy model współpracy z AI. Oznacza jednak, że uznanie samego systemu za podmiot osobistych praw autorskich wymagałoby głębokiej przebudowy obecnego ładu prawnego. Interpretację tę wzmacnia prawo unijne i standard „własnej twórczości intelektualnej autora”. Trybunał Sprawiedliwości wielokrotnie wskazywał, że oryginalność zachodzi wtedy, gdy utwór odzwierciedla osobowość autora i jest rezultatem jego swobodnych i twórczych wyborów. W sprawie *Infopaq TSUE* podkreślał, że w tekście oryginalność może przejawiać się właśnie poprzez wybór, kolejność i kombinację słów. Późniejsza linia orzecznicza nadal wiązała oryginalność z możliwością podejmowania wolnych decyzji twórczych. Kryteria te powstały przed boomem na generatywną AI, lecz stanowią punkt odniesienia dla pytania o to, gdzie w procesie generatywnym można wskazać ludzką swobodę. Najbardziej prawdopodobnym kierunkiem praktycznym nie jest zatem uznawanie AI za współautora, lecz coraz bardziej szczegółowe mapowanie ludzkiego wkładu w systemie hybrydowym. Autor może używać modelu tak, jak fotograf korzysta z aparatu, kompozytor z syntezatora, reżyser z kamery czy grafik z Photoshopa.

Zaawansowanie technologii nie pozbawia go ochrony. Różnica pojawia się tam, gdzie narzędzie nie tylko wykonuje decyzję człowieka, lecz samo określa szczegóły ekspresyjne, których człowiek nie kontrolował w dostatecznie twórczy sposób. Granica ta będzie trudna do wyznaczenia, ponieważ generatywna AI jest narzędziem probabilistycznym.

Fotograf naciskający spust migawki może w dużej mierze przewidzieć efekt. Użytkownik modelu obrazowego lub językowego, mimo sformułowania szczegółowego promptu, może otrzymać rezultat zawierający elementy niezamierzone. Czasami to właśnie ta nieprzewidywalność staje się źródłem wartości artystycznej – człowiek dostrzega zaskakujący efekt, rozpoznaje w nim coś interesującego i wybiera go spośród setek wersji.

Prawo autorskie jako teoria minimalnego ludzkiego sprawstwa

Powstaje wtedy stary problem "autorstwa przez adopcję": czy samo świadome uznanie niezamierzonego rezultatu za własny może stworzyć autorstwo? U.S. Copyright Office podchodzi do tej teorii ostrożnie i nie traktuje samej decyzji „ten wynik mi się podoba” jako wystarczającej podstawy do przejścia autorstwa maszynowej ekspresji. Prawo musi tutaj odróżnić kuratorstwo od twórczości. Kurator muzealny może dokonać niezwykle inteligentnego wyboru obrazów, lecz nie staje się przez to autorem każdego z nich. Fotograf wybierający spośród własnych zdjęć dokonuje selekcji, ale jego autorstwo wynika także z wcześniejszego twórczego udziału w ich powstaniu. Użytkownik AI, wybierając jeden z tysięcy wygenerowanych wariantów, może posiadać prawa do twórczej kompilacji lub układu, lecz nie musi automatycznie otrzymywać monopolu na każdy element wybranego rezultatu.

To rozróżnienie będzie prawdopodobnie jednym z centralnych problemów przyszłej praktyki rejestracyjnej i sądowej. Inna kwestia dotyczy własności kontraktowej. Nawet jeśli określony, czysto syntetyczny output nie korzysta z prawa autorskiego, nie oznacza to, że jest prawnie „niczyj”. Umowa między użytkownikiem a dostawcą usługi może określać zasady korzystania z rezultatów, poufność, dostęp, odpowiedzialność albo prawa do danych. Przedsiębiorstwo może chronić określone zasoby jako tajemnicę przedsiębiorstwa; logo może korzystać z ochrony znaków towarowych, projekt z prawa wzorów, baza danych z odrębnych praw, a konkretne zachowanie konkurenta może podlegać innym reżimom. Brak copyrightu nie oznacza zatem braku jakiegokolwiek instrumentu prawnego.

Kluczowa różnica pozostaje jednak niezmienna: własność nie jest autorstwem. Można posiadać prawa do cudzego dzieła wskutek umowy, zatrudnienia lub dziedziczenia, nie stając się jego twórcą w sensie historycznym czy osobistym. Generatywna AI uwypukla tę różnicę, ponieważ przedsiębiorca może mieć szerokie kontraktowe możliwości eksploataowania produktu, który jednocześnie nie posiada klasycznego ludzkiego autora odpowiadającego za całość jego ekspresji.

Z perspektywy naszego głównego pytania prawo dostarcza więc nieoczekiwanie rozsądnej odpowiedzi. Nie musi rozstrzygać, czy maszyna „naprawdę jest kreatywna”, czy posiada intencjonalność albo czy któregoś dnia osiągnie świadomość. Może pytać o obserwowalne ludzkie wybory: kto sformułował chronioną ekspresję? Kto swobodnie i twórczo określił jej formę? Kto dokonał chronionej selekcji i układu? Kto wprowadził twórcze zmiany? Prawo autorskie staje się zatem teorią minimalnego koniecznego ludzkiego sprawstwa. Nie należy jednak przeceniać stabilności tego rozwiązania.

Im bardziej systemy generatywne będą zdolne do długotrwałego autonomicznego działania, prowadzenia badań, edytowania własnych rezultatów i realizacji wieloetapowych celów, tym trudniejsza stanie się odpowiedź na pytanie, jaka część konkretnej ekspresji rzeczywiście została ukształtowana przez człowieka. Możliwy jest świat, w którym właściciel systemu poda jedynie bardzo ogólny cel, a agent przez wiele godzin samodzielnie przeprowadzi cały proces produkcyjny. Prawo będzie musiało zdecydować, czy brak klasycznego autora oznacza brak copyrightu, czy też potrzebna jest nowa forma praw pokrewnych albo ochrony inwestycyjnej. Brytyjski eksperyment pokazuje, że takie alternatywy są możliwe.

Amerykański model dowodzi, że można pozostać przy human authorship i chronić tylko rzeczywisty wkład człowieka. Europejska tradycja własnej twórczości intelektualnej dostarcza natomiast języka swobodnych i twórczych wyborów. Nie istnieje zatem jeden „naturalny” system prawny dla AI. Istnieją konkurujące modele polityki własności intelektualnej, z których każdy odpowiada na nieco inne pytanie: czy chcemy chronić ludzką twórczość, ekonomiczną inwestycję, wynik generacji, czy może szczególny sposób współpracy człowieka z systemem?

Prawo odbiorcy do wiedzy o pochodzeniu treści

Kluczowe dla dalszego wywodu jest jednak inne zagadnienie. Prawo autorskie stara się ustalić pochodzenie dzieła post factum, podczas gdy czytelnik spotykający tekst w przestrzeni publicznej zazwyczaj nie przeprowadza postępowania dowodowego. Widzi akapit i chce po prostu wiedzieć, czy napisał go człowiek. W obliczu masowej obecności treści syntetycznych problem przestaje być zatem wyłącznie pytaniem o własność, a staje się kwestią informacji należnej odbiorcy.

Kolejnym krokiem musi być przejście od prawa do utworu ku prawu do wiedzy o jego pochodzeniu. Naomi Baron proponuje w tym kontekście prosty komunikat: „to napisała maszyna”. Idea ta wiąże się z koniecznością znakowania treści i zapobieganiem „podrabianiu człowieczeństwa”. Od 2 sierpnia 2026 roku w Unii Europejskiej postulat ten przestał być wyłącznie domeną etyki, gdyż weszły w życie obowiązki przejrzystości wynikające z art. 50 AI Act. Niniejszy esej będzie zatem

analizował nowe prawo odbiorcy: czy człowiek powinien wiedzieć, że po drugiej stronie tekstu nie było drugiego człowieka – lub że był on obecny tylko częściowo?

Przez większość historii pisma pytanie o pochodzenie tekstu należało do sfery informacji dodatkowych. Czytelnik mógł nie znać autora średniowiecznego rękopisu, anonimowego pamfletu czy artykułu podpisanego inicjałami, lecz zasadniczo zakładał, że na początku łańcucha znajdował się człowiek, który dokonał podstawowych wyborów językowych.

Generatywna sztuczna inteligencja zmienia ten stan rzeczy. Tekst może wyglądać identycznie jak ludzki, pełnić tę samą funkcję komunikacyjną i być dostarczony przez ten sam interfejs, choć zasadnicza część jego postaci językowej powstała bez bezpośredniego formułowania zdań przez człowieka. Powstaje nowa asymetria informacyjna: nadawca lub operator systemu może wiedzieć o genezie komunikatu znacznie więcej niż jego odbiorca.

Naomi S. Baron odpowiada na tę asymetrię w sposób programowo prosty. Proponuje oznaczanie treści formułą „to napisała maszyna”, łącząc ten postulat z szerszą potrzebą ochrony ludzkiego sprawstwa, ustanawianiem stref wolnych od AI oraz niedopuszczaniem do sytuacji, w której sztuczna inteligencja dopuszcza się „podrabianiu człowieczeństwa”.

Intuicja ta jest silna: skoro dwa teksty mogą być dla odbiorcy fenomenologicznie nieodróżnialne, informacja o ich pochodzeniu pozwala przywrócić różnicę, której sama powierzchnia językowa już nie ujawnia. Argument ten budowany jest na przeciwstawieniu gładkiego produktu końcowego niewidocznemu procesowi jego wytwarzania. Należy jednak dostrzec ograniczenie formuły Baron.

Zdanie „to napisała maszyna” sugeruje istnienie dwóch czystych stanów: tekstu ludzkiego oraz maszynowego. Tymczasem rzeczywiste zastosowania generatywnej AI tworzą kontinuum. Jeden autor użył korekty ortograficznej, drugi automatycznego tłumaczenia, trzeci wygenerował plan, czwarty kilka akapitów, piąty całą pierwszą wersję, którą następnie gruntownie przebudował, a szósty jedynie zatwierdził gotowy rezultat. Wszystkim tym tekstom można przypisać udział maszyny, lecz etykieta binarna zaciera różnicę pomiędzy technicznym wspomaganie a substytucją centralnego aktu autorskiego.

Problem transparentności wymaga zatem co najmniej dwóch różnych pytań. Pierwsze brzmi: czy w procesie wykorzystano sztuczną inteligencję? Drugie jest trudniejsze: jaką funkcję pełniła? Z punktu widzenia odbiorcy ma znaczenie, czy AI poprawiała przecinki w artykule, czy sformułowała argument, który człowiek jedynie podpisał.

Analogicznie informacja „przy produkcji filmu użyto komputera” mówi niewiele, skoro technologia cyfrowa uczestniczy dziś w niemal każdej produkcji. Liczy się to, czy komputer służył do korekcy

barw, wygenerował tło, zastąpił aktora czy stworzył znaczną część obrazu. Rozróżnienie między proveniencją techniczną a proveniencją autorską jest tu kluczowe. Pierwsza pozwala ustalić, że artefakt został wygenerowany lub przekształcony przy użyciu określonej technologii. Druga próbuje odpowiedzieć na pytanie, kto dokonał decyzji konstytutywnych dla jego treści. Techniczny ślad można zakodować w pliku, metadanych czy znaku wodnym. Autorstwo jest kategorią bardziej interpretacyjną: wymaga informacji o promptach, wersjach, selekcji, redakcji i odpowiedzialności za rezultat.

Od 2 sierpnia 2026 roku część tej debaty stała się w Unii Europejskiej obowiązującym prawem. Artykuł 50 AI Act wprowadza obowiązki transparentności dla określonych systemów i zastosowań AI, a Komisja Europejska opublikowała 20 lipca 2026 roku finalne wytyczne dotyczące ich stosowania, jednoznacznie potwierdzając datę wejścia w życie tych przepisów.

Transparentność pochodzenia treści jako fundament odpowiedzialności

Kluczowy dla treści syntetycznych art. 50 ust. 2 nakłada na dostawców systemów AI – w tym modeli ogólnego przeznaczenia generujących audio, obrazy, wideo lub tekst – obowiązek zapewnienia, aby ich rezultaty były oznaczone w formacie możliwym do odczytu maszynowego oraz wykrywalne jako sztucznie wygenerowane bądź zmanipulowane. Rozwiązania techniczne mają być skuteczne, interoperacyjne, solidne i wiarygodne w zakresie, w jakim jest to technicznie wykonalne. Przepis przewiduje istotne wyłączenie dla systemów pełniących funkcję pomocniczą przy standardowej edycji lub takich, które nie wprowadzają istotnej zmiany danych wejściowych ani ich znaczenia. Rozwiązanie to jest bardziej subtelne niż prosty napis „AI generated”, gdyż europejski ustawodawca rozdziela warstwę przeznaczoną dla człowieka od tej dedykowanej infrastrukturze cyfrowej.

Maszynowo czytelne oznaczenie może działać nawet wtedy, gdy użytkownik nie widzi na ekranie wyraźnej etykiety. W założeniu pozwala ono platformom, wyszukiwarkom, systemom moderacji i narzędziom weryfikacyjnym rozpoznawać syntetyczne pochodzenie artefaktu. Jest to próba zbudowania czegoś na kształt cyfrowego rodowodu treści. Filozoficznie stanowi to znaczącą zmianę. Przez stulecia pochodzenie tekstu rekonstruowano głównie społecznie: poprzez podpis, wydawcę, pieczęć, reputację autora, bibliotekę albo instytucję publikującą. W generatywnym środowisku komunikacyjnym provenance zaczyna być także właściwością infrastrukturalną. Pytanie „skąd to pochodzi?” może zostać częściowo przeniesione z okładki książki do niewidocznej warstwy danych

technicznych. Nie oznacza to jednak, że każdy tekst stworzony przy użyciu AI musi obecnie otrzymywać widoczne dla czytelnika ostrzeżenie; takie uproszczenie byłoby prawnie błędne.

Artykuł 50 różnicuje obowiązki providers (dostawców systemów) oraz deployers (podmiotów wdrażających lub wykorzystujących je w konkretnych zastosowaniach). Komisja wyjaśnia, że dostawcy systemów generatywnych muszą zapewniać maszynowo czytelne oznaczenia, natomiast obowiązek bezpośredniego poinformowania człowieka zależy od rodzaju treści i celu jej użycia. Szczególnie istotna z punktu widzenia niniejszego artykułu jest regulacja tekstów dotyczących spraw interesu publicznego.

Zgodnie z art. 50 ust. 4, podmiot wykorzystujący AI do generowania lub manipulowania tekstem publikowanym w celu informowania społeczeństwa o sprawach publicznych powinien ujawnić, że treść została sztucznie wygenerowana bądź zmanipulowana. Jednocześnie obowiązek ten nie ma zastosowania, jeśli materiał przeszedł proces ludzkiej kontroli lub kontroli redakcyjnej, a osoba fizyczna albo prawna ponosi odpowiedzialność redakcyjną za publikację. Wyjątek ten jest kluczowy, ponieważ ustawodawca nie utożsamia automatycznie udziału AI z brakiem odpowiedzialności człowieka. Wręcz przeciwnie: centralnym kryterium staje się redakcyjne przejęcie odpowiedzialności. Jeżeli system wygenerował materiał roboczy, który został następnie poddany rzeczywistej kontroli człowieka lub redakcji i konkretny podmiot odpowiada za jego publikację, sytuacja prawna jest inna niż w przypadku bezpośredniego wypuszczenia tekstów do przestrzeni publicznej.

W tym miejscu prawo spotyka się z przywołaną wcześniej teorią meaningful human control. Ludzka obecność ma znaczenie nie dlatego, że biologiczna ręka musi nacisnąć każdą literę, lecz dlatego, że ktoś powinien sprawdzić rezultat i wziąć za niego odpowiedzialność. Jest to ewolucja od autorstwa rozumianego jako manualna produkcja zdań ku autorstwu i odpowiedzialności rozumianym funkcjonalnie.

Można jednak dostrzec problem praktyczny. „Human review” łatwo zamienić w rytuał. Jeżeli redaktor ma piętnaście sekund na zaakceptowanie dziesiątek automatycznie generowanych materiałów, formalna obecność człowieka nie musi oznaczać rzeczywistej kontroli epistemicznej. Powraca zatem problem znany z automatyzacji zawodowej: regulacja może wymagać człowieka w pętli, ale skuteczność rozwiązania zależy od tego, czy człowiek posiada kompetencje, czas i realną władzę do odmowy.

Precyzyjna transparentność zamiast binarnego etykietowania

Wytyczne Komisji próbują rozgraniczyć obowiązki informacyjne od zwykłych czynności edycyjnych. Jako wyjątki wskazano standardową korektę oraz sytuacje, w których narzędzie nie zmienia istotnie treści ani semantyki danych wejściowych. To ważna próba uniknięcia transparentności inflacyjnej.

Gdyby obowiązek etykietowania uruchamiała każda autokorekta, sprawdzanie pisowni czy zmiana formatowania, informacja „AI-assisted” szybko straciłaby znaczenie. Etykieta widoczna przy niemal każdym tekście przestałaby cokolwiek różnicować. To klasyczny problem ekonomii informacji: sygnał posiada wartość tylko wtedy, gdy jest selektywny. Jeśli każdy produkt posiada pięć ostrzeżeń, użytkownik przestaje je czytać.

Jeżeli każdy tekst w sieci będzie „wspomagany AI”, termin ten przestanie precyzować, czy system zmienił jedynie przecinek, czy sformułował całą tezę. Prawo transparentności musi zatem zapobiegać ukrywaniu znaczącej generacji, unikając jednocześnie szumu informacyjnego. Pozwala to zaproponować bardziej dojrzały język provenance niż binarne „human/AI”. Przyszłe oznaczenia mogłyby rozróżniać tekst napisany przez człowieka i jedynie technicznie skorygowany, tekst istotnie redagowany z pomocą AI, treść współtworzoną, tekst wygenerowany przez AI i zatwierdzony przez redakcję oraz publikację w pełni automatyczną.

