

Kryzys autorstwa w epoce sztucznej inteligencji: między funkcją a intencją w świetle myśli Naomi S. Baron



AUTOR

Fundacja Dobre Państwo

STRESZCZENIE / ABSTRACT

Artykuł analizuje transformację pojęcia autora w obliczu rozwoju generatywnej sztucznej inteligencji. Tradycyjne przekonanie, że za każdym tekstem stoi ludzki umysł i konkretna intencja, zostaje podważone przez modele językowe, które potrafią imitować sprawstwo bez posiadania wewnętrznych przeżyć. Autor wskazuje na rozpad figury autora na odrębne funkcje: od formułowania celu po redakcję końcową, co przesunęło pytanie o autorstwo z poziomu bibliograficznego na ontologiczny. Przywołując spór Turinga z Jeffersonem, tekst stawia kluczowe pytanie: czy w kulturze pisma liczy się jedynie funkcjonalny rezultat i skuteczność przekazu, czy też niezbywalnym elementem tekstu pozostaje autentyczne doświadczenie podmiotu, którego maszyna nie jest w stanie osiągnąć.

8. The article analyzes the transformation of the concept of the author in the face of the development of generative artificial intelligence. The traditional belief that a human mind and a specific intention stand behind every text is challenged by language models capable of imitating agency without possessing internal experiences. The author points to the fragmentation of the figure of the author into separate functions: from formulating the goal to final editing, which shifts the question of authorship from a bibliographic level to an ontological one. Invoking the dispute between Turing and Jefferson, the text poses a key question: in a writing culture, does only the functional result and effectiveness of communication matter, or does the authentic experience of the subject—which a machine cannot possess—remain an indispensable element of the text?

Artykuł analizuje fundamentalny kryzys figury autora w dobie generatywnej sztucznej inteligencji, stawiając tezę, że AI nie tyle konkuruje z pisarzem, co rozmontowuje tradycyjne rozumienie autorstwa. Autor dowodzi, że technologia ta oddziela zewnętrzną formę intencjonalnej wypowiedzi od rzeczywistego podmiotu przeżywającego znaczenie, co prowadzi do niebezpiecznego złudzenia: utożsamiania retorycznej płynności tekstu z kompetencją poznawczą i prawdą. Kluczowym punktem analizy jest rozróżnienie między pisanem jako mechaniczną produkcją znaków a pisanem jako procesem myślowym; delegowanie tego drugiego procesu na algorytmy może prowadzić do atrofii ludzkich zdolności analitycznych, gdyż trud formułowania myśli jest warunkiem rozwoju intelektualnego. Tekst ostrzega przed 'pokusą efektywności', która redukuje człowieka do roli post-edytora lub biologicznej pieczęci odpowiedzialności, i wzywa do redefinicji autorstwa jako struktury sprawstwa, gdzie wartość dzieła nie wynika z gładkości prozy, lecz z zakotwiczenia tekstu w autentycznym doświadczeniu, etyce i odpowiedzialności podmiotu.

SŁOWA KLUCZOWE

kryzys autorstwa

generatywna sztuczna inteligencja

intencjonalność tekstu

modele językowe

funkcja autora

proces poznawczy pisania

post-editing maszynowy

paradoks drabiny kompetencyjnej

desirable difficulties

efekt zakotwiczenia

odpowiedzialność normatywna

sprawstwo operacyjne

pokusa efektywności

imitacja intencji

struktura sprawstwa

AI rozmontowuje tradycyjną figurę autora

Kto to napisał? Człowiek, maszyna i kryzys autorstwa w epoce sztucznej inteligencji

Kto to napisał? To pytanie dotyka dziś kryzysu najstarszego założenia kultury pisma. Przez kilka tysięcy lat funkcjonowało przekonanie tak oczywiste, że niemal nigdy nie trzeba było go wypowiadać: jeśli istnieje tekst, za nim stoi człowiek. Ktoś chciał coś przekazać, dokonał wyboru słów, zaryzykował twierdzenie, świadomie przemilczał inne i nadał wypowiedzi formę, by w efekcie – przynajmniej w zasadzie – móc odpowiedzieć na pytanie, dlaczego napisał właśnie to. Autor nie

był jedynie producentem znaków; stanowił hipotetyczne centrum odpowiedzialności, intencji i znaczenia.

Nawet anonimowy pamflet zakładał istnienie anonimowego człowieka, fałszywy list – fałszerza, a propagandowy komunikat – propagandystę. Pismo oddzielało wypowiedź od ciała autora, lecz nie oddzielało jej od wyobrażenia ludzkiego umysłu. Generatywna sztuczna inteligencja zakłóca właśnie tę wielowiekową oczywistość. Po raz pierwszy na skalę przemysłową pojawiają się teksty, które noszą wszystkie zewnętrzne znamiona intencjonalnej wypowiedzi, choć mechanizm ich powstawania nie wymaga podmiotu przeżywającego znaczenie każdego napisanego zdania.