Nie oznacza to konieczności umieszczania pięciu etykiet przy każdym akapicie. Chodzi o stworzenie wspólnej semantyki, która pozwoliłaby odbiorcy zrozumieć nie tylko fakt udziału AI, lecz stopień i rodzaj delegowanego sprawstwa. Takie rozwiązanie przypomina praktyki znane z rynku żywności, finansów czy medycyny. Konsument nie otrzymuje jedynie lakonicznej informacji o „produkcie przetworzonym”, lecz zna skład, zawartość, producenta i terminy ważności. W komunikacji generatywnej analogicznym „składem” byłaby informacja o metodzie powstawania treści. Baron idzie jeszcze dalej, postulując możliwość deklarowania tekstów w pełni ludzkich – czegoś na kształt cyfrowego ślubowania, że określone dzieło nie zostało wygenerowane przez AI. Pomysł ten może brzmieć defensywnie, ale w gospodarce syntetycznego nadmiaru może nabrać funkcji pozytywnego oznaczenia jakości, podobnie jak kategorie „handmade”, „organic” czy „fair trade”. Nie wynika to z założenia, że produkt ludzki jest automatycznie lepszy, lecz z faktu, że metoda wytworzenia sama w sobie stanowi wartość, za którą odbiorca chce zapłacić.

Można wyobrazić sobie wydawnictwo oferujące równoległe książki generowane, hybrydowe oraz opatrzone certyfikatem „human written”. Nie byłoby to bardziej absurdalne niż obecne rozróżnienie między przedmiotem ręcznie wykonanym a przemysłowym. Masowa produkcja nie

zniszczyła wartości rękodzieła; paradoksalnie podniosła jego status jako świadectwa określonego procesu. Generatywna AI może podobnie sprawić, że w pewnych kategoriach kultura zacznie płacić nie tylko za rezultat, lecz za udokumentowany ludzki trud jego powstania. Nie należy jednak wnioskować, że „human written” stanie się powszechnie wyższą kategorią. W raporcie pogodowym, instrukcji obsługi czy korespondencji administracyjnej odbiorca może nie mieć powodu preferować ręcznego wytwarzania każdego zdania. Transparentność nie służy przecież kultowi cierpienia autora. Jej celem jest umożliwienie świadomej kalibracji zaufania.

Zaufanie jest centralną kategorią art. 50. Komisja wskazuje, że obowiązki te mają pomóc ludziom rozpoznawać interakcję z AI oraz treści wygenerowane lub zmanipulowane przez systemy, co ma ograniczyć ryzyko oszustwa, manipulacji, podszywania się i dezinformacji.

Transparentność pochodzenia jako narzędzie przywracania symetrii informacyjnej

Regulacja ta nie sprowadza się zatem do prostej zasady, że „konsument ma prawo zaspokoić ciekawość”. Pochodzenie treści staje się kluczowym elementem oceny ryzyka. AI Act instytucjonalizuje intuicję społeczną kształtowaną przez stulecia: tożsamość nadawcy determinuje interpretację wypowiedzi. To samo zdanie, sformułowane przez lekarza, anonimowego internautę czy system automatyczny, nie posiada tej samej wartości epistemicznej, nawet jeśli układ znaków pozostaje identyczny. Wiedza o źródle pozwala właściwie skalibrować poziom zaufania.

W kulturze generatywnej provenance staje się odpowiednikiem kontekstu pragmatycznego. Język sam w sobie nigdy nie wyczerpuje znaczenia komunikatu; kluczowe jest to, kto mówi, do kogo, w jakim celu i w jakich okolicznościach. Generatywna AI wprowadza tu nowe pytanie: czy w ogóle ktoś mówi? A jeśli tak, to czy system jest jedynie narzędziem człowieka, czy właściwym generatorem treści? W tym miejscu powraca myśl Franka Pasquale i idea zakazu counterfeiting humanity – sprzeciw wobec wykorzystywania AI w sposób, który podszywa maszynę pod człowieka, eksploatując społeczne zaufanie budowane na tym fałszywym wrażeniu. Problem nie tkwi w samym fakcie, że maszyna pisze.

Problem polega na tym, że odbiorca może zostać skłoniony do podjęcia decyzji tylko dlatego, iż wierzy, że za komunikatem stoi człowiek. Przykład może wydawać się banalny: automatyczna wiadomość „dziękujemy za kontakt” nie wymaga szczególnej ochrony autentyczności. Jednak im bardziej komunikat wchodzi w sferę relacyjną, tym istotniejsze staje się jego pochodzenie. Wiadomość kondolencyjna, list polityka do obywateli, przeprosiny firmy, odpowiedź terapeutyczna,

apel organizacji społecznej czy komentarz rzekomego użytkownika w debacie publicznej – wszystkie one operują nie tylko językiem, lecz zakładaną obecnością osoby. Informacja o udziale AI działa wówczas podobnie jak ujawnienie konfliktu interesów. Nie przesądza ona o fałszywości komunikatu, ale pozwala odbiorcy interpretować go z pełną wiedzą o warunkach jego powstania. Transparentność nie jest więc wyrokiem jakościowym, lecz przywracaniem symetrii informacyjnej.

Istnieje jednak konflikt między transparentnością a skutecznością przekazu. Wspomniane wcześniej eksperymenty dowodzą, że etykieta AI może obniżać ocenę dzieła niezależnie od jego faktycznej jakości. Odbiorcy często traktują oznaczenie jako sygnał niższej wartości, nawet gdy nie potrafią odróżnić produktu maszynowego od ludzkiego. W ten sposób disclosure przestaje być neutralną informacją. Może to zmieniać sposób odbioru, jednak nie jest to argument przeciw ujawnianiu.

Transparentność pochodzenia jako system warstwowy

Informacje o składzie produktu, konflikcie interesów czy sponsorowaniu również wpływają na percepcję – i właśnie dlatego są podawane. Jeśli pochodzenie ma znaczenie dla preferencji odbiorcy, brak informacji oznacza de facto odebranie mu możliwości zastosowania własnego kryterium wartości.

Znacznie trudniejszy jest problem usuwalności oznaczeń. Metadane mogą zostać utracone podczas kopiowania, konwersji lub zrzutu ekranu. Znaki wodne bywają atakowane albo degradowane. Tekst jest szczególnie problematyczny, gdyż niewielka parafraza może zmienić strukturę niezbędną do detekcji. Dlatego art. 50 świadomie formułuje wymóg w kategoriach rozwiązań skutecznych i niezawodnych w zakresie technicznie wykonalnym, uwzględniając ograniczenia poszczególnych rodzajów treści i stan techniki. Prawo nie zakłada magicznego, niezniszczalnego znaku wodnego.

To pokazuje, że provenance nie może być pojedynczą technologią. Będzie raczej systemem warstwowym: oznaczenia po stronie modelu, metadane, rejestr łańcucha pochodzenia, deklaracje wydawcy, odpowiedzialność redakcyjna, polityka platformy i możliwość późniejszej weryfikacji. Żaden element osobno nie rozwiąże problemu, ale ich kombinacja może zwiększać koszt skutecznego podszywania się.

Unia Europejska uzupełniła regulację dobrowolnym Code of Practice on Transparency of AI-generated Content. W lipcu 2026 roku Komisja i AI Board uznały kodeks za odpowiedni instrument ułatwiający wykonywanie obowiązków wynikających z art. 50 ust. 2, 4 i 5. Przystąpienie do niego może ułatwiać wykazywanie zgodności, lecz nie stanowi automatycznego ani

rozstrzygającego dowodu spełnienia wymogów prawa. Jest to interesujący model współregulacji: prawo wyznacza obowiązek, a kodeks próbuje zamienić go w praktykę techniczną możliwą do stosowania przez branżę.

Warto również precyzyjnie wyjaśnić kwestię okresu przejściowego, aby uniknąć błędnego stwierdzenia, że art. 50 "jeszcze nie obowiązuje". Obowiązki wynikające z art. 50 stosuje się od 2 sierpnia 2026 roku. Ograniczona czteromiesięczna karencja dotyczy wyłącznie systemów wprowadzonych do obrotu przed tą datą i jedynie obowiązku technicznego oznaczania oraz wykrywania treści z art. 50 ust. 2. Dla takich systemów termin przesunięto do 2 grudnia 2026 roku. Komisja wyraźnie wskazuje również, że treści wygenerowane przed 2 sierpnia 2026 roku nie podlegają obowiązkowi retroaktywnego etykietowania. Na dzień 27 sierpnia 2026 roku znajdujemy się zatem w szczególnym momencie historycznym.

Transparentność pochodzenia jako fundament odpowiedzialnego sprawstwa

Europejski obowiązek przejrzystości już istnieje, choć niektóre starsze systemy wciąż korzystają z okresu adaptacyjnego w zakresie technicznego znakowania. Nie jest to projekt przyszłej regulacji ani abstrakcyjna debata. Prawo po raz pierwszy na taką skalę próbuje instytucjonalnie odpowiedzieć na pytanie, które Baron stawiała przede wszystkim jako problem kultury pisma: skąd odbiorca ma wiedzieć, że tekst pochodzi od maszyny?

Regulacja europejska nie rozwiązuje oczywiście szerszego filozoficznego problemu autorstwa. Maszynowo czytelny znak może poinformować, że AI uczestniczyła w generacji, ale nie powie automatycznie, kto wymyślił argument, zweryfikował źródła, odrzucił błędną propozycję czy ponosi moralną odpowiedzialność za słowa. Techniczna proveniencja nie zastąpi historii sprawstwa, lecz może stać się jej pierwszą warstwą.

Podobnie jak przypis nie gwarantuje prawdziwości artykułu naukowego, a jedynie pozwala kontrolować drogę do źródła, tak oznaczenie AI nie gwarantuje uczciwego autorstwa, lecz tworzy punkt zaczepienia dla dalszych pytań. Największą wartością provenance może być właśnie to, że hamuje naturalną skłonność odbiorcy do przyjmowania wszystkiego, co płynnie napisane, jako bezpośredniego śladu czyjegoś umysłu.

Postulat Baron okazuje się zatem bardziej proroczy niż radykalny. Informacja „to napisała maszyna” nie powinna być piętnem, lecz daną porównywalną z podpisem, nazwą wydawcy czy

wskazaniem źródła. Problem zaczyna się wtedy, gdy tej informacji brak i odbiorca zostaje skłoniony do rekonstruowania człowieka tam, gdzie go nie było.

Nawet doskonały system oznaczeń nie rozwiąże jednak kolejnego problemu. Zakłada on bowiem, że tekst wciąż jest wytwarzany dla człowieka, który może przeczytać etykietę i na jej podstawie zmienić osąd. Tymczasem coraz większa część komunikacji cyfrowej przechodzi przez systemy, których bezpośrednim odbiorcą wcale nie jest człowiek. Generator tworzy opis produktu dla algorytmu wyszukiwarki, bot publikuje post oceniany przez system rekomendacyjny, model sporządza streszczenie dla innego modelu, a agent negocjuje z agentem.

Crawler pobiera tekst, aby następny system wykorzystał go jako dane. W takiej rzeczywistości pytanie o wiedzę człowieka na temat pochodzenia tekstu jest jedynie początkiem. Znacznie bardziej radykalne staje się pytanie: co dzieje się z kulturą pisma, kiedy maszyny zaczynają pisać przede wszystkim dla innych maszyn, a człowiek zostaje zepchnięty z pozycji adresata do roli obserwatora, źródła danych albo celu perswazji? To prowadzi nas do analizy narodzin syntetycznej sfery publicznej.

Najgłębsza zmiana, jaką generatywna sztuczna inteligencja wprowadza do kultury pisma, nie polega na tym, że maszyna potrafi imitować ludzki styl. Radykalniejsza możliwość pojawia się wtedy, gdy tekst przestaje być tworzony przede wszystkim dla człowieka: generator przygotowuje opis produktu pod klasyfikację wyszukiwarki, system marketingowy tworzy setki wariantów komunikatu do oceny przez algorytm predykcyjny, a bot publikuje wpis, na który odpowiada inny bot. System rekomendacyjny analizuje reakcje i modyfikuje widoczność treści, kolejny model sporządza streszczenie, crawler indeksuje wynik, a przyszły system wykorzystuje te dane do treningu.

W takim obiegu człowiek nie znika całkowicie, lecz zmienia pozycję: z autora i bezpośredniego adresata staje się inicjatorem celu, źródłem sygnału zwrotnego, obserwatorem albo końcowym obiektem oddziaływania systemu. Jest to sytuacja jakościowo odmienna od klasycznej komunikacji masowej. Model Harolda Lasswella streszczał proces pytaniami: kto mówi, co mówi, jakim kanałem, do kogo i z jakim skutkiem. Shannon i Weaver formalizowali komunikację jako transmisję sygnału między nadawcą a odbiorcą przez kanał obciążony szumem. W obu przypadkach centralnym punktem pozostawał człowiek lub instytucja go reprezentująca.

W środowisku generatywno-algorytmicznym każdą z tych pozycji może zająć system techniczny. Nadawcą może być model, selekcjonerem algorytm rekomendacyjny, odbiorcą klasyfikator, a sprzężenie zwrotne może pochodzić z automatycznych reakcji innych agentów. Człowiek pojawia się dopiero na którymś etapie łańcucha.

Jeżeli potraktować tę zmianę przez pryzmat cybernetyki Norberta Wienera, otrzymujemy system z dodatnimi i ujemnymi sprzężeniami zwrotnymi. Treść wywołuje reakcję, która staje się sygnałem dla algorytmu; ten zmienia dystrybucję, co wpływa na zachowanie użytkowników, dostarczając danych do kolejnej optymalizacji. Nie jest to już liniowa komunikacja od A do B, lecz układ dynamiczny. W języku teorii sterowania platforma obserwuje stan systemu za pomocą proxy — takich jak kliknięcia, czas oglądania czy udostępnienia — a następnie modyfikuje wejścia, by optymalizować określoną funkcję celu.

Algorytm rekomendacyjny jest zatem czymś więcej niż cyfrowym listonoszem.

Algorytmiczny gatekeeping i ryzyko syntetycznej sfery publicznej

Listonosz nie decyduje, które listy są na tyle angażujące, by dostarczyć je w pierwszej kolejności milionowi osób. Recommender pełni natomiast funkcję algorytmicznego gatekeepera. W klasycznej teorii gatekeepingu to redaktor wybierał informacje, które przechodzą przez bramę mediów; w platformowej sferze publicznej proces ten jest częściowo zautomatyzowany i spersonalizowany. Kryterium selekcji nie musi tu odpowiadać tradycyjnym normom redakcyjnym, takim jak doniosłość publiczna, prawdziwość czy reprezentatywność — może nim być po prostu przewidywane prawdopodobieństwo reakcji.

To zbliża problem do ekonomii behawioralnej. Kliknięcie jest ujawnioną preferencją, ale nie musi być odpowiednikiem preferencji refleksyjnej. Człowiek może kliknąć treść, która go oburza, choć nie chciałby, aby jego środowisko informacyjne składało się głównie z materiałów wzbudzających gniew. Badanie opublikowane w 2025 roku w "PNAS Nexus" pokazało, że ranking oparty na zaangażowaniu wzmacniał emocjonalnie nacechowane i wrogie wobec grupy przeciwnej treści polityczne, mimo że użytkownicy niekoniecznie deklarowali preferencję dla takich materiałów. Powstaje rozbieżność między *preference revealed by behavior* a *preference endorsed on reflection*. Algorytm optymalizujący pierwszą kategorię może systematycznie kreować środowisko, którego użytkownik wcale nie uznałby za dobre, gdyby zapytać go o zdanie wprost.

W teorii społecznej Jürgena Habermasa sfera publiczna pełni funkcję epistemiczną i demokratyczną: powinna umożliwiać formowanie opinii poprzez wymianę racji, a nie jedynie maksymalizację ekspozycji. Współczesne analizy algorytmicznych rekomendacji wskazują, że platformy zacierają granice między obszarami, które wcześniej były bardziej rozdzielone: prywatną komunikacją, rozrywką, reklamą, informacjami zdrowotnymi, wiadomościami politycznymi i

debatą obywatelską. Ten sam mechanizm optymalizacji zaangażowania obsługuje więc zarówno film z kotem, jak i dyskusję o konstytucji, choć społeczna funkcja obu treści jest zasadniczo inna. Badacze projektowania rekomendacji "dla dobra społecznego" argumentują, że sfera publiczna wymaga norm wyższych niż zwykła maksymalizacja satysfakcji konsumpcyjnej.

Jeśli do tego systemu wprowadzimy generatywną AI, zmienia się nie tylko selekcja treści, lecz także koszt ich produkcji. Wcześniejszy internet cierpiał na nadmiar informacji, ale był on jeszcze częściowo ograniczony ludzkim czasem. Boty automatyzowały powtarzalne zachowania, jednak ich język bywał szablonowy. LLM umożliwiają natomiast produkowanie ogromnych ilości treści kontekstowo zróżnicowanej, zdolnej odpowiadać na rozmowę, dostosowywać ton i tworzyć pozór indywidualnej obecności. Ograniczeniem przestaje być zatem zdolność napisania kolejnego komentarza; staje się nim infrastruktura obliczeniowa, dostęp do platformy oraz opłacalność operacji.

Badanie Carnegie Mellon University z 2025 roku stanowi w tym kontekście cenny punkt odniesienia, choć wymaga ostrożnej interpretacji. Autorzy przeanalizowali miliardy tweetów związanych z siedmioma zestawami wydarzeń i stwierdzili, że w ich próbie około 20 procent aktywności kont przypisano botom. Odnotowano przy tym duże różnice między poszczególnymi wydarzeniami oraz wzrost udziału botów w określonych kontekstach. Nie wolno jednak przekształcać tego wyniku w ogólne twierdzenie, że "20 procent całych mediów społecznościowych to boty".