Pytanie „kto to napisał?” przestaje być zatem pytaniem bibliograficznym, a staje się pytaniem ontologicznym. Nie chodzi już tylko o ustalenie nazwiska pod tekstem, lecz o określenie tego, co właściwie rozumiemy przez autorstwo. Czy autorem jest ten, kto fizycznie generuje ciąg znaków? Ten, kto formułuje cel lub wybiera spośród wariantów? Ten, kto posiada doświadczenie, do którego tekst się odnosi, czy może ten, kto sprawuje kontrolę nad ostatecznym kształtem wypowiedzi? A może dopiero ten, kto jest gotów ponieść społeczne, prawne i moralne konsekwencje własnych słów?

Epoka modeli językowych rozbija jedno pozornie proste pojęcie na wiązkę odrębnych funkcji. Dotychczas występowały one zazwyczaj razem: człowiek wymyślał, formułował, poprawiał, podpisywał i odpowiadał. Dziś funkcje te można rozdzielić między użytkownika, model, twórców systemu, dostawców danych, redaktorów, platformę dystrybucyjną oraz algorytm decydujący o tym, komu tekst zostanie wyświetlony.

Sztuczna inteligencja nie tylko konkuruje z pisarzem – ona rozmontowuje tradycyjną figurę autora na części składowe. Człowiek może dostarczyć zamiar, maszyna pierwszą wersję, człowiek kryteria wyboru, kolejny system korektę stylistyczną, a platforma automatycznie dopasować nagłówki do przewidywanego zachowania odbiorcy. Kto w takim łańcuchu „napisał” tekst? Problem zaczyna przypominać kwestię autorstwa katedry, filmu czy programu komputerowego, lecz z istotną różnicą: w klasycznej współpracy ludzkiej wszyscy uczestnicy są podmiotami zdolnymi do posiadania zamiarów. Model językowy uczestniczy w procesie inaczej. Potrafi wytworzyć formę, która wygląda jak rezultat zamiaru, choć nie musi posiadać intencji w ludzkim sensie tego pojęcia. Powstaje zatem osobliwy fenomen kulturowy: intencjonalność może być przekonująco imitowana na poziomie tekstu bez konieczności wykazania intencjonalności po stronie systemu, który go wygenerował.

W tym kontekście stary spór między Alanem Turingiem a Geoffreyem Jeffersonem nabiera po ponad siedemdziesięciu latach nowej ostrości. Turing proponował pragmatyczne przesunięcie

pytania z niedostępnego wnętrza maszyny na obserwowalne zachowanie. Zamiast zastanawiać się, czy maszyna „naprawdę myśli”, można sprawdzić, czy potrafi zachowywać się językowo w sposób nierozróżnialny od człowieka. Jefferson domagał się więcej. Nie wystarczał mu sonet jako poprawnie ułożony ciąg słów; maszyna musiałaby napisać go dlatego, że coś przeżyła, oraz wiedzieć, że to zrobiła. Oto dwa kryteria inteligencji, które do dziś prowadzą nas w przeciwnych kierunkach.

Różnica między opisem doświadczenia a jego posiadaniem

Pierwsze pytanie dotyczy skuteczności wykonania, drugie zaś wewnętrznego źródła działania. Im doskonalsze stają się modele generatywne, tym mniej ten spór przypomina historię filozofii techniki, a bardziej instrukcję obsługi współczesności. Można wyobrazić sobie sytuację graniczną: człowiek i maszyna tworzą dwa teksty, których kompetentny czytelnik nie potrafi odróżnić. Oba są logiczne, eleganckie, poprawne faktograficznie, stylistycznie interesujące i odpowiednio dostosowane do odbiorcy. Jeśli jakość produktu końcowego jest identyczna, czy różnica w pochodzeniu nadal ma znaczenie?

Dla skrajnego funkcjonalisty odpowiedź może brzmieć: nie. Jeżeli zadaniem tekstu jest skuteczne przekazanie informacji, a oba utwory realizują to zadanie równie dobrze, poszukiwanie metafizycznej "duszy autora" wydaje się zbędne. Tak rozumuje znaczna część gospodarki: liczy się rezultat, czas, koszt i niezawodność. Faktura nie zyskuje na wartości dlatego, że księgowy doznał egzystencjalnego olśnienia podczas jej sporządzania. Problem polega jednak na tym, że kultura pisma nigdy nie składała się wyłącznie z faktur. List miłosny, zeznanie świadka, autobiografia, rozprawa filozoficzna, polemika polityczna, wiersz, pamiętnik, wyznanie winy czy świadectwo prześladowania mają wartość częściowo dlatego, że łączymy wypowiedź z doświadczeniem konkretnej osoby. Gdy ktoś pisze „bałem się”, czytelnik zakłada istnienie podmiotu, który ten strach faktycznie odczuł. Gdy autor stwierdza „zmieniłem zdanie”, zakładamy historię wcześniejszego przekonania i jego późniejszej rewizji.