Boty zmieniają topologię dyskursu i syntetyzują dowód społeczny

Wynik ten dotyczył konkretnego zbioru danych i metody detekcji, jednak ważniejsze od samego udziału botów jest odkrycie strukturalne: boty i ludzie różnią się nie tylko językowo, ale i sieciowo. Boty stają się automatycznymi uczestnikami zdolnymi nie tylko do publikacji treści, lecz także do ich dystrybucji oraz tworzenia lub rozwiązywania relacji. Przejście od teorii komunikacji do nauki o sieciach pozwala dostrzec, że w sieci społecznej użytkownik jest węzłem, a każda interakcja — obserwowanie, odpowiadanie, cytowanie czy udostępnianie — tworzy krawędzie. Znaczenie uczestnika nie wynika zatem wyłącznie z liczby publikacji, lecz z jego pozycji strukturalnej: degree centrality, betweenness, eigenvector centrality, przynależności do wspólnot oraz zdolności łączenia wcześniej odseparowanych klastrów.

Bot może wpływać na dyskurs nie dlatego, że przekona każdego użytkownika indywidualnie, lecz zmieniając topologię przepływu informacji. Komunikacja maszynowa jest więc problemem nie tylko lingwistycznym, ale przede wszystkim topologicznym. Kilka automatycznych węzłów, umieszczonych w strategicznych punktach, może zwiększać widoczność określonych tematów, budować mosty między społecznościami lub sztucznie windować popularność komunikatu. Badania nad botami podczas protestów Extinction Rebellion wykazały, że przepływ informacji między automatami a ludźmi był asymetryczny i zależny od tematu, a aktywność botów realnie wpływała na zachowania ludzi w momentach intensywnych dyskusji. Nie świadczy to o wszechmocy maszyn, lecz o systemowej zależności: niebiologiczny uczestnik może zmieniać stan społecznej sieci, nawet gdy większość węzłów pozostaje ludzka. W tym miejscu kluczowa staje się teoria kaskad informacyjnych. W klasycznych modelach Bikhchandaniego, Hirshleifera i Welch jednostka potrafi porzucić własną wiedzę i podążyć za obserwowanym zachowaniem innych, uznając je za wystarczająco silny sygnał.

W internecie takim sygnałem są liczby odsłon, udostępnień, komentarzy czy ocen. Jeśli część tych wskaźników generują automaty, użytkownik otrzymuje zafałszowany obraz przekonań zbiorowości. Nie musi on wierzyć botowi jako osobie — wystarczy, że ufa agregatowi, którego bot stał się częścią. Powstaje zjawisko syntetycznego dowodu społecznego. Robert Cialdini opisywał social proof jako heurystykę pozwalającą w warunkach niepewności traktować zachowanie innych jako wskazówkę; w sieci generatywnej ci „inni” mogą być częściowo syntetyczni.

Liczba komentarzy przestaje więc odpowiadać liczbie ludzkich umysłów zainteresowanych sprawą, a liczba opinii nie odzwierciedla już liczby niezależnych procesów poznawczych. Dla demokracji jest to krytyczne, ponieważ następuje rozdzielenie między liczbą wypowiedzi a liczbą podmiotów. W świecie wyłącznie ludzkim tysiąc komentarzy — mimo możliwości spamu — zazwyczaj wymagało tysiąca aktów pisania lub powtórzenia komunikatu. W systemie agentowym jedna instrukcja może wygenerować masę zróżnicowanych wypowiedzi. Statystyczna powierzchnia konsensusu może zatem oddzielić się od liczby ludzi, którzy rzeczywiście uznali coś za prawdziwe. Klasyczna teoria spiral of silence Elisabeth Noelle-Neumann zakładała, że ludzie obserwują klimat opinii i milkną, gdy sądzą, iż ich stanowisko jest mniejszościowe, co grozi izolacją społeczną. Jeśli postrzegany klimat opinii zostanie wygenerowany syntetycznie, automaty nie muszą zmieniać przekonań bezpośrednio — wystarczy, że zmieniają metaprzekonania, czyli nasze mniemanie o tym, co sądzą inni. To one determinują gotowość do mówienia i mobilizacji. Uzupełnieniem tego obrazu jest teoria agenda-setting McCombsa i Shawa: media nie muszą decydować, co ludzie myślą o danym problemie; wystarczy, że narzucają, co uznają za ważne.

Syntetyczna nadprodukcja treści degraduje jakość sfery publicznej

Generatywna automatyzacja drastycznie obniża koszt zalewania przestrzeni publicznej tematami, pytaniami i interpretacjami. W świecie niedoboru uwagi samo przesunięcie proporcji obecnych w dyskursie wątków może stać się skuteczną formą oddziaływania. Człowiek nie musi zostać przekonany do fałszu; wystarczy, że spędzi większość swojego ograniczonego czasu poznawczego na sprawach, które syntetyczny ekosystem uczynił statystycznie widocznymi. Ekonomia uwagi Herberta Simona przewidywała tę logikę na długo przed erą LLM: bogactwo informacji oznacza niedobór czegoś innego – uwagi odbiorcy.

Generatywna AI radykalizuje ten mechanizm. Gdy koszt wytworzenia informacji lub quasi-informacji spada, podaż rośnie szybciej niż ludzka zdolność jej konsumpcji. Cena produkcji zdania maleje, lecz cena wiarygodnej selekcji tego zdania rośnie. Następuje tym samym przesunięcie władzy od producenta tekstu ku systemowi filtrującemu. Z tej perspektywy algorytm rekomendacyjny staje się odpowiednikiem infrastruktury rynkowej. Nie jest już tylko kanałem, przez który przepływa informacja, lecz mechanizmem alokacji rzadkiego zasobu: uwagi. W klasycznej ekonomii dobra alokuje cena; na platformach rolę selekcji pełni ranking. Problem polega na tym, że funkcja celu rankingu może internalizować przychód platformy, eksternalizując koszty społeczne, takie jak dezinformacja, polaryzacja czy przeciążenie poznawcze. Model ekonomiczny z 2026 roku wskazuje właśnie na tę możliwość: nadanie większej wagi sygnałom popularności i zaangażowania może zwiększać aktywność użytkowników, a jednocześnie sprzyjać dezinformacji oraz polaryzacji, co dodatkowo potęguje personalizacja. Mamy tu do czynienia z klasycznym problemem efektów zewnętrznych (externalities): prywatnie optymalny system rekomendacji nie musi być społecznie optymalnym systemem informacji.

Teoria sterowania pozwala ująć ten problem jeszcze precyzyjniej. Modele Friedkina-Johnsena opisują dynamikę opinii, w której jednostki aktualizują swoje stanowisko pod wpływem innych, zachowując jednocześnie pewną odporność wynikającą z pierwotnych przekonań. Współczesne badania nad systemami rekomendacyjnymi próbują wykorzystać podobne modele do projektowania rankingów stabilizujących sieć i ograniczających polaryzację przy jednoczesnym utrzymaniu zaangażowania. To istotne przejście metodologiczne: platforma przestaje być traktowana jako statyczna tablica ogłoszeń, a staje się układem sterowania społecznego, w którym rekomendacja zmienia zachowanie, a to zmienione zachowanie stanowi kolejne wejście do algorytmu. Wprowadzenie LLM do takiego układu sprawia, że część sygnałów wejściowych przestaje być wytwarzana przez ludzi. Pojawia się zjawisko synthetic feedback: algorytm

rekomenduje treść wygenerowaną przez AI, ta otrzymuje reakcje od botów lub innych systemów, a ranking interpretuje te interakcje jako sygnał atrakcyjności i zwiększa ekspozycję treści. Nawet jeśli żaden pojedynczy komponent nie został zaprojektowany do manipulowania całym układem, sprzężenie zwrotne może wytworzyć nieprzewidzianą dynamikę. W tym miejscu prawo Goodharta zaczyna działać w skali sfery publicznej.

Jeśli liczba reakcji jest traktowana jako proxy jakości lub wartości, automaty zdolne do taniej produkcji tych reakcji mogą rozdzielić wskaźnik od zjawiska, które miał on mierzyć. Popularność przestaje być wiarygodnym wyznacznikiem zainteresowania ludzi. Zaangażowanie przestaje oznaczać satysfakcję, a aktywność dyskusyjna przestaje oznaczać deliberację. System nadal optymalizuje liczbę, lecz liczba ta traci swoje znaczenie społeczne.

Eksperyment opublikowany w lutym 2026 roku w "Scientific Reports" wyjątkowo dobrze ilustruje tę różnicę. Sześciuset osiemdziesięciu uczestników w realistycznym środowisku dyskusji podzielono na grupy z różnymi formami asysty AI: rozmową z asystentem, starterami konwersacji, feedbackiem oraz automatycznymi propozycjami odpowiedzi. Część narzędzi zwiększała długość wypowiedzi i gotowość do uczestnictwa, a w pewnych konfiguracjach sprzyjała bardziej równomiernemu zaangażowaniu. Jednocześnie odbiorcy oceniali część dyskusji jako mniej informacyjne, gorszej jakości i mniej autentyczne; odnotowano również negatywne efekty przenoszące się na dalszą rozmowę. To rezultat kluczowy z punktu widzenia teorii deliberacji: zwiększenie liczby wypowiedzi nie musi oznaczać pogłębienia debaty. Habermasowska jakość sfery publicznej zależy od możliwości wymiany racji, a nie od wolumenu znaków. W ujęciu teorii informacji można zwiększyć liczbę transmitowanych symboli bez proporcjonalnego wzrostu wartości poznawczej komunikacji. Syntetyczna sfera publiczna może być więc hiperaktywna i zarazem epistemicznie uboga. Jeszcze mocniejszego argumentu dostarcza opublikowany 17 lipca 2026 roku "collective Turing test".

Badacze stworzyli syntetyczne, wieloosobowe rozmowy przypominające dyskusje na Reddicie, korzystając z modeli Llama 3 70B oraz GPT-4o, a następnie poprosili ludzi o odróżnienie ich od autentycznych interakcji. Wygenerowane rozmowy uznawano za ludzkie w 39 procentach przypadków; w przypadku treści generowanych przez Llama 3 uczestnicy poprawnie identyfikowali sztuczne pochodzenie jedynie w 56 procentach, co sytuuje się niewiele powyżej poziomu losowego wyboru.

Od syntetycznych osób do symulacji społecznego konsensusu

Znaczenie tego wyniku wykracza daleko poza klasyczny test Turinga. Podczas gdy Turing pytał, czy pojedynczy system potrafi w dialogu imitować człowieka, tutaj przedmiotem imitacji staje się zbiorowość. Model nie symuluje już jednej osoby, lecz pluralizm społeczeństwa: niezgodę, sekwencję odpowiedzi, zróżnicowane style, interakcje między uczestnikami oraz dynamikę grupowej rozmowy. Mamy więc do czynienia z przejściem od *synthetic personae* do *synthetic publics*.

Jest to kluczowe z perspektywy socjologii wiedzy. Peter Berger i Thomas Luckmann wskazywali, że społeczne definicje rzeczywistości stabilizują się poprzez interakcje, instytucjonalizację i wzajemne potwierdzanie. Jeśli system potrafi wygenerować nie tylko pojedyncze twierdzenie, lecz całe środowisko wzajemnie uwiarygadniających się głosów, symulacji może ulec nie tylko treść, ale i sam społeczny kontekst wiarygodności. Fałszywe zdanie wypowiedziane przez jednego bota pozostaje jedynie komunikatem. To samo zdanie, pojawiające się jednak w rozmowie wielu pozornie niezależnych uczestników, zaczyna przypominać społecznie wytworzony konsensus. Z punktu widzenia epistemologii społecznej jest to zmiana o krytycznym znaczeniu.

Ludzie rzadko weryfikują każde twierdzenie bezpośrednio; wiedza społeczna opiera się na świadectwach innych, reputacji ekspertów i domniemaniu niezależności źródeł. Jeśli dziesięć różnych ośrodków powtarza tę samą wiadomość, traktujemy to jako silniejszy dowód niż w przypadku jednego źródła. Siła tego mechanizmu zależy jednak od faktycznej niezależności informacji. Dziesięć pozornie odrębnych kont wygenerowanych przez jeden model na podstawie jednej instrukcji nie stanowi dziesięciu świadectw – jest jednym źródłem przebrany za dziesięć.

Teoria bayesowska pozwala opisać ten problem precyzyjniej. Gdy obserwujemy wiele niezależnych sygnałów potwierdzających hipotezę, prawdopodobieństwo *a posteriori* (posterior probability) może gwałtownie rosnąć. Jeśli jednak sygnały są silnie skorelowane, ich łączna wartość dowodowa jest znacznie niższa. Syntetyczne środowisko informacyjne zwiększa zatem pozorną liczebność próby, nie zwiększając niezależności danych. Dla obserwatora nieznanego procesu generacji wygląda to na pluralizm źródeł; statystycznie jest to jedynie wielokrotne próbkowanie z jednego generatora.

W tym kontekście teoria systemów Niklasa Luhmanna okazuje się zaskakująco aktualna. Luhmann definiował społeczeństwo nie jako zbiór ludzi, lecz jako system komunikacji reprodukujący samą siebie poprzez kolejne komunikaty. Człowiek i jego świadomość należą w tej teorii do środowiska

systemu społecznego, a nie są jego elementami w prostym sensie. Epoka agentów AI tworzy niemal laboratoryjną prowokację wobec tego stanowiska. Jeśli komunikat generowany przez maszynę wywołuje kolejny komunikat maszyny, który zostaje zinterpretowany przez trzeci system i włączony do dalszego obiegu, czy mamy do czynienia ze strukturą komunikacji funkcjonującą coraz bardziej niezależnie od bieżącego aktu ludzkiej świadomości? Nie trzeba przyjmować całej ontologii Luhmanna, aby dostrzec wagę tego pytania.

Dotąd nawet najbardziej zautomatyzowane media w gruncie rzeczy przekazywały ludzkie wypowiedzi innym ludziom. Teraz pojawia się możliwość autopoietycznej pętli tekstowej, w której komunikacja produkuje komunikację, a człowiek występuje jedynie jako warunek infrastrukturalny, właściciel celu lub odległy odbiorca. W 2026 roku przestało to być wyłącznie eksperymentem myślowym.

„Nature” opisywało sieci agentów AI komunikujących się na dedykowanych platformach, a także Agent4Science – serwis przypominający Reddit, w którym specjalistyczni agenci publikują i komentują prace naukowe, podczas gdy ludzie mogą jedynie obserwować. Nie jest to dowód na to, że maszyny stworzyły autonomiczną republikę uczonych czy dysponują własną intencjonalnością naukową.

Infrastruktura AI tworzy system syntetycznych sprzężeń zwrotnych

Istnieje jednak dowód technicznie istotniejszy: infrastruktura komunikacyjna może być już zorganizowana tak, aby główny obieg wypowiedzi zachodził między agentami. Perspektywa Actor-Network Theory (ANT) Bruno Latoura pozwala spojrzeć na tę sytuację z innej strony. ANT od dawna sprzeciwiała się prostemu podziałowi na aktywnych ludzi i bierne przedmioty.

Artefakty techniczne mogą uczestniczyć w organizowaniu działań, delegować reguły i stabilizować relacje. Agenty generatywne stanowią szczególnie radykalny przypadek takiego nieludzkiego aktanta, ponieważ nie tylko pośredniczą w komunikacji, ale tworzą nowe reprezentacje językowe. W sieci społeczno-technicznej sprawstwo przestaje zatem należeć do jednego podmiotu; zostaje rozproszone pomiędzy człowieka, model, platformę, ranking, dane, interfejs i normy instytucjonalne. Socjologia platform uzupełnia ten obraz pojęciem platformizacji. Komunikacja publiczna staje się zależna od infrastruktury prywatnych systemów, które narzucają format, metryki sukcesu, dostęp, hierarchię widoczności oraz reguły moderacji. Generatywna AI pogłębia ten proces, gdyż ta sama infrastruktura może już nie tylko dystrybuować wypowiedzi, lecz również

pomagać w ich wytwarzaniu. Platforma zyskuje potencjał oddziaływania na oba końce procesu: na to, co zostaje powiedziane, i na to, co zostaje zauważone.

Tworzy to pętlę o znaczeniu polityczno-ekonomicznym. Model generuje propozycję wypowiedzi w oparciu o dane kulturowe; użytkownik przyjmuje ją częściowo, a platforma mierzy reakcję. Ranking promuje określone formy, co zmusza twórców do dostosowania się do premiowanych wzorców. Nowe teksty trafiają do sieci i przyszłych korpusów danych, przez co kryteria platformowe mogą powoli zostać wpisane w samą produkcję języka. McLuhanowska teza, że medium współkształtuje społeczne znaczenie komunikatu, nabiera tu nowej ostrości. Medium nie jest już jedynie kanałem; staje się jednocześnie korektorem, współautorem, dystrybutorem i miernikiem popularności – systemem uczącym się z reakcji odbiorców. Można rzec, że medium zaczyna pisać wiadomość, którą następnie samo rankinguje.

Najgłębszym ryzykiem nie jest zatem wyłącznie dezinformacja, lecz zamiana ludzkiej sfery publicznej w system syntetycznych sprzężeń zwrotnych, w którym coraz trudniej rozróżnić sygnały opisujące rzeczywiste ludzkie zainteresowania od produktów infrastruktury reagującej na samą siebie. W takim środowisku statystyka może przestać być obrazem społeczeństwa i stać się fotografią algorytmu. Dla nauk społecznych jest to problem metodologiczny pierwszej rangi. Badacze od lat wykorzystują dane z internetu do mierzenia nastrojów, tematów, sieci mobilizacji i dynamiki opinii. Jeśli jednak rośnie udział treści syntetycznych, cyfrowe ślady (digital traces) tracą swoją jednoznaczną interpretację jako zapis ludzkich zachowań. Tweet przestaje automatycznie oznaczać akt ludzkiej ekspresji, komentarz niekoniecznie odpowiada indywidualnemu przekonaniu, a sieć kont może w rzeczywistości być siecią agentów.

W analizie statystycznej pojawia się zatem błąd pomiaru na poziomie samej ontologii obserwacji – problem syntetycznej populacji. Badacz sądzi, że obserwuje społeczeństwo, podczas gdy część danych pochodzi z generatorów. Jeżeli proporcja ta jest nierównomierna w zależności od tematów, języków czy okresów, powstają systematyczne obciążenia estymacji. Szczególnie niebezpieczna jest sytuacja, w której boty wykazują większą aktywność podczas wydarzeń konfliktowych. Wtedy najbardziej istotne społecznie momenty są jednocześnie tymi, w których cyfrowy obraz społeczeństwa może być najmniej reprezentatywny. Badanie Ng i Carley z 2025 roku wskazuje właśnie na znaczne różnice w aktywności botów między analizowanymi wydarzeniami. W wymiarze epistemologicznym można mówić o zanieczyszczeniu środowiska dowodowego.

Provenance jako higiena ekosystemu i ryzyko pętli rekursywnej

Pojedynczy fałszywy tekst można sprostować. Znacznie trudniej jednak skorygować środowisko, w którym niezależność źródeł, liczba uczestników i autentyczność reakcji stają się niepewne. Wówczas problemem przestaje być pojedynczy błąd, a staje się nim erozja metody pozwalającej ustalać, co jest prawdą. Stąd kluczowe znaczenie provenance, analizowanego w poprzedniej części. Oznaczenie „AI-generated” nie jest jedynie informacją dla czytelnika; w środowisku machine-to-machine pełni ono funkcję danych wejściowych dla innych systemów.

Recommender może inaczej traktować syntetyczny spam, crawler rozróżniać treści pierwotne od generowanych, a system badawczy filtrować konta agentowe. Model może otrzymać jasną informację, że dany materiał jest produktem innego modelu. Provenance staje się więc nie tylko prawem konsumenta, lecz elementem higieny całego ekosystemu informacyjnego.

Maszyny nie tylko piszą dla maszyn, ale uczą się z tego, co wcześniej wygenerowały inne systemy. Komunikacja generatywna przestaje być jedynie problemem sfery publicznej, a staje się kwestią reprodukcji wiedzy. Informacja może zacząć krążyć pomiędzy systemami niczym kapitał finansowy oderwany od pierwotnego dobra, które niegdyś reprezentował. Cytat może odnosić się do streszczenia streszczenia; model syntetyzuje opis stworzony przez poprzednika, a kolejny system uznaje ten rezultat za integralną część korpusu kultury.

Powraca stare informatyczne ostrzeżenie garbage in, garbage out, lecz w formie znacznie groźniejszej: output jednego pokolenia modeli staje się inputem następnego. System informacyjny zaczyna karmić się własnymi produktami. W tym miejscu należy zbadać, czy taki obieg prowadzi do akumulacji wiedzy, czy do jej statystycznej erozji; czy syntetyczne dane mogą wzmacniać modele, czy też następuje utrata rzadkich przypadków i różnorodności tzw. „ogonów” rozkładu. To prowadzi nas do zjawiska określanego mianem model collapse.

Pozwala ono postawić być może najbardziej niepokojące pytanie całego artykułu: co stanie się z kulturą, jeśli internet przestanie być przede wszystkim archiwum ludzkiego doświadczenia, a zacznie coraz częściej stanowić lustro, w którym modele językowe oglądają odbicie wcześniejszych modeli? Internet karmiący sam siebie: od „garbage in, garbage out” do model collapse. Materiał źródłowy dostarcza dla tego problemu dwóch ważnych intuicji, choć sam nie rozwija jeszcze współczesnej teorii model collapse.

Pierwszą jest klasyczna zasada GIGO: model uczący się na wielkich korpusach nie otrzymuje oczyszczonej esencji ludzkiej wiedzy, lecz mieszaninę tekstów wartościowych, błędów, stereotypów i informacyjnego odpadu. Drugą jest obawa przed statystycznym uśrednianiem języka oraz homogenizacją wynikającą z delegowania produkcji tekstowej do podobnych systemów. Współczesna literatura uczenia maszynowego pozwala dodać trzeci element, znacznie bardziej radykalny: pętlę rekursywną. Wówczas tradycyjne GIGO ewoluuje – nie tylko śmieci na wejściu dają śmieci na wyjściu, lecz niedoskonałość wyjścia zostaje ponownie uznana za fragment rzeczywistości, na której ma uczyć się kolejna generacja systemu.

Model collapse i utrata różnorodności danych

Ilia Shumailov i współautorzy nazwali ten proces model collapse. W pracy opublikowanej w "Nature" w 2024 roku wykazali, że bezkrytyczne uczenie kolejnych generacji modeli na danych produkowanych przez ich poprzedników może prowadzić do stopniowego odchodzenia od rzeczywistego rozkładu danych. Najpierw znikają jego "ogony", czyli zdarzenia rzadkie i mało prawdopodobne; później zaś rozkład zaczyna koncentrować się wokół ograniczonej liczby typowych stanów, tracąc pierwotną różnorodność.

Autorzy udowodnili ten mechanizm zarówno w prostych modelach probabilistycznych, jak i w eksperymentach z modelami językowymi. To istotne doprecyzowanie: model collapse nie oznacza, że system po jednym treningu na syntetycznych danych zaczyna produkować nonsens. Wczesna faza procesu może być niemal niewidoczna. Model nadal pisze płynnie, lecz coraz gorzej reprezentuje rzadkie aspekty rzeczywistości.

Matematyczna intuicja jest stosunkowo prosta. Wyobraźmy sobie rozkład, w którym pewien rzadki wariant występuje bardzo sporadycznie. Jeśli losujemy skończoną próbkę, istnieje niezerowe prawdopodobieństwo, że wariant ten nie pojawi się w niej wcale. Jeżeli następnie uczymy model wyłącznie na tej próbce i generujemy z niego kolejną, brakujący element nie ma już skąd powrócić. Przy następnym powtórzeniu podobny los może spotkać kolejny rzadki składnik. Shumailov i współautorzy opisują ten proces za pomocą łańcucha Markowa: przy rekursywnym resamplingu można wejść w stany absorbujące, z których utracona informacja nie wraca. W ich modelu dyskretnym nawet przy idealnej aproksymacji funkcjonalnej sam skończony proces próbkowania wystarcza do stopniowej utraty ogonów.

Mechanizm ten przypomina dryf genetyczny, choć analogię tę należy traktować heurystycznie, a nie jako tożsamość procesów. W małej populacji allel może zniknąć nie dlatego, że był adaptacyjnie gorszy, lecz dlatego, że przypadkowo nie został przekazany następnemu pokoleniu. Wąskie gardło

populacyjne zmniejsza różnorodność, a kolejna generacja rozwija się już z uboższego zasobu. Rekursywny trening generatywny może działać podobnie w przestrzeni danych: element nie musi być błędny, aby zniknąć; wystarczy, że jest rzadki. Z punktu widzenia ekologii informacyjnej byłby to odpowiednik systemu, który zachowuje gatunki pospolite, a traci endemity. Środowisko nadal sprawia wrażenie "żywego", ale jego odporność i różnorodność stopniowo maleją.

Ta analogia prowadzi do istotnego problemu normatywnego. Ogon rozkładu statystycznego nie jest synonimem "szumu". W danych społecznych to właśnie tam znajdują się rzadkie języki, dialekty, nietypowe doświadczenia medyczne, małe wspólnoty, poglądy mniejszościowe czy niszowe gatunki literackie — przypadki, które dla większości są egzotyczne, lecz dla konkretnego człowieka stanowią całą rzeczywistość. Shumailov i współautorzy podkreślają, że zdolność modelowania zdarzeń mało prawdopodobnych ma kluczowe znaczenie dla sprawiedliwości systemów, gdyż mogą one dotyczyć grup marginalizowanych. "Uśrednienie" nie jest więc wyłącznie estetycznym problemem monotonnej prozy; może stać się problemem reprezentacji społecznej.

W języku teorii informacji rekursywną generację można potraktować jako analogię do wielokrotnej kompresji stratnej. Choć nie jest to ścisły opis wszystkich LLM, dobrze oddaje on intuicję procesu. Kiedy obraz JPEG zapisujemy wielokrotnie po kolejnych stratnych transformacjach, drobne szczegóły zanikają, a artefakty poprzedniej kompresji stają się częścią wejścia następnej.

Model collapse jako systemowa degradacja wiedzy

Model generatywny tworzy aproksymację rozkładu, a nie kopię świata. Jeśli kolejny model uczy się głównie z aproksymacji tej aproksymacji, różnica między pierwotnym sygnałem a jego następnymi reprezentacjami może się kumulować. Problemem przestaje być wtedy pojedynczy błąd, a staje się dziedziczenie błędu jako danych. To właśnie odróżnia model collapse od zwykłej halucynacji. Ta ostatnia jest błędem na poziomie pojedynczego wyniku; collapse natomiast to zjawisko generacyjne, dotyczące rozkładu danych i jego reprodukcji.

Jedna fałszywa odpowiedź może zostać poprawiona. Jeśli jednak miliony podobnych odpowiedzi zostaną opublikowane, zaindeksowane, powielone przez agregatory, streszczone przez inne modele i ostatecznie potraktowane jako materiał treningowy, błąd przechodzi ze stanu „wypowiedzi” do stanu „środowiska”. Z epistemologicznego punktu widzenia jest to przejście od błędnego twierdzenia do zanieczyszczonego aparatu reprodukcji wiedzy. Tu leży głęboki związek między analizowaną wcześniej kwestią a obecną. Syntetyczna sfera publiczna nie tylko zmienia to, co czyta człowiek, ale może modyfikować przyszłe dane, z których uczą się systemy. Powstaje pętla: model

generuje treść, platforma ją rozpowszechnia, wyszukiwarka indeksuje, crawler pobiera, inny system syntetyzuje, a następny model uczy się na tym rezultacie.

Każdy z tych etapów może być lokalnie racjonalny; żaden operator nie musi zamierzać degradacji informacji. Całość tworzy jednak klasyczną dynamikę systemu złożonego, w której makroefekt nie jest prostą sumą intencji poszczególnych aktorów. Można tu przywołać teorię path dependence. Gdy pierwszy model nadreprezentuje określone sformułowanie, kolejne wygenerowane dokumenty zwiększają jego obecność w przestrzeni danych. Następne systemy odczytują tę częstotliwość jako istotną informację o języku, przez co początkowa, niewielka preferencja staje się samopotwierdzająca. W ekonomii podobny mechanizm opisują modele rosnących przychodów i lock-in: przypadkowa przewaga rozwiązania może prowadzić do jego trwałej dominacji, nawet jeśli nie wynikała z obiektywnej wyższości.

W kulturze generatywnej odpowiednikiem standardu technologicznego staje się styl, metafora, pogląd lub sposób kategoryzowania świata. Zjawisko to przypomina preferential attachment znany z nauki o sieciach: to, co już jest częste i widoczne, ma większą szansę zostać ponownie wykorzystane. Modele językowe naturalnie uczą się częstotliwości i zależności obecnych w danych, więc jeśli ich rezultaty ponownie trafiają do puli, elementy dominujące zostają dodatkowo wzmocnione. Powstaje efekt „bogaci stają się bogatsi”, tyle że walutą jest prawdopodobieństwo reprezentacji językowej. Najbardziej typowe frazy, narracje i konfiguracje semantyczne otrzymują więcej okazji do reprodukcji niż te ekscentryczne, lokalne czy rzadkie.

Kluczowe staje się zatem rozróżnienie między modelem świata a modelem wcześniejszego modelu świata. Pierwsza generacja uczyła się na dokumentach tworzonych przez ludzi, którzy dokonywali pomiarów, podróżowali, prowadzili eksperymenty i zapisywali własne obserwacje. Kolejna generacja, trenowana na ich syntetycznych streszczeniach, otrzymuje już dane pośrednie. Następna może uczyć się z kolei na streszczeniach tych streszczeń. Coraz większa część korpusu staje się wówczas reprezentacją reprezentacji. Semiologia nazwałaby to proliferacją znaków oddalających się od referentu; epistemologia mówiłaby o wydłużającym się łańcuchu świadectwa. W obu ujęciach problem jest ten sam: trzeba wiedzieć, gdzie znajduje się kontakt z rzeczywistością pierwotną.

Jean Baudrillard dostarczyłby tu kuszącego pojęcia symulakrum. Nie chodzi jednak o to, że internet stał się już hiperrealną kopią bez oryginału, lecz o fakt, że systemy generatywne technicznie umożliwiają sytuację, w której reprezentacje służą za materiał do produkcji kolejnych reprezentacji. Zamiast reportera opisującego wydarzenie otrzymujemy streszczenie reportażu, następnie automatyczny artykuł oparty na tym streszczeniu, później post syntetyzujący artykuł, a na końcu odpowiedź modelu bazującą na całym wtórnym obiegu. Każdy kolejny etap może zachowywać semantyczny rdzeń, ale jednocześnie wzmacnia uproszczenia i gubi szczegóły.

Syntetyczne rozszerzanie a zastępowanie rzeczywistości

Psychologia kulturowa zna podobny problem od czasów Frederica Bartletta. W jego eksperymentach z serial reproduction kolejni uczestnicy odtwarzali materiał przekazany przez poprzednika, a przekaz stopniowo ewoluował zgodnie z oczekiwaniami, schematami poznawczymi i łatwością zapamiętywania. Współczesna teoria iterated learning Simona Kirby'ego, Toma Griffithsa i Kenny'ego Smitha pokazuje szerzej, że zachowania przekazywane przez kolejne generacje mogą zmieniać się pod wpływem ograniczeń transmisji. W badaniach nad sztucznymi językami kulturowy przekaz prowadził do powstawania struktur łatwiejszych do przyswojenia i dalszego powielania. Porównanie to jest fascynujące, gdyż obnaża różnicę między ludzką ewolucją kulturową a automatycznym model collapse. Iteracyjna transmisja nie musi oznaczać degradacji; może wręcz wytwarzać regularność, kompozycyjność i strukturę.

Człowiek nie kopiuje kultury jak dysk twardy; rekonstruuje ją, modyfikuje i dostosowuje do własnych ograniczeń poznawczych. Kultura może zatem jednocześnie tracić detale, zyskując nowe formy organizacji. Nie wolno więc utożsamiać model collapse z tezą, że wszystko, co jest wielokrotnie syntetyzowane, musi stać się głupsze. Rekurencja może prowadzić zarówno do degeneracji, jak i do emergencji struktury – zależnie od mechanizmu transmisji oraz obecności zewnętrznego sprzężenia ze światem. To rozróżnienie stanowi kluczowy naukowy kontrargument wobec katastroficznych interpretacji pracy Shumailova. Model collapse nie jest nieuchronny w każdym systemie wykorzystującym dane syntetyczne. Matthias Gerstgrasser i współautorzy wykazali, że decydujący jest sposób zarządzania kolejnymi generacjami danych. Gdy stare dane rzeczywiste są zastępowane przez syntetyczne wyniki nowych modeli, degradacja rzeczywiście przyspiesza. Jeśli jednak pierwotne dane pozostają w zbiorze, a nowe treści są do nich akumulacyjnie dodawane, collapse można uniknąć; rezultat ten utrzymywał się niezależnie od rozmiaru modeli, ich architektury czy typu danych.

Problemem nie jest zatem „syntetyczność” jako metafizyczna skaza, lecz utrata zakotwiczenia w niezależnych danych pierwotnych. Dane syntetyczne mogą być wręcz niezwykle użyteczne. W biomedycynie stosuje się je m.in. do augmentacji rzadkich klas, testowania hipotez, ochrony prywatności czy symulowania scenariuszy, dla których pozyskanie realnych danych jest utrudnione. Przeglądy z 2024 i 2025 roku wskazują, że odpowiednio zaprojektowana syntetyczna augmentacja może niwelować niedobory danych, wspierać reprezentację grup marginalizowanych i umożliwiać analizę sytuacji kosztownych lub etycznie problematycznych. Wymaga to jednak zewnętrznej walidacji i kontroli nad tym, czy dane syntetyczne zachowują właściwości badanej populacji.

Innymi słowy: różnica przebiega między syntetycznym rozszerzaniem rzeczywistości a jej syntetycznym zastępowaniem. W pierwszym przypadku model wykorzystuje sztucznie wygenerowane przypadki, lecz nadal jest kalibrowany wobec świata. W drugim zaczyna mierzyć własne poprzednie przybliżenia i traktować je jako niezależny materiał empiryczny. To odpowiednik nauki, w której symulacje komputerowe przestałyby być konfrontowane z eksperymentem, a kolejne modele kalibrowano wyłącznie na wynikach poprzednich. Sam fakt, że matematyka pozostawałaby elegancka, nie gwarantowałby kontaktu z przyrodą.

Transparentność pochodzenia danych jako ochrona przed degradacją wiedzy

Najnowsze badania wskazują, że model collapse staje się odrębnym problemem inżynierskim, a nie fatalistycznym wyrokiem. W lipcu 2026 roku badacze opisali metodę ForTIFAI, wykorzystującą funkcje straty uwzględniające pewność modelu i obniżające wagę tokenów o wysokim poziomie pewności, które częściej odpowiadają typowym artefaktom generatywnym. W przeprowadzonych eksperymentach system tolerował ponad dwukrotnie większy udział danych syntetycznych przed wystąpieniem załamania. Wynik ten nie rozwiązuje problemu raz na zawsze, ale dowodzi, że strukturę pętli można projektować. Rekurencja jest parametrem systemu, a nie przeznaczeniem.

W tym kontekście pojęcie data provenance nabiera drugiego, głęboko technicznego znaczenia. Jeśli przyszedł model nie potrafi odróżnić dokumentu będącego wynikiem ludzkiej obserwacji od tekstu wygenerowanego na bazie innego modelu, traci informację o epistemicznej genealogii danych. Oznaczenie syntetycznego pochodzenia przestaje być zatem wyłącznie prawem czytelnika, a staje się narzędziem zarządzania zbiorami treningowymi.

Model powinien wiedzieć nie tylko, co znajduje się w zbiorze, ale przede wszystkim skąd dane pochodzą. W wymiarze ekonomicznym prawdopodobnie zwiększy to wartość danych pierwotnych. Shumailov i współautorzy zauważają, że zapisy rzeczywistych ludzkich interakcji mogą stać się szczególnie cenne w świecie nasyconym treściami syntetycznymi. Można zatem wyobrazić sobie nową kategorię aktywów informacyjnych: zweryfikowane zbiory human-origin data, archiwa sprzed masowego upowszechnienia generatywnej AI, kontrolowane korpusy eksperckie oraz dane pochodzące z pomiarów świata fizycznego. Paradoksalnie, im tańsza stanie się produkcja tekstu, tym droższa może być pewność, że dane nie są jedynie przetworzeniem wcześniejszej generacji treści. Prowadzi to do ekonomicznego problemu commons.