Gdy człowiek przeprosza, liczy się nie tylko poprawność składni aktu przeprosin, ale relacja między słowem a odpowiedzialnością mówiącego. Maszyna potrafi znakomicie odwzorować język lęku, skruchy, tęsknoty czy zachwyty, lecz nie wynika z tego automatycznie, że posiada biografię, której tekst miałby być ekspresją. Różnica między opisaniem doświadczenia a posiadaniem doświadczenia okazuje się jednym z najważniejszych pęknięć współczesnej teorii autorstwa. Naomi S. Baron

ujmuje ten problem za pomocą sugestywnej figury "maszynki do mielenia języka". Metafora ta jest celowo nieelegancka, by przeciwstawić ją elegancji produktu. Na zewnątrz widzimy gładkie zdanie.

Wewnątrz kryje się jednak ogromny proces przekształcania wcześniejszych form językowych, korelacji, regularności i wzorców. To, że osłonka jest estetyczna, nie mówi nam nic o tym, co znajdowało się w środku ani jak rozumieć relację między wejściem a wynikiem. W epoce modeli językowych szczególnie niebezpiecznym błędem poznawczym staje się utożsamienie płynności z prawdą, elegancji z wiedzą i językowej pewności z rzeczywistą pewnością epistemiczną. Człowiek odruchowo odczytuje dobry styl jako sygnał obecności kompetentnego umysłu. Przez tysiąclecia była to heurystyka dostatecznie użyteczna, lecz generatywna AI zmienia warunki jej działania. Nie oznacza to, że człowiek jest z definicji wiarygodny, oryginalny czy głęboki.

Pisanie jako proces poznawczy a pułapka efektywności

Historia literatury, polityki i administracji dostarcza nadmiaru dowodów przeciwnych. Ludzie również posługują się kliszami, kłami, powtarzają cudze poglądy, piszą rzeczy, których nie rozumieją i podpisują dokumenty przygotowane przez innych. Ghostwriting wyprzedza sztuczną inteligencję o stulecia, propaganda o tysiąclecia, a biurokratyczna bezmyślność jest prawdopodobnie niemal tak stara jak sama biurokracja. Technologia nie wynalazła zatem nieautentyczności; zrobiła coś znacznie ciekawszego: drastycznie obniżyła koszt produkowania jej przekonującej imitacji.

Gdy stworzenie poprawnej, rozbudowanej i stylistycznie wiarygodnej wypowiedzi wymagało czasu oraz kompetencji, sam trud jej wytworzenia działał jako niedoskonały filtr. W momencie jednak, gdy koszt marginalny kolejnego tekstu zbliża się do zera, produkcja pozoru intencjonalności może osiągnąć skalę, której dawne instytucje komunikacyjne nie zostały zaprojektowane obsługiwać. Debata o autorstwie styka się tu bezpośrednio z ekonomią. Generatywna AI jest potężna nie dlatego, że napisała jeden imponujący tekst — człowiek potrafiłby odpowiedzieć na taki sukces własnym dziełem. Jej przewaga tkwi w możliwości tworzenia milionów wariantów, ich personalizacji, testowania i dystrybucji z prędkością nieosiągalną dla biologicznego autora.

Roald Dahl w "The Great Automatic Grammatizator" rozpoznał ten problem zadziwiająco wcześnie. Najbardziej przenikliwym elementem jego satyry nie była wizja maszyny konstruującej opowieści, lecz ekonomiczna konsekwencja jej pojawienia się. Jeśli twórczość można uprzemysłowić, rynek zaczyna premiować nie tego, kto najgłębiej przeżywa, lecz tego, kto dysponuje maszyną, skalą dystrybucji i zdolnością do przejęcia uwagi. Pytanie o autorstwo staje się wówczas jednocześnie pytaniem o własność infrastruktury wytwarzania języka.

Zasadniczy problem pojawia się jednak wcześniej, zanim tekst trafi na rynek. Dotyczy tego, co dzieje się z człowiekiem w trakcie samego pisania. W potocznym wyobrażeniu myśl poprzedza zdanie: najpierw wiemy, co chcemy powiedzieć, a później po prostu przelewamy to na papier. Doświadczenie pisarzy, filozofów i badaczy przeczy tej prostocie. Flannery O'Connor ujmowała to paradoksalnie: pisała między innymi po to, by odkryć, co właściwie myśli. Podobną intuicję mają autorzy, dla których pierwsza wersja tekstu nie jest opakowaniem gotowej idei, lecz laboratorium, w którym ta idea dopiero nabiera kształtu.

Skreślenie słowa może oznaczać korektę sądu. Zmiana kolejności akapitów potrafi ujawnić wadę rozumowania. Próba znalezienia odpowiedniej metafory często obnaża fakt, że autor nie rozumie jeszcze zjawiska, które chciał opisać. Pisanie nie jest więc wyłącznie urządzeniem wyjściowym mózgu; jest częścią obwodu poznawczego. To rozpoznanie zasadniczo zmienia sposób oceny sztucznej inteligencji. Gdyby pisanie było tylko mechanicznym etapem ekspresji gotowych myśli, jego automatyzacja przypominałaby zastąpienie ręcznego przepisywania drukarką. Skoro jednak formułowanie zdań uczestniczy w powstawaniu przekonań, delegowanie pisania może oznaczać delegowanie części operacji poznawczej.

Nie zawsze musi to być strata. Kalkulator uwolnił człowieka od rutynowych rachunków, nie niszcząc matematyki; edytor tekstu usunął konieczność ręcznego przepisywania stron po drobnej korekcie, nie unicestwiając literatury. Historia technologii to historia owocnego odciążania człowieka. Dlatego uczciwe pytanie nie brzmi: "czy należy delegować cokolwiek?", lecz: które elementy wysiłku są kosztownym balastem, a które stanowią proces, dzięki któremu rozwija się kompetencja?

Rozróżnienie to będzie kluczowe dla dalszej analizy. W pedagogice dobrze znamy sytuacje, w których trudność nie jest wadą procesu, lecz jego mechanizmem. Uczeń, który samodzielnie próbuje przypomnieć sobie informację, wykonuje inną pracę poznawczą niż ten, który natychmiast odczytuje odpowiedź. Człowiek streszczający skomplikowany argument własnymi słowami sprawdza zarazem, czy go zrozumiał. Autor zmuszony uporządkować sprzeczne intuicje wykonuje pracę, której rezultatem nie jest jedynie bardziej elegancki akapit, lecz bardziej uporządkowany umysł. W tym świetle "pokusę efektywności", o której pisze Baron, można odczytać nie jako naiwną krytykę postępu, lecz jako ostrzeżenie przed błędnym rachunkiem kosztów. Narzędzie może skracać czas wykonania zadania, a jednocześnie usuwać z niego właśnie ten fragment, który miał wartość rozwojową. Nie każda pomoc AI prowadzi więc do poznawczej kapitulacji.

Autorstwo jako struktura sprawstwa i odpowiedzialności

Między autorem samotnie zapisującym każde słowo a człowiekiem bezrefleksyjnie podpisującym cudzy tekst rozciąga się szerokie kontinuum. System można wykorzystać do wyszukania kontrargumentów, sprawdzenia spójności struktury, wskazania niejasnych fragmentów lub stworzenia wariantów, które autor następnie świadomie oceni. Można jednak poprosić o „napisanie wszystkiego”, pobieżnie przejrzeć wynik i zaakceptować go tylko dlatego, że brzmi wystarczająco profesjonalnie.

Technicznie oba zachowania klasyfikujemy jako „korzystanie z AI”. Poznawczo są one niemal przeciwieństwami. W pierwszym przypadku maszyna staje się instrumentem zwiększającym tarcie intelektualne, gdyż podsuwa pytania i konkurencyjne perspektywy. W drugim usuwa tarcie całkowicie, redukując człowieka do operatora przycisku akceptacji. Stąd autorstwo w epoce generatywnej AI należy analizować nie binarnie, lecz funkcjonalnie. Nie wystarczy pytać, czy maszyna uczestniczyła w powstawaniu tekstu; trzeba ustalić, które funkcje twórcze zostały delegowane.

Kto określił problem? Kto zdecydował o wiarygodności źródeł? Kto stworzył strukturę argumentu i rozpoznał sprzeczność? Kto dokonał wyboru wartościującego lub odrzucił atrakcyjną, lecz fałszywą tezę? Kto wie, dlaczego poszczególne zdania znalazły się w tekście i kto jest w stanie obronić tekst, gdy odbiorca zacznie zadawać pytania? Autorstwo okazuje się nie tyle własnością dokumentu, ile strukturą sprawstwa rozłożoną na kolejne etapy powstawania znaczenia.

Można sformułować pierwszą roboczą zasadę: im bardziej wartość tekstu zależy od intencji, doświadczenia, rozumienia i odpowiedzialności autora, tym mniej wystarczającym kryterium jego jakości jest sam wygląd produktu końcowego. W instrukcji technicznej bezbłądność może dominować nad biografią twórcy; w osobistym świadectwie pochodzenie wypowiedzi staje się częścią jej znaczenia. Między tymi biegunami rozciąga się domena tekstów naukowych, dziennikarskich, prawnych, literackich, edukacyjnych i politycznych, w których znaczenie autorstwa trzeba dopiero określić. Nie istnieje jedna odpowiedź dla całej kultury pisma, tak jak nie istnieje jeden sposób użycia języka.

Technologia generatywna zmusza nas paradoksalnie do ponownego odkrycia czegoś bardzo starego: autorstwo nigdy nie było tylko produkowaniem słów. Było zobowiązaniem pomiędzy człowiekiem, jego wypowiedzią i odbiorcą. Autor mówił w istocie: to zdanie przedstawiam jako

moje; za tym rozumowaniem stoję; tego doświadczenia jestem świadkiem; tej fikcji nadaję formę; tę interpretację proponuję innym.