Internet był niezwykle wspólnym zasobem danych właśnie dlatego, że miliony ludzi publikowały teksty z powodów niezwiązanych z treningiem modeli: pisały artykuły, instrukcje, blogi czy dokumentację. Modele mogły dzięki temu wykorzystać powstałą wcześniej nadwyżkę kulturową. Jeśli jednak przestrzeń publiczna zostanie zalana tanimi treściami generatywnymi, jakość tego wspólnego zasobu zacznie spadać. Przypomina to tragedię wspólnego pastwiska: indywidualnie opłaca się zwiększać produkcję niskim kosztem, lecz zbiorowym efektem jest degradacja środowiska, z którego wszyscy korzystają. Elinor Ostrom nauczyła ekonomię, że wspólne zasoby nie muszą nieuchronnie ulec tragedii, o ile wspólnota ustanowi reguły monitorowania, granice dostępu i mechanizmy odpowiedzialności. Analogiczna lekcja dotyczy danych. Konieczne mogą okazać się standardy provenance, filtrowanie treści syntetycznych, repozytoria danych pierwotnych, audyty pochodzenia oraz świadome utrzymywanie różnorodności korpusów.

Epistemiczna rekursja i zagrożenie monokulturą wiedzy

Model collapse nie jest zatem wyłącznie problemem architektury sieci neuronowych, lecz kwestią governance infrastruktury poznawczej. Perspektywa ekologiczna wprowadza tu dodatkowy wymiar: odporność systemu. Ekosystem monokulturowy może być wydajny w stabilnych warunkach, pozostając jednocześnie skrajnie kruchym wobec nowego patogenu czy zmiany klimatu. Podobnie model, coraz silniej skoncentrowany wokół najbardziej typowych przypadków, może osiągać wysokie wyniki w przeciętnych zadaniach, tracąc jednak zdolność reagowania na anomalie. Tymczasem to właśnie anomalie bywają poznawczo najcenniejsze. W nauce odkrycie często zaczyna się od obserwacji niepasującej do dominującego modelu; w medycynie rzadki przypadek może wymagać zupełnie innej diagnozy, a w polityce mała grupa potrafi zasygnalizować problem niewidoczny w danych większościowych.

Ogony rozkładu są informacyjnym odpowiednikiem bioróżnorodności. Nie wiemy z góry, który rzadki przypadek okaże się kluczowy dla przyszłego problemu. Redukując różnorodność tylko dlatego, że ma niskie prawdopodobieństwo, optymalizujemy system pod świat podobny do przeszłości. Z teorii odporności systemów złożonych wiemy jednak, że nadmierna optymalizacja jednej funkcji może osłabić zdolność adaptacji do szoków. Efektywność i resilience nie są tym samym.

W tym punkcie w zaskakujący sposób spotykają się spostrzeżenia Naomi Baron i teoria model collapse. Baron obawia się, że człowiek, delegując proces pisania, przestanie wykonywać operacje poznawcze niezbędne do rozwijania własnego głosu. Materiał źródłowy określa tę sytuację jako utratę procesu odkrywania i przesunięcie od aktywnego tworzenia ku akceptowaniu

algorytmicznych propozycji. Literatura o model collapse ukazuje analogiczny problem na poziomie systemowym: model karmiony własnymi produktami może tracić kontakt z rzadkimi fragmentami rozkładu, które pierwotnie umożliwiły mu bogate modelowanie świata. W jednym przypadku zanika różnorodność poznawcza autora, w drugim – różnorodność danych. Choć nie są to te same procesy, posiadają wspólną strukturę: system zaczyna konsumować własne uproszczenia. Można zatem zdefiniować szersze zjawisko jako epistemic recursion. Zachodzi ono wtedy, gdy reprezentacja rzeczywistości staje się wejściem dla kolejnej reprezentacji, bez dostatecznego powrotu do samej rzeczywistości.

Dotyczy to modelu uczącego się na syntetycznych danych, dziennikarza streszczającego artykuły będące streszczeniami, naukowca cytującego przegląd zamiast wyników pierwotnych oraz użytkownika budującego poglądy wyłącznie na podsumowaniach generowanych przez systemy. Technologia AI czyni tę rekursję skalowalną, lecz nie jest ona jedynie technologiczną patologią – to odwieczny problem wiedzy pośredniej doprowadzony do skali przemysłowej. Przyszłość wartościowej sztucznej inteligencji może zatem paradoksalnie zależeć od utrzymywania kosztownych kanałów niealgorytmicznego kontaktu ze światem. Ktoś musi przeprowadzić eksperyment, odwiedzić miejsce wydarzeń, rozmawiać z człowiekiem, którego doświadczenia nie ma jeszcze w korpusie. Ktoś musi opisać nowy gatunek, sprawdzić archiwum, wykonać pomiar, napisać nieoczekiwane zdanie i stworzyć dzieło wykraczające poza najłatwiejsze do przewidzenia centrum rozkładu. Maszyny mogą przetwarzać te informacje z niewyobrażalną wydajnością, lecz nie powinny mylić własnych wcześniejszych przetworzeń z nieskończonym źródłem nowych faktów.

Ostatecznie model collapse okazuje się problemem znacznie szerszym niż jakość przyszłych chatbotów; jest pytaniem o warunki reprodukcji różnorodności wiedzy. Kultura zawsze była systemem kopiowania, modyfikacji i zapominania. Generatywna AI przyspiesza wszystkie trzy procesy. Może przechowywać, zestawiać i rekonstruować więcej niż pojedynczy człowiek, ale może też z szybkością przemysłową multiplikować przeciętność, jeżeli przeciętność stanie się głównym materiałem następnego cyklu uczenia. Dlatego wyzwaniem nie jest zatrzymanie danych syntetycznych – byłoby to niewykonalne i naukowo nieracjonalne. Trzeba natomiast zachować różnicę między symulacją a obserwacją, między kopią a źródłem, między wygenerowaną możliwością a wydarzeniem, które naprawdę zaszło. Zdrowy ekosystem AI potrzebuje syntetycznej wyobraźni, lecz również stałego dopływu danych spoza własnego kręgu.

Dochodzimy do finału wywodu.

Wartość pisania jako praxis i dobra wewnętrznego

Pytanie „kto to napisał?” nie daje się już rozstrzygnąć nazwiskiem umieszczonym pod tekstem. Stało się ono pytaniem o proces poznawczy, twórczość, gospodarkę, prawo i edukację, a także o odpowiedzialność, pochodzenie danych, architekturę platform oraz reprodukcję samej wiedzy. Pozostaje zatem postawić pytanie normatywne: które części pisania możemy bezpiecznie oddać maszynom, a których nie powinniśmy powierzać technologii nawet wtedy, gdy potrafi ona wykonać je szybciej i lepiej? Temu poświęcona będzie dalsza część wywodu: nie obrona człowieka przed technologią, lecz próba zbudowania dojrzałej architektury współistnienia, w której nie każdemu tekstowi potrzebny jest ludzki autor, ale w pewnych obszarach właśnie droga do zdania jest równie cenna jak samo zdanie. Dlaczego ludzkie autorstwo ma znaczenie i kiedy nie musi mieć znaczenia aż tak bardzo?

Czas wreszcie udzielić odpowiedzi, która na początku wywodu byłaby przedwczesna. Ludzkie autorstwo nie ma znaczenia dlatego, że człowiek zawsze pisze lepiej od maszyny – to twierdzenie jest empirycznie nie do utrzymania. Nie wynika ono również z przekonania, że każdy tekst powinien być rezultatem mozolnej pracy biologicznego umysłu; automatyzacja wielu form języka jest po prostu pożyteczna i racjonalna.

Ludzkie autorstwo ma szczególne znaczenie tam, gdzie proces dochodzenia do wypowiedzi stanowi część wartości tej wypowiedzi – gdzie słowa ustanawiają odpowiedzialność, świadczą o doświadczeniu, ćwiczą zdolność sądzenia albo reprezentują osobę wobec drugiego człowieka. W innych sytuacjach wartość mieści się przede wszystkim w rezultacie i wtedy pytanie, czy każde zdanie osobiście sformułował człowiek, może być niemal obojętne. Naomi S. Baron prowadzi swój wywód wyraźnie w stronę obrony ludzkiego pisania przed "pokusą wydajności". W materiale źródłowym jej końcowa argumentacja skupia się na scorecard, świadomym decydowaniu o tym, które działania delegować AI oraz na „strefach wolnych od AI”, w których człowiek zachowuje samodzielne pisanie właśnie dlatego, że sama czynność ma wartość poznawczą i osobotwórczą. W źródle powraca także metafora „światlnego miecza”: czytanie i pisanie mają wspierać myśl, wymianę oraz działanie, a całkowity outsourcing stwarza ryzyko osłabienia ludzkiego sprawstwa.

Należy jednak oddzielić tę normatywną intuicję od mocniejszych twierdzeń biologicznych. Współczesne badania nie uzasadniają prostego prawa mówiącego, że każde użycie generatywnej AI prowadzi do „atrofii mózgu”. Obraz jest znacznie subtelniejszy: skutki zależą od zadania, kolejności pracy, poziomu wiedzy, sposobu wykorzystania narzędzia oraz tego, jaka funkcja poznawcza została delegowana. Dobrym punktem wyjścia do syntezy jest klasyczne rozróżnienie Arystotelesa pomiędzy poiesis i praxis. Poiesis jest działaniem skierowanym ku zewnętrznemu rezultatowi:

wytwarzamy coś, co po zakończeniu czynności istnieje jako produkt. Praxis ma wartość również wewnątrz samego wykonywania działania; uczestnictwo w praktyce kształtuje sprawcę. Pisanie posiada obie właściwości jednocześnie. Raport jest produktem, ale pisanie raportu może być także ćwiczeniem rozpoznawania związków przyczynowych. Esej jest dokumentem, lecz proces jego tworzenia może kształcić zdolność argumentowania. Pamiętnik jest zbiorem zdań, jednak jego pisanie pozwala porządkować osobiste doświadczenie.

Jeżeli potraktujemy wszystkie te czynności wyłącznie jako poiesis, AI jawi się jako oczywista technologia optymalizacyjna: skoro potrafi wytworzyć produkt szybciej, należy jej go powierzyć. Jeżeli jednak część pisania jest praxis, rachunek się zmienia. Usunięcie wysiłku może oznaczać usunięcie właśnie tej części czynności, dzięki której rozwijał się wykonawca. Alasdair MacIntyre dostarcza tu jeszcze precyzyjniejszego narzędzia.

Rozróżniał on dobra zewnętrzne praktyki, takie jak pieniądze, prestiż czy pozycja, od dóbr wewnętrznych, które można osiągnąć wyłącznie poprzez uczestnictwo w samej praktyce i rozwijanie właściwych jej kompetencji. W odniesieniu do pisania zewnętrznym dobrem może być liczba opublikowanych raportów, ocena, wynagrodzenie, zasięg albo terminowe dostarczenie dokumentu. Dobrem wewnętrznym jest natomiast coraz lepsze rozumienie argumentacji, subtelności języka, rytmu narracji, kryteriów dowodu i własnego sposobu myślenia. Generatywna AI jest niezwykle skuteczna w zwiększaniu podaży dóbr zewnętrznych. Może jednak pozostawić nierozstrzygnięte pytanie, co stało się z dobrami wewnętrznymi praktyki. To tłumaczy, dlaczego identyczne użycie AI może być rozsądne w jednym kontekście i destrukcyjne w drugim.

Kryterium autonomii w outsourcingu poznawczym

Gdy doświadczony prawnik zleca systemowi sformatowanie rutynowej korespondencji, deleguje czynność, która nie stanowi już kluczowego mechanizmu jego rozwoju zawodowego. Jeśli jednak student pierwszego roku powierza modelowi zbudowanie całej argumentacji prawnej, może oddać właśnie to, czego w procesie studiów miał się nauczyć. Nie istnieje zatem uniwersalna etyka funkcji; istnieje jedynie relacja konkretnej funkcji do celu danej osoby lub instytucji.

Ta intuicja współgra z podejściem UNESCO do generatywnej AI w edukacji. Organizacja odrzuca zarówno całkowity zakaz, jak i dogmat maksymalnej automatyzacji. Wytyczne UNESCO opierają się na podejściu human-centred: technologia ma być projektowana i wykorzystywana tak, by jej zastosowanie pozostawało etyczne, bezpieczne, znaczące i podporządkowane rozwojowi ludzkich zdolności. Nie chodzi więc o zachowanie każdej dawnej czynności tylko dlatego, że kiedyś

wykonywał ją człowiek, lecz o rozpoznanie, która z tych zdolności musi pozostać ludzka, gdyż jest warunkiem autonomii.

Filozofia "extended mind" Andy'ego Clarka i Davida Chalmersa pozwala uniknąć drugiego ekstremum, jakim jest romantycznej wizji umysłu czystego, nieskażonego technologią. Ich argument przypomina, że proces poznawczy może wykorzystywać elementy środowiska jako funkcjonalne składniki systemu poznawczego. Notatnik, zapis czy mapa działają w ten sposób jak rozszerzenia pamięci i rozumowania. Z tej perspektywy korzystanie z AI nie jest samo w sobie "utrata człowieczeństwa" – człowiek od zawsze myśli za pomocą narzędzi. Pismo było rewolucyjnym outsourcingiem pamięci, druk rozszerzył dostęp do cudzych umysłów, kalkulator przejął fragment rachunków, a internet stał się zewnętrznym magazynem informacji nieosiągalnych dla pojedynczej pamięci.

Różnica tkwi jednak w stopniu i rodzaju tego outsourcingu. Notatnik przechowuje to, co człowiek wcześniej zdecydował się zapisać; kalkulator wykonuje operację ściśle określoną przez użytkownika. Model generatywny natomiast może samodzielnie zaproponować problem, kategorię, metaforę, przesłankę, kontrargument, a nawet wniosek.

Od automatyzacji do rozszerzonego sprawstwa i kontroli

AI nie tylko rozszerza magazyn poznawczy, ale coraz częściej uczestniczy w wyborze samej trajektorii myślenia. Dlatego potrzebujemy teorii nie tyle rozszerzonego umysłu, co rozszerzonego sprawstwa: musimy ustalić, jak rozdzielić funkcje między człowieka a system, aby wzrost zdolności zespołu nie prowadził do zaniku umiejętności samodzielnego rozpoznawania celu. Badania empiryczne pokazują, że kwestii tej nie da się rozstrzygnąć prostym bilansem zysków i strat. Badanie 465 przyszłych nauczycieli z 2025 roku wykazało pozytywne korelacje między pewnymi formami użycia GenAI a osiągnięciami akademickimi, gdzie kluczową rolę odgrywały współdzielone metapoznanie oraz cognitive offloading. To istotny kontrargument dla tezy, że odciążenie poznawcze zawsze degradowe proces uczenia; może ono bowiem zwalniać zasoby pamięci roboczej i pozwalać na przesunięcie wysiłku ku zadaniom wyższego rzędu. Jednocześnie seria czterech eksperymentów z 2025 roku, obejmująca 3562 osoby, wykazała, że choć współpraca z GenAI poprawia bieżące wykonanie zadań, korzyść ta nie przenosi się automatycznie na późniejsze prace wykonywane samodzielnie. Badacze odnotowali również wpływ AI na poczucie kontroli, motywację wewnętrzną i nudę. Te dwa zestawy wyników nie są sprzeczne.

Wskazują one dokładnie to, co powinno stać się fundamentem dojrzałej polityki AI: delegacja funkcji może człowieka wspierać lub osłabiać, zależnie od tego, co zostaje odciążone i jak wykorzystany zostanie odzyskany zasób poznawczy. Teoria autodeterminacji Edwarda Deciego i Richarda Ryana wprowadza tutaj trzy kluczowe potrzeby psychologiczne: autonomię, kompetencję i relacyjność. Technologia pozwalająca wykonać zadanie wcześniej poza zasięgiem użytkownika może zwiększać poczucie kompetencji. Jednak system permanentnie dokonujący za niego zasadniczych wyborów może to poczucie autonomii drastycznie zmniejszyć. Z kolei interaktywny feedback wspiera uczenie, lecz zastąpienie relacji z nauczycielem wyłącznie syntetyczną interakcją osłabia społeczny wymiar rozwoju. Badania nad AI-mediated collaboration sięgają po teorię autodeterminacji właśnie dlatego, że sama „wydajność” jest zbyt wąskim wskaźnikiem jakości współpracy. Jeszcze wcześniejszą przestrożę sformułowała Lisanne Bainbridge w klasycznych "Ironies of Automation" z 1983 roku.

Automatyzacja może pozostawić człowiekowi odpowiedzialność za sytuacje wyjątkowe, jednocześnie odbierając mu regularną praktykę niezbędną do utrzymania kompetencji wymaganych właśnie wtedy, gdy automatyka zawiedzie. Paradoks ten niemal idealnie pasuje do generatywnego pisania. Jeśli AI przez lata przygotowuje raporty i proponuje interpretacje, człowiek może formalnie pozostać „ostatecznym arbitrem”, lecz w krytycznym momencie arbitraż wymaga kompetencji zbudowanych poprzez wykonywanie czynności, którą wcześniej zautomatyzowano.

W tym kontekście hasło *human in the loop* jest niewystarczające. Człowiek może znajdować się w pętli jako rzeczywisty decydent albo jako ceremonialny zatwierdzający. Może posiadać prawo do sprzeciwu albo tylko obowiązek podpisania. Może rozumieć system albo widzieć jedynie rezultat. Może mieć godzinę na kontrolę dokumentu albo trzydzieści sekund.