Pisanie jako proces kształtowania myśli

Maszyna może uczestniczyć w każdym z tych aktów na poziomie językowym, lecz nie wszystkie ich funkcje dają się sprowadzić do automatycznego generowania znaków. To właśnie w przestrzeni między tekstem, który jedynie przypomina wypowiedź, a wypowiedzią, za którą stoi konkretny podmiot, rozgrywa się właściwy spór o przyszłość pisania.

Zanim ocenimy, które z tych funkcji można bezpiecznie powierzyć algorytmom, musimy przyjrzeć się temu, co samo pisanie robiło z człowiekiem, zanim pojawiły się systemy gotowe przejąć ten trud. Historia edukacji, psychologia poznawcza i neurobiologia dowodzą bowiem, że ręka przesuwająca się po kartce, autor wracający po raz dziesiąty do nieudanego akapitu czy student próbujący nazwać własnymi słowami niejasną jeszcze ideę, wykonują pracę znacznie wykraczającą poza samą produkcję zdań. W tej pozornie nieefektywnej walce z językiem kryje się mechanizm, dzięki któremu człowiek nie tylko zapisuje swoje myśli, lecz uczy się je posiadać. Skoro autorstwo ma znaczenie, ponieważ łączy tekst z intencją, rozumieniem i odpowiedzialnością, należy zapytać, skąd właściwie bierze się owo rozumienie. Intuicja, że najpierw powstaje kompletna myśl, którą następnie ubieramy w słowa, jest kusząca, lecz poznawczo zbyt uproszczona. W rzeczywistym procesie tworzenia granica między myśleniem a pisanem okazuje się przepuszczalna.

Człowiek rozpoczyna zdanie z przybliżonym zamiarem, by w połowie odkryć, że wybrane pojęcie jest niewystarczające. Poprawia konstrukcję, dostrzega ukrytą sprzeczność, zmienia przykład, a wraz z nim koryguje samą tezę. Tekst nie jest zatem jedynie śladem zakończonego rozumowania, lecz środowiskiem, w którym to rozumowanie się dokonuje. Z tej perspektywy kartka, ekran czy edytor tekstu nie są neutralnymi pojemnikami na gotową treść; uczestniczą one aktywnie w organizowaniu uwagi, pamięci roboczej i relacji między pojęciami.

Ta właściwość pisania tłumaczy powracającą w literaturze maksymę, że autor pisze także po to, by dowiedzieć się, co właściwie ma do powiedzenia. Flannery O'Connor ujmowała tę intuicję z właściwą sobie przewrotnością: nie zna w pełni własnej myśli, dopóki nie zobaczy jej zapisanej. Sens tej deklaracji nie tkwi w romantycznym kulcie natchnienia, lecz jest niemal laboratoryjny.

Myśl istniejąca wyłącznie w głowie może pozostawać mglista, gdyż pamięć robocza toleruje skróty i przeskoki, których nie wybacza zapisane zdanie. Gdy argument zostaje przeniesiony do zewnętrznej przestrzeni znaków, staje się obiektem: można porównać jego początek z końcem, znaleźć lukę,

dostrzec powtórzenie czy sprawdzić następstwo przesłanek. Pisanie działa więc jak urządzenie zwiększające rozdzielczość samokontroli intelektualnej. Myśl przekonująca w formie intuicji może rozpaść się podczas próby ujęcia jej w dwóch precyzyjnych akapitach. Można powiedzieć, że człowiek posiada dwa warsztaty myślenia jednocześnie. Pierwszym jest biologiczny układ poznawczy: pamięć, uwaga, język wewnętrzny, wyobraźnia i doświadczenie. Drugim jest zewnętrzne środowisko symboliczne: znaki na kartce, notatki, diagramy, skreślenia i archiwa wcześniejszych prób. Ta druga przestrzeń rozszerza możliwości pierwszej.

Matematyczka nie musi utrzymywać całego równania w pamięci, ponieważ część struktury znajduje się przed nią. Historyk zestawia źródła, pisarka przesuw scenę, filozof oddziela przesłankę od konsekwencji. W tym szerokim sensie kultura piśmienna od dawna praktykuje poznawcze „odciążanie”. Problem nie leży zatem w samym przeniesieniu części pracy poza mózg – cywilizacja rozwijała się właśnie dzięki takim transferom. Trudności zaczynają się wtedy, gdy narzędzie zewnętrzne przestaje jedynie przechowywać lub porządkować elementy rozumowania, a zaczyna wykonywać za nas operację, której samodzielne przeprowadzenie było warunkiem zdobycia kompetencji. Tu pojawia się zasadnicza różnica między kalkulatorem a generatorem argumentów. Kalkulator przejmuje obliczenie, lecz użytkownik wciąż definiuje problem, wybiera operację i ocenia wynik. Model językowy wchodzi znacznie głębiej: może zaproponować sam problem, jego interpretację, hierarchię argumentów i konkluzję. Odciążenie nie dotyczy wówczas pojedynczej technicznej czynności, lecz całej sekwencji procesów, z których tradycyjnie składało się uczenie myślenia poprzez pisanie. Różnica jest jakościowa.

Trud pisania jako warunek rozwoju kompetencji poznawczych

Zewnętrzny notatnik pomaga zapamiętać własną myśl; zewnętrzny generator może dostarczyć myśl, której użytkownik nigdy sam nie musiał skonstruować. Otrzymuje on produkt poznawczy bez pokonywania drogi, która zazwyczaj prowadzi do jego wytworzenia. Nie oznacza to jednak, że człowiek staje się przez to mniej inteligentny – takie wnioskowanie byłoby zbyt pochopne.

Historia technologii ostrzega przed nostalgią podszywającą się pod teorię poznania. Sokrates w platońskim "Fajdroście" przywołuje obawę egipskiego króla Tamusa, że pismo osłabi pamięć, gdyż ludzie zaczną ufać zewnętrznym znakom zamiast własnej zdolności przypominania. Miał rację: kultury piśmienne nie muszą trenować pamięci w taki sposób jak kultury oralne. Zarazem właśnie to osłabienie jednej sprawności stworzyło warunki do rozwoju innych: skomplikowanego prawa,

nauki kumulatywnej, archiwów, rachunku, krytyki tekstu czy filozofii prowadzonej przez stulecia. Technologia poznawcza rzadko jedynie "dodaje" zdolności; ona przekształca ich strukturę.

Dlatego pytanie o generatywną AI powinno brzmieć nie: "czy odbierze nam inteligencję?", lecz: "które zdolności będzie wzmacniać, które osłabiać i jakie nowe zależności między nimi stworzy?".

Szczególnego znaczenia nabiera tu koncepcja *desirable difficulties*, czyli trudności pożądaných. Psychologia uczenia od dawna wskazuje, że łatwość wykonania zadania nie jest tożsama z jakością nauki. Informacja odczytana ponownie może wydawać się znajoma, podczas gdy próba samodzielnego jej odtworzenia ujawnia, czy rzeczywiście została przyswojona. Gotowe objaśnienie daje poczucie płynności, natomiast samodzielne wyprowadzenie odpowiedzi zmusza do aktywizacji relacji między pojęciami. Podobny mechanizm pojawia się w *generative learning*: uczeń, który streszcza, wyjaśnia i łączy nowe treści z istniejącą wiedzą, nie jest biernym odbiorcą informacji, lecz rekonstruuje ich reprezentację. Trudność jest tu kosztowna, ale właśnie ten koszt stanowi część mechanizmu uczenia. Nie każda przeszkoda jest wartościowa; nie ma poznawczej chwały w walce z zepsutym długopisem. Istotne są te trudności, które wymuszają operację intelektualną niezbędną do zdobycia późniejszej samodzielności.

Pisanie stanowi niezwykle zbiór takich operacji. Dobór słowa wymaga rozróżnienia znaczeń. Budowanie akapitu wymusza ustalenie hierarchii twierdzeń. Redagowanie każe konfrontować zamierzony sens z rzeczywistym efektem językowym, a skracanie zmusza do wskazania tego, co naprawdę konieczne. Parafrazowanie ujawnia, czy autor rozumie cudzy argument, czy jedynie pamięta jego brzmienie. Konstrukcja kontrargumentu wymusza chwilowe opuszczenie własnej pozycji poznawczej. Nawet błąd może pełnić funkcję rozwojową, jeżeli jego rozpoznanie prowadzi do przebudowy modelu problemu.

Wartość procesu pisania leży w wysiłku poznawczym

Właśnie dlatego tekst akademicki był tradycyjnie czymś więcej niż tylko sposobem przedstawienia wiedzy nauczycielowi. Stanowił urządzeniem pedagogicznym służącym do wymuszenia na studencie ciągu niewidocznych operacji umysłowych. Końcowy esej był produktem, lecz prawdziwa wartość edukacyjna tkwiła w samym procesie jego powstawania.

Automatyzacja pisania stawia przed edukacją problem poważniejszy niż klasyczne oszustwo. Student zawsze mógł przepisać cudzą pracę, zapłacić ghostwriterowi lub skorzystać z gotowego wzoru. Generatywna AI zmienia jednak ekonomię i dostępność tych praktyk. Skraca dystans między poleceniem a gotową odpowiedzią do kilkunastu sekund, nie wymagając przy tym

współpracy drugiego człowieka. Co ważniejsze, potrafi tworzyć teksty na tyle dopasowane do konkretnego zadania, by imitować wykonanie wymaganej operacji poznawczej. Instytucja otrzymuje produkt podobny do efektu uczenia się, podczas gdy sam proces przyswajania wiedzy mógł w ogóle nie zajść. To sytuacja analogiczna do urządzenia, które nie tylko podało by wynik zadania matematycznego, ale wygenerowało również przekonujące ślady rozumowania ucznia. W takim układzie ocena produktu przestaje wystarczać do wnioskowania o kompetencjach.