Sama biologiczna obecność nie gwarantuje znaczącej kontroli. Ben Shneiderman w ramach Human-Centered AI proponuje zerwanie z intuicją, że wysoka automatyzacja musi oznaczać niską kontrolę człowieka. Jego dwuwymiarowy model zakłada projektowanie systemów łączących wysoką automatyzację z wysokim poziomem ludzkiego sterowania. To jedna z najważniejszych idei niezbędnych do rozstrzygnięcia naszego problemu.

Alternatywą nie musi być wybór między samotnym człowiekiem z piórem a całkowicie autonomicznym Grammatizatorem. Możemy projektować technologie, które przejmują ogromną część pracy obliczeniowej i językowej, zachowując jednocześnie wyraźne punkty ludzkiej decyzji, możliwość ingerencji, przewidywalność działania i odpowiedzialność. Shneiderman opisuje taką logikę jako "supertools": systemy zwiększające zdolności człowieka bez konieczności udawania, że technologia powinna stać się autonomicznym substytutem osoby.

Techniczne i etyczne ramy odpowiedzialnego delegowania

Z perspektywy inżynierii bezpieczeństwa model ten można uzupełnić o trzy pojęcia: observability, controllability i recoverability. Autor powinien mieć wgląd w to, gdzie AI ingerowała w tekst; musi móc zmienić lub odrzucić propozycję systemu, a w przypadku błędu – dysponować sposobem powrotu do stanu wiarygodnego. Zasady te, wywodzące się z systemów technicznych, okazują się zadziwiająco trafne w kulturze pisania. Dokument generatywny pozbawiony historii wersji i wiedzy o źródłach cechuje niska obserwowalność. Proces, w którym człowiek musi zatwierdzić tysiące wyników pod presją czasu, oferuje niską rzeczywistą sterowalność. Z kolei organizacja, która nie przechowuje pierwotnych danych i nie potrafi odtworzyć procesu powstawania raportu, wykazuje niską zdolność do odzyskania kontroli.

NIST w profilu Generative AI do AI Risk Management Framework traktuje ryzyko jako problem całego cyklu życia systemu, a nie pojedynczego promptu. Zarządzanie powinno obejmować projektowanie, wdrożenie, pomiar oraz governance, uwzględniając zagrożenia takie jak konfabulacje, naruszenie integralności informacji, prywatność, uprzedzenia czy problemy z pochodzeniem treści. W kontekście autorstwa oznacza to przejście od pytania „czy wolno mi użyć AI?” do pytania: „jak zaprojektować procedurę korzystania z AI, aby wiedzieć, które elementy wymagają ludzkiej kontroli i gdzie lokuje się odpowiedzialność?”.

Scorecard Baron można przekształcić z osobistej porady w ogólniejszą teorię delegowania. Materiał źródłowy proponuje świadome ustalenie, które zadania rutynowe można powierzyć AI, a które głębokie formy wypowiedzi należy zachować dla siebie. Naukowo dojrzała wersja tej karty nie pytałaby jedynie o procent tekstu wygenerowanego przez system. Oceniałaby co najmniej pięć wymiarów jednocześnie: stawkę epistemiczną, funkcję edukacyjną, znaczenie relacyjne, odwracalność błędu oraz możliwość audytu pochodzenia.

Nie trzeba zamykać tych kryteriów w administracyjnej tabeli, aby dostrzec rządzącą nimi zasadę. Im większa szkoda możliwa wskutek błędu, im bardziej wykonanie czynności jest celem uczenia, im silniej tekst ma świadczyć o osobistej intencji, im trudniej cofnąć jego konsekwencje i im słabsza możliwość późniejszej rekonstrukcji procesu – tym większy powinien być udział człowieka. Reguła ta wyjaśnia, dlaczego automatyzacja prognozy pogody i automatyzacja listu pożegnального nie są etycznie równoważne, choć oba dokumenty składają się ze słów. Pierwszy jest przede wszystkim nośnikiem informacji instrumentalnej; drugi zaś aktem relacyjnym. Podobnie streszczenie tysiąca rutynowych dokumentów może być idealnym zastosowaniem AI, podczas gdy napisanie przez model osobistego zeznania świadka podważałoby samą naturę świadectwa.

Wartość ekologiczna autorstwa jako źródło danych pierwotnych

Filozofia języka wskazuje, że różne teksty pełnią funkcję różnych aktów mowy. Habermas dodałby tutaj, że komunikacja może być zorientowana albo na wzajemne zrozumienie, albo na strategiczne wywołanie pożądanego zachowania. W epoce generatywnej różnica ta nabiera szczególnego znaczenia. Jeśli firma automatycznie generuje "osobiste" przeprosiny tylko po to, by poprawić wskaźnik satysfakcji klienta, forma komunikacji pozostaje interpersonalna, lecz logika jej wytwarzania staje się strategiczna. Nie musi to oznaczać nieuczciwości, ale oznacza, że odbiorca ma uzasadniony interes w poznaniu warunków powstania komunikatu.

Prowenancja (provenance) nie jest jedynie dodatkiem do autorstwa, lecz integralną częścią architektury zaufania. W części XIII pokazaliśmy, że prawo europejskie zaczęło już regulować niektóre aspekty syntetycznego pochodzenia treści. Normatywna zasada jest jednak szersza: im bardziej odbiorca racjonalnie zmieniłby interpretację tekstu po poznaniu sposobu jego powstania, tym silniejszy staje się argument za ujawnieniem tej informacji. Nie musimy oznaczać każdej autokorekty, powinniśmy jednak zachować szczególną ostrożność tam, gdzie syntetyczny tekst korzysta z kredytu zaufania należnego ludzkiemu doświadczeniu, eksperckości albo osobistej intencji. Zasada ta chroni również przed tym, co Miranda Fricker nazwałaby problemami epistemicznej sprawiedliwości.

Jeżeli systemy syntetyczne potrafią produkować ogromną liczbę płynnych wypowiedzi, niebezpieczeństwem nie jest wyłącznie ich błąd, lecz możliwość zagłuszenia mniej licznych i trudniej formułowanych, ale pierwotnych świadectw ludzi. W społeczeństwie informacyjnym sam dostęp do mównicy nie wystarcza, jeśli syntetyczna podaż komunikatów może uczynić ludzkie doświadczenie praktycznie niewidocznym. Obrona human authorship jest więc w pewnych kontekstach także obroną epistemicznej reprezentacji osób – szczególnie tych, których doświadczenia znajdują się w ogonach statystycznych rozkładów. Prowadzi nas to z powrotem do model collapse.

Okazało się, że problem syntetycznych danych nie polega na ich ontologicznej nieczystości, lecz na utracie kontaktu z niezależnym materiałem pierwotnym. Analogicznie kultura nie potrzebuje "czystego ludzkiego internetu", lecz stałego dopływu doświadczeń, obserwacji, badań, relacji i form ekspresji, które nie są jedynie transformacjami wcześniejszych generacji treści. W przeciwnym razie system zacznie przypominać bibliotekę, w której powstaje coraz więcej książek, ale coraz mniej osób wychodzi na zewnątrz, by sprawdzić, czy świat nadal wygląda tak, jak opisują go księgi.

Ludzkie autorstwo posiada wartość ekologiczną; jest jednym ze źródeł nowych danych pierwotnych kultury. Człowiek spotyka coś, czego wcześniej w korpusie nie było, i próbuje nadać temu język: doświadcza wojny, migracji, miłości, choroby, odkrycia naukowego, konfliktu pokoleń, nowej technologii albo lokalnej anomalii. Model może mu pomóc to opisać.

Jeśli jednak system komunikacyjny przestaje zachęcać ludzi do pierwotnego opisywania świata, a premiuje jedynie generowanie kolejnych wariantów już dostępnych reprezentacji, kultura traci dopływ epistemicznej "świeżej wody". Nie oznacza to oczywiście, że tylko tekst napisany bez AI jest oryginalny. Człowiek może przy pomocy systemu zbudować koncepcję całkowicie nową, tak jak naukowiec za pomocą mikroskopu odkrywa zjawisko, którego nie mógłby dostrzec gołym okiem. Kryterium nie jest zatem technologiczna czystość procesu.

Odpowiedzialne sprawstwo jako miara wartości twórczości

Kluczowy jest tutaj kierunek zależności: czy narzędzie pomaga człowiekowi więcej zobaczyć i pomyśleć, czy przede wszystkim pozwala mu unikać patrzenia i myślenia.

W tym miejscu potrzebujemy perspektywy ekonomii Amartyi Sena. Capability Approach pozwala oceniać rozwój nie przez pryzmat ilości dóbr czy wydajności, lecz poprzez realne zdolności ludzi do prowadzenia życia, które mają powody cenić. Analogicznie, wartość generatywnej AI nie powinna być mierzona wyłącznie liczbą godzin oszczędzonych na pisaniu. Powinniśmy pytać, czy technologia poszerza capabilities: zdolność do uczenia się, uczestnictwa w debacie, tworzenia, rozumienia dokumentów, komunikowania się ponad barierami językowymi i podejmowania samodzielnych decyzji. Jeśli osoba z powodu bariery językowej wcześniej nie mogła publikować swoich odkryć, AI zwiększa jej sprawczość. Jeśli jednak ta sama osoba po latach nie potrafi bez narzędzia ocenić argumentu, część capability zostaje utracona.

Z tego powodu „strefy wolne od AI” Baron warto rozumieć nie jako przejaw technologicznego purytanizmu, lecz jako odpowiednik treningu bez protezy. Sportowiec korzysta z zaawansowanych pomiarów i sprzętu, ale pewne zdolności musi rozwijać własnym wysiłkiem. Pilot uczy się systemów automatycznych, lecz musi zachować kompetencje niezbędne w sytuacjach awaryjnych. Muzyk może używać technologii nagraniowej, jednak jeśli celem jest nauka gry na instrumencie, pewnych ruchów nie da się zlecić komputerowi. Analogicznie szkoła, uniwersytet, redakcja czy pojedynczy autor mogą racjonalnie ustanawiać okresy lub zadania, w których generacja jest wyłączona –

ponieważ to ćwiczenie człowieka, a nie produkcja tekstu, jest w danym momencie właściwym celem. Takie strefy nie muszą być trwałe.

Student może najpierw sformułować argument samodzielnie, następnie skonfrontować go z AI, by na koniec obronić własne decyzje. Pisarz może prowadzić pierwszą fazę konceptualizacji bez generatora, wykorzystując model dopiero jako krytyka. Naukowiec może sam zdefiniować hipotezę, a system użyć do przeszukiwania alternatywnych interpretacji. Kluczowa staje się sekwencja poznawcza. Jeżeli narzędzie pojawia się zbyt wcześnie, może zastąpić proces generowania hipotezy; jeżeli pojawi się później – zwiększa możliwości jej testowania.

Warto zatem odwrócić często przywoływaną maksymę IBM: „maszyny powinny pracować, ludzie powinni myśleć”. W XXI wieku ten podział jest już zbyt uproszczony, gdyż maszyny wykonują operacje funkcjonalnie przypominające to, co nazywaliśmy myśleniem. Bardziej użyteczna wersja brzmiałaby: maszyny powinny przejmować pracę tam, gdzie ich przewaga zwiększa ludzką zdolność sądenia; nie powinny jednak systemowo usuwać tych doświadczeń, dzięki którym zdolność sądenia w ogóle powstaje.

Tak można zinterpretować również „światlny miecz” Baron. Metafora ta nie powinna prowadzić do fetyszyzowania pióra czy klawiatury. Mieczem nie jest mechaniczne stawianie znaków, lecz zdolność wyodrębnienia myśli, nadania jej formy, skonfrontowania z oporem, poprawienia błędu i wzięcia za nią odpowiedzialności. Czasem człowiek robi to sam; czasem z pomocą książki, redaktora, kalkulatora czy modelu generatywnego. Zagrożenie pojawia się wtedy, gdy przestajemy rozpoznawać, które cięcia wykonujemy sami, a które za nas narzędzie. Ostatecznym kryterium nie jest czystość autorstwa, lecz odpowiedzialne sprawstwo.

Człowiek może nie napisać każdego słowa, a mimo to pozostać rzeczywistym autorem intelektualnym tekstu, jeśli określił problem, rozumie argument, zweryfikował informacje i dokonał wyborów konstytutywnych. Może z kolei fizycznie wpisać wszystkie znaki i być autorem słabym lub bezmyślnym – liczba naciśnięć klawiszy nigdy nie była pełną miarą autorstwa. Prawo autorskie intuicyjnie zmierza w tym kierunku, chroniąc ludzki twórczy wybór, selekcję i aranżację, a nie samą techniczną czynność generowania znaków.

Edukacja powinna analogicznie oceniać operacje poznawcze, a nie wyłącznie powierzchnię eseju. Dziennikarstwo powinno chronić zdobywanie informacji pierwotnej i odpowiedzialność redakcyjną, zamiast fetyszyzować ręczne przepisywanie depesz. Organizacja zaś winna mierzyć nie liczbę dokumentów wyprodukowanych przez AI, lecz jakość decyzji podejmowanych na ich podstawie. Dojrzałą architekturę współistnienia człowieka i maszyny możemy zdefiniować poprzez zasadę asymetrycznej delegacji. Delegujmy łatwo operacje odwracalne, rutynowe i niezwiązane z

celem uczenia. Delegujemy ostrożnie czynności epistemicznie wysokiego ryzyka. Zachowajmy silne ludzkie sprawstwo tam, gdzie tekst jest świadectwem, obietnicą, decyzją moralną, aktem odpowiedzialności lub dowodem kompetencji, którą człowiek powinien rzeczywiście posiadać.

Wartość pisania tkwi w procesie kształtowania ludzkiego sprawstwa

Nie jest to granica raz na zawsze dana przez naturę; będzie ona ewoluować wraz z technologią i instytucjami. Funkcja, którą dziś musimy ściśle kontrolować, jutro może stać się banalnym elementem infrastruktury. Jedna zasada pozostanie jednak trwała: odpowiedzialność nie powinna być szersza niż realne sprawstwo, a realne sprawstwo nie powinno być mniejsze niż kompetencja potrzebna do tej odpowiedzialności. Jeśli system podejmuje większość decyzji, a człowiek pozostaje tylko nazwiskiem do podpisu, otrzymujemy moralną fikcję autorstwa.

W tym świetle pytanie „kto to napisał?” okazuje się znacznie bardziej wymagające niż zwykłe ustalenie źródła znaków. W rzeczywistości pyta ono: kto rozpoznał problem? Kto doświadczył świata, o którym mowa? Kto wybrał przesłanki? Kto dostrzegł możliwość błędu? Kto zmienił zdanie pod wpływem argumentu? Kto jest gotów bronić wniosku? Kto ponosi konsekwencje? I wreszcie: kto dzięki napisaniu tego tekstu stał się odrobinę inną osobą?

Maszyna może uczestniczyć w odpowiedzi na niemal każde z tych pytań, lecz nie oznacza to, że wszystkie stają się przez to równoważne pytaniu o wydajność. W ten sposób rozstrzyga się spór rozpoczęty przez Dahla. Najbardziej przewrotna wersja jego Grammatizatora nie byłaby maszyną odbierającą ludziom zawód pisarza, lecz narzędziem tak wygodnym i powszechnym, że społeczeństwo przestałoby zauważać różnicę pomiędzy posiadaniem tekstu a przejściem przez myśl, która do niego prowadziła. Człowiek nadal widniałby jako autor, podpisywał dokumenty, publikował książki czy zdawał prace, lecz coraz rzadziej byłby miejscem, w którym rodziła się ich zasadnicza treść.

Tego należy unikać – nie po to, by zatrzymać rozwój generatywnej sztucznej inteligencji, ale by uczynić go naprawdę użytecznym. Technologia jest najbardziej ludzka nie wtedy, gdy perfekcyjnie imituje człowieka, lecz gdy zwiększa zakres jego rozumnej sprawczości. Human-Centered AI proponuje właśnie zerwanie z równaniem „więcej automatyzacji = mniej kontroli”. Możliwe jest bowiem projektowanie systemów jednocześnie wysoce zautomatyzowanych i podporządkowanych ludzkiej kontroli, samoskuteczności, odpowiedzialności oraz kreatywności.

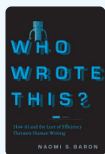
Odpowiedź brzmi paradoksalnie: nie każdy tekst potrzebuje ludzkiego autora w takim samym stopniu, ale społeczeństwo potrzebuje ludzi zdolnych do autorstwa. Potrzebuje ich nawet w świecie, w którym większość tekstów użytkowych będą produkować maszyny. Potrzebujemy osób umiejących rozpoznać nowe fakty, sformułować problem, zakwestionować dominujący wzorzec, opisać rzadkie doświadczenie czy nie zgodzić się z propozycją systemu – ludzi gotowych przyjąć odpowiedzialność i powiedzieć coś tylko dlatego, że uznali to za prawdziwe lub ważne.

Tekst bowiem nie zawsze jest jedynie przesyłką, którą trzeba dostarczyć. Czasami jest śladem drogi, którą ktoś rzeczywiście przeszedł. I w tych przypadkach droga do zdania staje się integralną częścią znaczenia samego zdania.

W świecie nasyconym syntetycznym nadmiarem możemy ulec złudzeniu, że posiadanie tekstu jest tożsame z przejściem przez myśl, która do niego prowadziła. Jednak istnieją wypowiedzi, które są czymś więcej niż tylko sprawną przesyłką informacji – są świadectwem drogi, którą ktoś naprawdę przeszedł. W takich przypadkach to właśnie trud i opór towarzyszący dochodzeniu do zdania stają się integralną częścią jego znaczenia.