Nie należy jednak utożsamiać każdej technologicznej pomocy z usunięciem pożądaney trudności. Program sprawdzający pisownię pozwala osobie z dysleksją skupić się na argumentacji, a tłumaczenie maszynowe umożliwia badaczowi uczestnictwo w anglojęzycznym obiegu naukowym bez konieczności inwestowania ogromnych zasobów w korektę idiomatyczną. Z kolei generator może przedstawić pięć kontrargumentów i w ten sposób wystawić użytkownika na intelektualny opór, którego sam by nie stworzył. Narzędzie eliminujące jedną trudność może zatem otworzyć drogę do zmierzenia się z wyzwaniem wyższego rzędu. Dlatego analiza użycia AI wymaga precyzyjnego określenia celu zadania.

Jeżeli celem jest ćwiczenie ortografii, autokorekta sabotuje proces. Jeśli jednak celem jest analiza konstytucyjna, ta sama funkcja staje się poznawczo nieistotnym udogodnieniem. Nie istnieje zatem technologia z natury „ogłupiająca” – istnieje jedynie niedopasowanie między funkcją narzędzia a kompetencją, którą chcemy rozwijać.

W kontekście pisania często powraca kwestia zapisu ręcznego. Intuicja podpowiada, że ruch dłoni, zróżnicowany kształt liter i informacje proprioceptywne tworzą bogatsze środowisko sensoryczno-ruchowe niż mechaniczne uderzanie w identyczne klawisze. Badania neurofizjologiczne rzeczywiście wskazują na różnice w synchronizacji sieci mózgowych podczas pisania ręcznego, jednak nie można przekształcić tej obserwacji w proste prawo: „pióro rozwija mózg, klawiatura go osłabia”. Aktywacja neuronalna nie jest bowiem samym w sobie dowodem lepszego uczenia; większa aktywność może oznaczać zarówno głębsze przetwarzanie, jak i po prostu wyższy koszt wykonania zadania.

Znaczenie ma także sposób korzystania z klawiatury: doświadczony autor piszący dziesięcioma palcami wykonuje inną czynność niż uczestnik eksperymentu zmuszony do pisania jednym palcem. Podobnej ostrożności wymaga teza, że „pióro jest potężniejsze od klawiatury”. Wczesne badania nad notowaniem sugerowały, że osoby korzystające z laptopów częściej przepisywały wykład dosłownie, podczas gdy notujący ręcznie musieli kondensować treść, co sprzyjało lepszemu rozumieniu pojęciowemu. Mechanizm jest logiczny: skoro ręka nie nadąża za mową, student musi selekcjonować. Jednak późniejsze replikacje nie potwierdziły uniwersalnej przewagi odręcznego zapisu w testach. Wartość pierwotnego eksperymentu nie znika, zmienia się jedynie jego

interpretacja. Kluczową zmienną nie jest materiał – papier czy komputer – lecz strategia poznawcza. Osoba bezmyślnie stenografująca wykład ręcznie może uczyć się słabo, podczas gdy student tworzący na laptopie strukturę i własne parafrazy może przetwarzać materiał bardzo głęboko.

Pióro nie posiada właściwości magicznych. Jego przewaga pojawia się tylko wtedy, gdy fizyczne ograniczenie tempa wymusza korzystną operację poznawczą. Ta uwaga jest szczególnie istotna w debacie o AI. Łatwo byłoby stworzyć prostą analogię: skoro wolniejsze pisanie ręczne zmusza do myślenia, a generowanie tekstu jest najszybsze, to proces ten musi prowadzić do największej powierzchowności.

Wartość poznawczego tarcia w procesie pisania

Taki sylogizm ignorowałby rolę projektu interakcji. System AI może być przecież zaprojektowany nie jako automat do dostarczania odpowiedzi, lecz jako sokratejski przeciwnik: narzędzie żądające uzasadnienia, wskazujące niespójności, podsuwające kontrprzykłady czy odmawiające napisania gotowego akapitu, dopóki użytkownik nie określi własnego stanowiska. W takim układzie technologia nie znosi poznawczego tarcia, lecz może je celowo konstruować. Badania dydaktyczne nad wykorzystaniem generatywnej AI przesuwają obecnie punkt ciężkości z prostego pytania „używać czy zakazywać?” w stronę projektowania sekwencji, w której student najpierw myśli, następnie korzysta z modelu, a potem sprawdza wynik, kwestionuje go i integruje z własnym tekstem.