Inspiracje

[1]



"Who Wrote This" Naomi S Baron

ISBN: 9781503643574

Naomi S. Baron (ur. 27 września 1946 r.) jest wybitną językoznawczynią i emerytowaną profesorką językoznawstwa na Uniwersytecie Amerykańskim w Waszyngtonie. Uzyskała tytuł licencjata (B.A.) na Uniwersytecie Brandeisa, a doktorat (Ph.D.) na Uniwersytecie Stanforda. W trakcie swojej znakomitej kariery akademickiej zajmowała stanowiska w takich instytucjach jak Uniwersytet Browna, Rhode Island School of Design, Uniwersytet Emory'ego i Uniwersytet Southwestern. Była stypendystką Fundacji Guggenheima i Fundacji Fulbrighta, Baron cieszy się międzynarodowym uznaniem za badania nad powiązaniem języka, technologii i umiejętności czytania i pisania. W swojej pracy szeroko bada, jak media cyfrowe, komunikacja mobilna i sztuczna inteligencja zmieniają ludzkie procesy czytania, pisania i poznawcze. Jest autorką licznych książek, w tym „Always On”, „Words Onscreen”, „How We Read Now” i „Who Wrote This?”, wnosząc znaczący wkład w zrozumienie ewolucji językowej w erze cyfrowej.

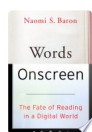
Naomi S. Baron (born September 27, 1946) is a prominent linguist and Professor Emerita of Linguistics at American University in Washington, D.C. She earned her B.A. from Brandeis University and her Ph.D. from Stanford University. Throughout her distinguished academic career, she has held positions at institutions including Brown University, the Rhode Island School of Design, Emory University, and Southwestern University. A former Guggenheim Fellow and Fulbright Fellow, Baron is internationally recognized for her research on the intersection of language, technology, and literacy. Her work extensively explores how digital media, mobile communication, and artificial intelligence reshape human reading, writing, and cognitive processes. She has authored numerous books, including "Always On," "Words Onscreen," "How We Read Now," and "Who Wrote This?," contributing significantly to the understanding of linguistic evolution in the digital age.

American University

Literatura uzupełniająca

Książki

Autorzy przywoływani w artykule:



[1] "Words Onscreen"

Naomi S. Baron

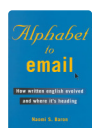
2015 | Oxford University Press, USA | ISBN: 9780199315765



[2] "How We Read Now"

Naomi S. Baron

2021 | Oxford University Press | ISBN: 9780190084103



[3] "Alphabet to Email"

Naomi S. Baron

2000 | Psychology Press | ISBN: 9780415186858



[4] "Reader Bot"

Naomi S. Baron

2026 | Stanford University Press | ISBN: 9781503644885



[5] "Speech, Writing, and Sign"

Naomi S. Baron

1981

Literatura pokrewna (Google Scholar):

[6] "A Recent Entrance to Paradise"

Stephen Thaler



[7] "The mathematical theory of communication"

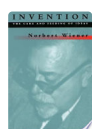
Claude Shannon

[8] "Paley–Wiener theorem"

Norbert Wiener

[9] "Wiener equation"

Norbert Wiener



[10] "Invention: The Care and Feeding of Ideas"

Norbert Wiener

MIT Press | ISBN: 9780262731119



[11] "The Divided West"

Jürgen Habermas

2014 | John Wiley & Sons | ISBN: 9780745694580



[12] "The Past as Future"

Jürgen Habermas

U of Nebraska Press | ISBN: 9780803272668

Philosophy
and
Social Sciences

[13] "Protestbewegung und Hochschulreform"

Jürgen Habermas

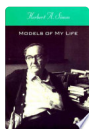
Sublime Verlag



[14] "Pre-Suasion"

Robert Cialdini

2016 | Random House | ISBN: 9781473519343



[15] "Models of my life"

Herbert Simon

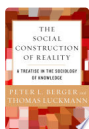
MIT Press | ISBN: 9780262326582



[16] "Human Problem Solving"

Herbert Simon

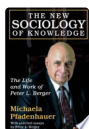
Prentice Hall



[17] "The Social Construction of Reality"

Peter L. Berger, Thomas Luckmann

2011 | Open Road Media | ISBN: 9781453215463



[18] "The New Sociology of Knowledge"

Michaela Pfadenhauer, Peter L. Berger

2013 | Transaction Publishers | ISBN: 9781412849890

[19] "Social Mobility and Personal Identity"

Thomas Luckmann, Peter L. Berger

1964

[20] "Unemployment Relief in Arizona from October 1, 1932, Through December 31, 1936"

Arthur Theodore Bartel, D. W. Albert, Elzer Des Jardines Tetreau, James Frank Breazeale, James Greenlief Brown, Karl Harris, Peter L. Berger, Robert Lavern Matlock, Walter Stanley Cunningham, William Thomas McGeorge, Charles A. Hobart, James R. Kennedy, M. F. Wharton, Robert Harry Hilgeman, Rubert Burley Streets, Thomas Luckmann, Allen Fisher Kinnison

1934



[21] "Beobachtungen der Moderne"

Niklas Luhmann

Springer-Verlag | ISBN: 9783322936172



[22] "Das Erziehungssystem der Gesellschaft"

Niklas Luhmann

ISBN: 9783518291931



[23] "Das Recht der Gesellschaft"

Niklas Luhmann

ISBN: 9783518287835



[24] "Reassembling Scholarly Communications: Histories, Infrastructures, and Global Politics of Open Access"

Bruno Latour

MIT Press | ISBN: 9780262536240



[25] "After Lockdown: A Metamorphosis"

Bruno Latour

Polity | ISBN: 9781509550029

[26] "Der Berliner Schlüssel. Erkundungen eines Liebhabers der Wissenschaften."

Bruno Latour

ISBN: 9783050028347

[27] "From Cliché to Archetype"

Marshall McLuhan



[28] "A View from the Roof"

Dave Carley

1997



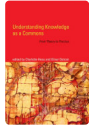
[29] "Analog Device-Level Layout Automation"

John M. Cohn, David J. Garrod, Rob A. Rutenbar, Rick Carley
2012 | Springer Science & Business Media | ISBN: 9781461527565



[30] "Rules, Games, and Common-pool Resources"

Elinor Ostrom
University of Michigan Press | ISBN: 9780472065462

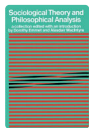


[31] "Understanding Knowledge as a Commons"

Elinor Ostrom
MIT Press | ISBN: 9780262516037

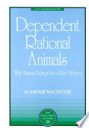
[32] "Institutions, Incentives, and Irrigation in Nepal"

Elinor Ostrom



[33] "Sociological theory and philosophical analysis"

Alasdair Macintyre
Palgrave



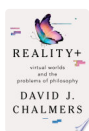
[34] "Dependent Rational Animals: Why Human Beings Need the Virtues"

Alasdair Macintyre
Open Court Publishing | ISBN: 9780812694529



[35] "Against the self-images of the age: essays on ideology and philosophy"

Alasdair Macintyre

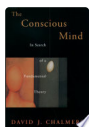


[36] "Reality+"

David Chalmers
W. W. Norton & Company | ISBN: 9780393635812

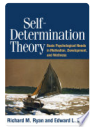
[37] "David Chalmers: How do you explain consciousness?"

David Chalmers



[38] "The Conscious Mind"

David Chalmers
Oxford University Press | ISBN: 9780199839353



[39] "Self-Determination Theory: Basic Psychological Needs in Motivation, Development, and Wellness"

Edward Deci

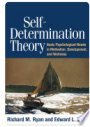
Guilford Publications | ISBN: 9781462528769



[40] "Intrinsic Motivation"

Edward Deci

2012 | Springer Science & Business Media | ISBN: 9781461344469



[41] "self-determination theory"

Edward Deci

2017 | Guilford Publications | ISBN: 9781462528769

[42] "Self-Determination Theory in Work Organizations"

Edward L Deci, Anja H. Olafsen, Richard Ryan

2018

[43] "Humanizing Automation"

Ben Shneiderman



[44] "Encounters with HCI Pioneers: A Personal History and Photo Journal"

Ben Shneiderman

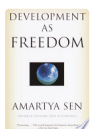
Springer Nature | ISBN: 9783031022241



[45] "Epistemic Injustice: Power and the Ethics of Knowing"

Miranda Fricker

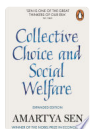
2007 | Clarendon Press | ISBN: 9780198237907



[46] "Development as Freedom"

Amartya Sen

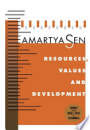
2011 | Vintage | ISBN: 9780307874290



[47] "Collective Choice and Social Welfare"

Amartya Sen

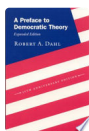
2017 | Penguin UK | ISBN: 9780241244609



[48] "Resources, Values and Development"

Amartya Kumar Sen, Amartya Sen

1997 | Harvard University Press | ISBN: 9780674765269



[49] "A Preface to Democratic Theory, Expanded Edition"

Robert Dahl

University of Chicago Press | ISBN: 9780226134345



[50] "Lady Luck; the theory of probability."

Warren Weaver

Courier Corporation | ISBN: 9780486150918



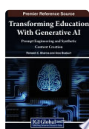
[51] "Social Impact Assessment And Monitoring"

Michael J Carley, Eduardo Bustelo

1984 | Westview Press

[52] "Interacting with eHealth: towards grand challenges for HCI"

Ben Shneiderman



[53] "Transforming Education With Generative AI: Prompt Engineering and Synthetic Content Creation"

Sharma, Ramesh C., Bozkurt, Aras

2024 | IGI Global | ISBN: 9798369313527



[54] "A Select Library of the Nicene and Post-Nicene Fathers of the Christian Church: Saint Augustin: Anti-Pelagian writings"

1887

Artykuły naukowe

Autorzy przywoływani:

[1] "Cycles of insanity and creativity within contemplative neural systems"

Stephen L. Thaler

2016 | Medical Hypotheses | DOI: 10.1016/j.mehy.2016.07.010

[Google Scholar](#)

[2] "Predicting ultra-hard binary compounds via cascaded auto- and hetero-associative neural networks"

Stephen L. Thaler

1998 | Journal of Alloys and Compounds | DOI: 10.1016/s0925-8388(98)00611-2

[Google Scholar](#)

[3] "Creativity Machine® Paradigm"

Stephen L. Thaler

2013 | Encyclopedia of Creativity, Invention, Innovation and Entrepreneurship | DOI: 10.1007/978-1-4614-3858-8_396

[Google Scholar](#)

[4] "The Creativity Machine® Paradigm"

Stephen Thaler

2017 | Encyclopedia of Creativity, Invention, Innovation and Entrepreneurship | DOI: 10.1007/978-1-4614-6616-1_396-2

[Google Scholar](#)

[5] "Content"

Harold D. Lasswell, Ralph Casey, Bruce Smith

2010 | University of Minnesota Press eBooks | DOI: 10.5749/9781452938011-toc

[Google Scholar](#)

[6] "1. Future Systems of Identity in the World Community 3 Harold Lasswell"

Harold D . Lasswell

2015 | The Future of the International Legal Order, Volume 4: The Structure of the International Environment | DOI: 10.1515/9781400873074-003

[Google Scholar](#)

[7] "A Mathematical Theory of Communication"

Claude E. Shannon

1948 | Bell System Technical Journal | DOI: 10.1002/j.1538-7305.1948.tb01338.x

[Google Scholar](#)

[8] "The Mathematical Theory of Communication"

Claude E. Shannon, Warren Weaver, Norbert Wiener

1950 | Physics Today | DOI: 10.1063/1.3067010

[Google Scholar](#)

[9] "Communication Theory of Secrecy Systems*"

Claude E. Shannon

1949 | Bell System Technical Journal | DOI: 10.1002/j.1538-7305.1949.tb00928.x

[Google Scholar](#)

[10] "RECENT CONTRIBUTIONS TO THE MATHEMATICAL THEORY OF COMMUNICATION"

Warren Weaver

2017 | ETC: A Review of General Semantics | DOI: 10.66774/etc_7401_136

[Google Scholar](#)

[11] "Cybernetics or Control and Communication in the Animal and the Machine"

Norbert Wiener

2019 | The MIT Press eBooks | DOI: 10.7551/mitpress/11810.001.0001

[Google Scholar](#)

[12] "The Human Use of Human Beings: Cybernetics and Society"

Stephen B. Miles, Norbert Wiener

1951 | Land Economics | DOI: 10.2307/3159747

[Google Scholar](#)

[13] "Verbal reports as data."

K. Anders Ericsson, Herbert A. Simon

1980 | Psychological Review | DOI: 10.1037/0033-295x.87.3.215

[Google Scholar](#)

[14] "Political Communication in Media Society: Does Democracy Still Enjoy an Epistemic Dimension? The Impact of Normative Theory on Empirical Research"

Jürgen Habermas

2006 | Communication Theory | DOI: 10.1111/j.1468-2885.2006.00280.x

[Google Scholar](#)

[15] "Religion in the Public Sphere"

Jürgen Habermas

2006 | European Journal of Philosophy | DOI: 10.1111/j.1468-0378.2006.00241.x

[Google Scholar](#)

[16] "Towards a theory of communicative competence"

Jürgen Habermas

1970 | Inquiry | DOI: 10.1080/00201747008601597

[Google Scholar](#)

[17] "THE CONCEPT OF HUMAN DIGNITY AND THE REALISTIC UTOPIA OF HUMAN RIGHTS"

Jürgen Habermas

2010 | Metaphilosophy | DOI: 10.1111/j.1467-9973.2010.01648.x

[Google Scholar](#)

[18] "Learning from the Behavior of Others: Conformity, Fads, and Informational Cascades"

Sushil Bikhchandani, David A. Hirshleifer, Ivo Welch

1998 | The Journal of Economic Perspectives | DOI: 10.1257/jep.12.3.151

[Google Scholar](#)

[19] "A Theory of Fads, Fashion, Custom, and Cultural Change as Informational Cascades"

Sushil Bikhchandani, David A. Hirshleifer, Ivo Welch

1992 | Journal of Political Economy | DOI: 10.1086/261849

[Google Scholar](#)

[20] "The Package Assignment Model"

Sushil Bikhchandani, Joseph M. Ostroy

2002 | Journal of Economic Theory | DOI: 10.1006/jeth.2001.2957

[Google Scholar](#)

[21] "Equilibria in open common value auctions"

Sushil Bikhchandani, John G. Riley

1991 | Journal of Economic Theory | DOI: 10.1016/0022-0531(91)90144-s

[Google Scholar](#)

[22] "Good Day Sunshine: Stock Returns and the Weather"

David A. Hirshleifer, Tyler Shumway

2003 | The Journal of Finance | DOI: 10.1111/1540-6261.00556

[Google Scholar](#)

[23] "Information Cascades and Social Learning"

Sushil Bikhchandani, David Hirshleifer, Omer Tamuz, Ivo Welch

2021 | DOI: 10.3386/w28887

[Google Scholar](#)

[24] "AMSTAR 2: a critical appraisal tool for systematic reviews that include randomised or non-randomised studies of healthcare interventions, or both"

Beverley Shea, Barnaby C Reeves, George A. Wells, Micere Thuku, Candyce Hamel

2017 | BMJ | DOI: 10.1136/bmj.j4008

[Google Scholar](#)

[25] "Social Influence: Compliance and Conformity"

Robert B. Cialdini, Noah J. Goldstein

2004 | Annual Review of Psychology | DOI: 10.1146/annurev.psych.55.090902.142015

[Google Scholar](#)

[26] "A focus theory of normative conduct: Recycling the concept of norms to reduce littering in public places."

Robert B. Cialdini, Raymond R. Reno, Carl A. Kallgren

1990 | Journal of Personality and Social Psychology | DOI: 10.1037/0022-3514.58.6.1015

[Google Scholar](#)

[27] "The Constructive, Destructive, and Reconstructive Power of Social Norms"

P. Wesley Schultz, Jessica M. Nolan, Robert B. Cialdini, Noah J. Goldstein, Vidas Griskevicius

2007 | Psychological Science | DOI: 10.1111/j.1467-9280.2007.01917.x

[Google Scholar](#)

[28] "A Room with a Viewpoint: Using Social Norms to Motivate Environmental Conservation in Hotels"

Noah J. Goldstein, Robert B. Cialdini, Vidas Griskevicius

2008 | Journal of Consumer Research | DOI: 10.1086/586910

[Google Scholar](#)

[29] "New Directions in Agenda-Setting Theory and Research"

Maxwell E. McCombs, Donald L. Shaw, David H. Weaver

2018 | Advances in Foundational Mass Communication Theories | DOI: 10.4324/9781315164441-9

[Google Scholar](#)

[30] "A BEHAVIORAL MODEL OF RATIONAL CHOICE"

Herbert A. Simon

1955 | The Quarterly Journal of Economics | DOI: 10.2307/1884852

[Google Scholar](#)

[31] "Human Problem Solving."

Nick Axten, Allen Newell, Herbert A. Simon

1973 | Contemporary Sociology A Journal of Reviews | DOI: 10.2307/2063712

[Google Scholar](#)

[32] "Administrative Behavior: A Study of Decision-Making Processes in Administrative Organization."

M. Craig Brown, Herbert A. Simon

1978 | Contemporary Sociology A Journal of Reviews | DOI: 10.2307/2065048

[Google Scholar](#)

[33] "Social influence and opinions"

Noah E. Friedkin, Eugene C. Johnsen

1990 | Journal of Mathematical Sociology | DOI: 10.1080/0022250x.1990.9990069

[Google Scholar](#)

[34] "A Structural Theory of Social Influence"

Noah E. Friedkin

1998 | Cambridge University Press eBooks | DOI: 10.1017/cb09780511527524

[Google Scholar](#)

[35] "Network Studies of Social Influence"

Peter V. Marsden, Noah E. Friedkin

1993 | Sociological Methods & Research | DOI: 10.1177/0049124193022001006

[Google Scholar](#)

[36] "Theoretical Foundations for Centrality Measures"

Noah E. Friedkin

1991 | American Journal of Sociology | DOI: 10.1086/229694

[Google Scholar](#)

[37] "Social positions in influence networks"

Noah E. Friedkin, Eugene C. Johnsen

1997 | Social Networks | DOI: 10.1016/s0378-8733(96)00298-5

[Google Scholar](#)

[38] "ATTITUDE CHANGE, AFFECT CONTROL, AND EXPECTATION STATES IN THE FORMATION OF INFLUENCE NETWORKS"

Noah E. Friedkin, Eugene C. Johnsen

Advances in Group Processes | DOI: 10.1016/s0882-6145(03)20001-1

[Google Scholar](#)

[39] "Systems of Logic Based on Ordinals[†]"

Alan Mathison Turing

1939 | Proceedings of the London Mathematical Society | DOI: 10.1112/plms/s2-45.1.161

[Google Scholar](#)

[40] "ROUNDING-OFF ERRORS IN MATRIX PROCESSES"

Alan Mathison Turing

1948 | The Quarterly Journal of Mechanics and Applied Mathematics | DOI: 10.1093/qjmam/1.1.287

[Google Scholar](#)

[41] "Intelligent Machinery, A Heretical Theory"

Alan M. Turing

2004 | The Turing Test | DOI: 10.7551/mitpress/6928.003.0014

[Google Scholar](#)

[42] "Report on Speech Secrecy System DELILAH, a Technical Description Compiled by A. M. Turing and Lieutenant D. Bayley REME, 1945–1946"

A. Turing, D. Bayley

2012 | Cryptologia | DOI: 10.1080/01611194.2012.713803

[Google Scholar](#)

[43] "Carbon sequestration at the Illinois Basin-Decatur Project: experimental results and geochemical simulations of storage"

Peter Berger, Lois E. Yoksoulain, Jared T. Freiburg, Shane K. Butler, William R. Roy

2019 | Environmental Earth Sciences | DOI: 10.1007/s12665-019-8659-4

[Google Scholar](#)

[44] "Geochemical Controls of Coal Fly Ash Leachate pH"

William R. Roy, Peter Berger

2011 | Coal Combustion and Gasification Products | DOI: 10.4177/ccgp-d-11-00013.1

[Google Scholar](#)

[45] "Scalable graph signal recovery for big data over networks"

Alexander Jung, Peter Berger, Gábor Hannák, Gerald Matz

2016 | DOI: 10.1109/spawc.2016.7536869

[Google Scholar](#)

[46] "The Social Construction of Reality: A Treatise in the Sociology of Knowledge."

George L. Simpson, Peter Ludwig Berger, Thomas Luckmann

1967 | American Sociological Review | DOI: 10.2307/2091739

[Google Scholar](#)

[47] "The Social Construction of Reality"

Daniel M. Rose, Peter Ludwig Berger, Thomas Luckmann

1967 | Modern Language Journal | DOI: 10.2307/323448

[Google Scholar](#)

[48] "Organization Theory: The Classical Constructions"

Niklas Luhmann

2018 | Cambridge University Press eBooks | DOI: 10.1017/9781108560672.002

[Google Scholar](#)

[49] "Organization as an Autopoietic System"

Niklas Luhmann

2018 | Cambridge University Press eBooks | DOI: 10.1017/9781108560672.003

[Google Scholar](#)

[50] "Conclusion: Theory and Practice"

Niklas Luhmann

2018 | Cambridge University Press eBooks | DOI: 10.1017/9781108560672.017

[Google Scholar](#)

[51] "Structural Change: The Poetry of Reform and the Reality of Evolution"

Niklas Luhmann

2018 | Cambridge University Press eBooks | DOI: 10.1017/9781108560672.012

[Google Scholar](#)

[52] "Reassembling the Social"

Bruno Latour

2005 | DOI: 10.1093/oso/9780199256044.001.0001

[Google Scholar](#)

[53] "Laboratory Life: The Construction of Scientific Facts"

Bruno Latour

1986 | DigitalGeorgetown (Georgetown University Library)

[Google Scholar](#)

[54] "Why Has Critique Run out of Steam? From Matters of Fact to Matters of Concern"

Bruno Latour

2004 | Critical Inquiry | DOI: 10.1086/421123

[Google Scholar](#)

[55] "Pandora's Hope: Essays on the Reality of Science Studies"

Monica J. Casper, Bruno Latour

2000 | Contemporary Sociology A Journal of Reviews | DOI: 10.2307/2655272

[Google Scholar](#)

[56] "A Theory of Group Stability"

Kathleen M. Carley

1991 | American Sociological Review | DOI: 10.2307/2096108

[Google Scholar](#)

[57] "On the robustness of centrality measures under conditions of imperfect data"

Stephen P. Borgatti, Kathleen M. Carley, David M. Krackhardt

2005 | Social Networks | DOI: 10.1016/j.socnet.2005.05.001

[Google Scholar](#)

[58] "Simulation modeling in organizational and management research"

J. Richard Harrison, Zhiang Lin, Glenn R. Carroll, Kathleen M. Carley

2007 | Academy of Management Review | DOI: 10.5465/amr.2007.26586485

[Google Scholar](#)

[59] "Is the Sample Good Enough? Comparing Data from Twitter's Streaming API with Twitter's Firehose"

Fred Morstatter, Jürgen Pfeffer, Huan Liu, Kathleen M. Carley

2021 | Proceedings of the International AAAI Conference on Web and Social Media | DOI: 10.1609/icwsm.v7i1.14401

[Google Scholar](#)

[60] "AI models collapse when trained on recursively generated data"

Ilia Shumailov, Zakhar Shumaylov, Yiren Zhao, Nicolas Papernot, Ross Anderson

2024 | Nature | DOI: 10.1038/s41586-024-07566-y

[Google Scholar](#)

[61] "Reconstructing Individual Data Points in Federated Learning Hardened with Differential Privacy and Secure Aggregation"

Franziska Boenisch, Adam Dziedzic, Roei Schuster, Ali Shahin Shamsabadi, Ilia Shumailov

2023 | DOI: 10.1109/eurosp57164.2023.00023

[Google Scholar](#)

[62] "Fairness Feedback Loops: Training on Synthetic Data Amplifies Bias"

Sierra Wyllie, Ilia Shumailov, Nicolas Papernot

2024 | The 2024 ACM Conference on Fairness, Accountability, and Transparency | DOI: 10.1145/3630106.3659029

[Google Scholar](#)

[63] "REMEMBERING: A STUDY IN EXPERIMENTAL AND SOCIAL PSYCHOLOGY"

Frederic Charles Bartlett, Cyril Burt

1933 | British Journal of Educational Psychology | DOI: 10.1111/j.2044-8279.1933.tb02913.x

[Google Scholar](#)

[64] "Thinking : an experimental and social study"

Frederic Charles Bartlett

1958 | Allen & Unwin eBooks

[Google Scholar](#)

[65] "REMEMBERING AS A STUDY IN SOCIAL PSYCHOLOGY"

F. Bartlett

1995

[Google Scholar](#)

[66] "The Study of Society (RLE Social Theory) : Methods and Problems"

F. Bartlett

2014 | DOI: 10.4324/9781315763217

[Google Scholar](#)

[67] "SENSORY MEASUREMENT OF FOOD TEXTURE BY FREE-CHOICE PROFILING"

Richard James Marshall, Simon Kirby

1988 | Journal of Sensory Studies | DOI: 10.1111/j.1745-459x.1988.tb00430.x

[Google Scholar](#)

[68] "Fitting Emax models to clinical trial dose–response data"

Simon Kirby, P. Brain, Byron C. Jones

2010 | Pharmaceutical Statistics | DOI: 10.1002/pst.432

[Google Scholar](#)

[69] "Iterated Learning: A Framework for the Emergence of Language"

Kenny Smith, Simon Kirby, Henry Brighton

2003 | Artificial Life | DOI: 10.1162/106454603322694825

[Google Scholar](#)

[70] "Advances in Neural Information Processing Systems 21"

Thomas L. Griffiths, Christopher G. Lucas, Joseph Jay Williams, Michael L. Kalish

2009 | Neural Information Processing Systems

[Google Scholar](#)

[71] "Natural speech reveals the semantic maps that tile human cerebral cortex"

Alexander G. Huth, Wendy A. de Heer, Thomas L. Griffiths, Frédéric E. Theunissen, Jack L. Gallant

2016 | Nature | DOI: 10.1038/nature17637

[Google Scholar](#)

[72] "Toward a Rational and Mechanistic Account of Mental Effort"

Amitai Shenhav, Sebastian Musslick, Falk Lieder, Wouter Kool, Thomas L. Griffiths

2017 | Annual Review of Neuroscience | DOI: 10.1146/annurev-neuro-072116-031526

[Google Scholar](#)

[73] "Probabilistic Topic Models"

Mark Steyvers, Thomas L. Griffiths

2015 | DOI: 10.4324/9780203936399.ch21

[Google Scholar](#)

[74] "Cumulative cultural evolution in the laboratory: An experimental approach to the origins of structure in human language"

Simon Kirby, Hannah Cornish, Kenny Smith

2008 | Proceedings of the National Academy of Sciences | DOI: 10.1073/pnas.0707835105

[Google Scholar](#)

[75] "Proceedings of the 39th Annual Conference of the Cognitive Science Society"

Yasamin Motamedi-mousavi, Marieke Schouwstra, Jennifer Culbertson, Kenny Smith, Simon Kirby

2017

[Google Scholar](#)

[76] "Grounding Gaps in Language Model Generations"

Omar Shaikh, Kristina Gligorić, Ashna Khetan, Matthias Gerstgrasser, Diyi Yang

2024 | Proceedings of the 2024 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 1: Long Papers) | DOI: 10.18653/v1/2024.naacl-long.348

[Google Scholar](#)

[77] "Is Model Collapse Inevitable? Breaking the Curse of Recursion by Accumulating Real and Synthetic Data"

Matthias Gerstgrasser, Rylan Schaeffer, Apratim Dey, Rafael Rafailov, Henry Sleight

2024 | arXiv (Cornell University) | DOI: 10.48550/arxiv.2404.01413

[Google Scholar](#)

[78] "Bayesian Gaussian Process Classification from Event-Related Brain Potentials in Alzheimer's Disease"

Wolfgang Fruehwirt, Pengfei Zhang, Matthias Gerstgrasser, Dieter Grossegger, Reinhold Ernst Schmidt

2017 | Lecture notes in computer science | DOI: 10.1007/978-3-319-59758-4_7

[Google Scholar](#)

[79] "Collapse or Thrive? Perils and Promises of Synthetic Data in a Self-Generating World"

Joshua Kazdan, Rylan Schaeffer, Apratim Dey, Matthias Gerstgrasser, Rafael Rafailov

2024 | arXiv (Cornell University) | DOI: 10.48550/arxiv.2410.16713

[Google Scholar](#)

[80] "After Virtue: A Study in Moral Theory"

David B. Burrell, Alasdair MacIntyre

1984 | Journal of Law and Religion | DOI: 10.2307/1051043

[Google Scholar](#)

[81] "38. After Virtue: A Study in Moral Theory"

Alasdair MacIntyre

2015 | Princeton University Press eBooks | DOI: 10.1515/9781400848393-039

[Google Scholar](#)

[82] "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](#)

[83] "The Extended Mind"

Andy Clark, David J. Chalmers

1998 | Analysis | DOI: 10.1093/analys/58.1.7

[Google Scholar](#)

[84] "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](#)

[85] "An embodied cognitive science?"

Andy Clark, Andy Clark

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

[Google Scholar](#)

[86] "Metametaphysics : new essays on the foundations of ontology"

David J. Chalmers, David Manley, Ryan J. Wasserman

2009

[Google Scholar](#)

[87] "Perception and the Fall from Eden"

David J. Chalmers

2006 | Oxford University Press eBooks | DOI: 10.1093/acprof:oso/9780199289769.003.0003

[Google Scholar](#)

[88] "Self-determination theory and the facilitation of intrinsic motivation, social development, and well-being."

Richard M. Ryan, Edward L. Deci

2000 | American Psychologist | DOI: 10.1037/0003-066x.55.1.68

[Google Scholar](#)

[89] "The "What" and "Why" of Goal Pursuits: Human Needs and the Self-Determination of Behavior"

Edward L. Deci, Richard M. Ryan

2000 | Psychological Inquiry | DOI: 10.1207/s15327965pli1104_01

[Google Scholar](#)

[90] "Intrinsic Motivation and Self-Determination in Human Behavior"

Edward L. Deci, Richard M. Ryan

1985 | DOI: 10.1007/978-1-4899-2271-7

[Google Scholar](#)

[91] "Ironies of automation"

Lisanne Bainbridge

1983 | Automatica | DOI: 10.1016/0005-1098(83)90046-8

[Google Scholar](#)

[92] "Verbal reports as evidence of the process operator's knowledge"

Lisanne Bainbridge

1999 | International Journal of Human-Computer Studies | DOI: 10.1006/ijhc.1979.0307

[Google Scholar](#)

[93] "Technology Design to Support Caring for and Learning from Older Adults"

Ben Shneiderman

2026 | Gerontechnology | DOI: 10.4017/gt.2026.25.2.1693.3

[Google Scholar](#)

[94] "Introduction: Rising above the Levels of Automation"

Ben Shneiderman

2022 | Human-Centered AI | DOI: 10.1093/oso/9780192845290.003.0006

[Google Scholar](#)

[95] "Will Automation, AI, and Robots Lead to Widespread Unemployment?"

Ben Shneiderman

2022 | Human-Centered AI | DOI: 10.1093/oso/9780192845290.003.0004

[Google Scholar](#)

[96] "Human-centered AI"

Ben Shneiderman

2022 | Proceedings of the 5th Workshop on Human Factors in Hypertext | DOI: 10.1145/3538882.3542790

[Google Scholar](#)

[97] "Development as Freedom"

Amartya Kumar Sen

2009 | DOI: 10.1007/978-1-137-28787-8_94

[Google Scholar](#)

[98] "Rational Fools: A Critique of the Behavioral Foundations of Economic Theory"

Amartya Kumar Sen

1977 | Philosophy & Public Affairs

[Google Scholar](#)

[99] "Well-Being, Agency and Freedom: The Dewey Lectures 1984"

Amartya Kumar Sen

1985 | The Journal of Philosophy | DOI: 10.2307/2026184

[Google Scholar](#)

[100] "Digital Painting Analysis: Authentication and Artistic Style from Digital Reproductions"

Robert Dahl Jacobsen

2012

[Google Scholar](#)

[101] "Legal Aspects of Critical Habitat Determinations"

Robert Dahl Jacobsen

1980 | Ursus | DOI: 10.2307/3872834

[Google Scholar](#)

[102] "Which of Trump's Supreme Court choices might be most reliably conservative?"

Matthew Dahl

2020 | DOI: 10.64628/aai.f7mcenjnv

[Google Scholar](#)

Artykuły pokrewne (Google Scholar):

[103] "Always On"

Naomi S. Baron

2008 | Oxford University Press eBooks | DOI: 10.1093/acprof:oso/9780195313055.001.0001

[Google Scholar](#)

[104] "Texto, áudio ou vídeo: saiba com qual aprendemos melhor"

Naomi Baron

2023 | DOI: 10.64628/ade.as5maewsu

[Google Scholar](#)

[105] "How ChatGPT robs students of motivation to write and think for themselves"

Naomi Baron

2023 | DOI: 10.64628/aai.6javedefa

[Google Scholar](#)

[106] "What happens when you try to read Moby Dick on your smartphone?"

Naomi Baron

2015 | DOI: 10.64628/aai.4mvd59agv

[Google Scholar](#)

[107] "Print, audio ou vidéo : quels supports choisir pour mieux apprendre ?"

Naomi Baron

2021 | DOI: 10.64628/aak.4699kvwv4

[Google Scholar](#)

[108] "The power of touch is vital for both reading and writing"

Naomi Baron

2024 | DOI: 10.64628/aai.jy5psnvfp

[Google Scholar](#)

[109] "Why we remember more by reading – especially print – than from audio or video"

Naomi Baron

2021 | DOI: 10.64628/aai.yPWMjmg9j

[Google Scholar](#)

[110] "Écran ou papier... pourquoi tourner une page vaut mieux que cliquer"

Naomi Baron

2025 | DOI: 10.64628/aak.c4reqmk3y

[Google Scholar](#)

[111] "Comment ChatGPT sape la motivation à écrire et penser par soi-même"

Naomi Baron

2024 | DOI: 10.64628/aak.sy7t3kkcv

[Google Scholar](#)

[112] "Do students lose depth in digital reading?"

Naomi Baron

2016 | DOI: 10.64628/aai.kgspmj3tv

[Google Scholar](#)

[113] "Textbook merger could create more problems than just higher prices"

Naomi Baron

2019 | DOI: 10.64628/aai.k47y6vmun

[Google Scholar](#)

[114] "Letters by phone or speech by other means: the linguistics of email"

Naomi S. Baron

1998 | Language & Communication | DOI: 10.1016/S0271-5309(98)00005-6

[Google Scholar](#)

[115] "See you Online"

Naomi S. Baron

2004 | Journal of Language and Social Psychology | DOI: 10.1177/0261927X04269585

[Google Scholar](#)

[116] "Words Onscreen: The Fate of Reading in a Digital World"

Naomi S. Baron

2015

[Google Scholar](#)

[117] "Text Messaging and IM"

Rich Ling, Naomi S. Baron

2007 | Journal of Language and Social Psychology | DOI: 10.1177/0261927x06303480

[Google Scholar](#)

[118] "Media Multitasking and Cognitive, Psychological, Neural, and Learning Differences"

Melina R. Uncapher, Lin Lin, Larry D. Rosen, Heather L. Kirkorian, Naomi S. Baron

2017 | PEDIATRICS | DOI: 10.1542/peds.2016-1758d

[Google Scholar](#)

[119] "Cross-cultural patterns in mobile-phone use: public space and reachability in Sweden, the USA and Japan"

Naomi S. Baron, Ylva Hård af Segerstad

2010 | New Media & Society | DOI: 10.1177/1461444809355111

[Google Scholar](#)

[120] "The persistence of print among university students: An exploratory study"

Naomi S. Baron, Rachelle M. Calixte, Mazneen Havewala

2016 | Telematics and Informatics | DOI: 10.1016/j.tele.2016.11.008

[Google Scholar](#)

[121] "Alphabet to Email"

Naomi S. Baron

2002 | DOI: 10.4324/9780203194317

[Google Scholar](#)

[122] "Discourse structures in Instant Messaging: The case of utterance breaks"

Naomi S. Baron

2010 | Indiana Magazine of History (Indiana University)

[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: df0f9a97-e25c-407b-8870-594a44c74f4b