Pojęcie tarcia jest tutaj wyjątkowo użyteczne. W ekonomii i projektowaniu interfejsów tarcie zazwyczaj traktowane jest jako koszt: dodatkowe kliknięcie, opóźnienie czy zbędny etap. Cyfrowa gospodarka została zbudowana na jego eliminacji. „One click”, autouzupełnianie, rekomendacje, bezszwowe płatności i predykcyjny tekst realizują tę samą filozofię: im mniej momentów wymagających świadomej decyzji, tym sprawniejszy system. Edukacja i twórczość mają jednak osobliwą właściwość – część ich wartości rodzi się właśnie w momencie zatrzymania. Autor, który nie wie, jak dokończyć zdanie, może doświadczać nie awarii procesu, lecz jego najbardziej produktywnego momentu. Niedopasowanie dwóch akapitów bywa sygnałem konfliktu między przekonaniami. Godzina spędzona nad definicją może być „nieefektywna” pod względem liczby słów na minutę, lecz niezwykle efektywna dla jakości rozumienia. Optymalizacja pisania wyłącznie według kryterium szybkości przypominałaby optymalizowanie siłowni przez usuwanie ciężarów. Anegdota Naomi Baron o J.V. Cunninghamie pozwala uchwycić społeczną odmianę tego samego mechanizmu.

Surowa ocena studenckiego tekstu była dla niej bolesnym doświadczeniem, ale stała się korekcyjnym oporem – komunikatem, że dotychczasowa jakość myślenia i pisania jest niewystarczająca. Należy jednak wystrzegać się romantyzowania pedagogicznej brutalności; upokorzenie samo w sobie nie stanowi metody nauczania. Istotna pozostaje jednak funkcja wymagającej informacji zwrotnej.

Dobry mentor nie zawsze wygładza drogę ucznia. Czasem odmawia zaakceptowania argumentu, którego autor nauczył się bronić jedynie retoryczną pewnością. Zmusza do powrotu do źródeł, zmiany definicji czy odróżnienia dowodu od sugestii. Asystent technologiczny podporządkowany wyłącznie satysfakcji użytkownika działa odwrotnie: obniża próg oporu, natychmiast naprawia nieporadność i tym samym ukrywa przed człowiekiem obszary, w których jego kompetencje jeszcze nie dojrzały. W tej perspektywie kluczową umiejętnością przyszłości może okazać się nie samo pisanie bez AI ani sprawne posługiwanie się modelem, lecz zarządzanie własnym wysiłkiem poznawczym. Człowiek będzie musiał nauczyć się rozpoznawać, które trudności warto delegować maszynie, a które należy zachować właśnie dlatego, że są trudne. Można bez żalu zautomatyzować mechaniczną zmianę formatowania, ale warto samemu rozstrzygnąć konflikt argumentów. Można poprosić system o wykrycie powtórzeń, lecz trzeba wiedzieć, dlaczego kluczowy akapit powinien znaleźć się przed innym. Można wykorzystać AI do wygenerowania przeciwnych stanowisk, ale to człowiek musi podjąć decyzję, które z nich uznaje za przekonujące i z jakiego powodu. Dojrzałość poznawcza w epoce sztucznej inteligencji będzie więc polegać częściowo na umiejętności odmawiania sobie wygody.

Aktywne konstruowanie znaczenia zamiast bierności

"Karta wyników" Naomi Baron przestaje w tym świetle wyglądać jak prywatna deklaracja konserwatywnego przywiązania do dawnego rzemiosła. Można ją odczytać raczej jako prototyp polityki zarządzania cognitive offloading. Pytanie nie powinno brzmieć jedynie: "czy użyłem AI?", lecz: "co dzięki niej przestałem robić samodzielnie?". Czy system usunął czynność rutynową, czy zastąpił proces odkrywania? Czy pomógł mi lepiej dostrzec własną myśl, czy dostarczył mi myśl wystarczająco elegancką, abym nie zauważył jej obcości? Czy po ukończeniu tekstu wiem więcej niż przed jego rozpoczęciem?

Czy umiem obronić argument bez ponownego pytania maszyny? Czy potrafiłbym rozpoznać, że wygenerowana odpowiedź jest błędna? Te pytania są bardziej użyteczne niż proste przeciwstawienie "naturalnego" człowieka "sztucznej" technologii, ponieważ dotyczą faktycznego rozkładu pracy poznawczej. Współczesne eksperymenty edukacyjne zaczynają prowadzić właśnie w

tę stronę. Zamiast traktować model jako automatycznego autora, próbuje się budować środowiska, w których student najpierw formułuje własny zamiar, a następnie iteracyjnie rozmawia z systemem, weryfikując odpowiedzi i zastanawiając się nad rolą AI w końcowym rezultacie. Charakterystyczne jest nawet celowe wprowadzanie "tarcia" do interfejsu: utrudnienie bezmyślnego kopiowania po to, aby człowiek musiał wykonać dodatkowy akt wyboru. To paradoks godny epoki automatyzacji. Najbardziej dojrzałym projektem technologicznym może okazać się czasami narzędzie świadomie mniej wygodne, ponieważ projektant rozumie, że użytkownik nie zawsze potrzebuje krótszej drogi do odpowiedzi. Czasem potrzebuje dłuższej drogi do zrozumienia.

Z tego punktu widzenia nie ma potrzeby czynienia z pióra relikwii oporu przeciwko sztucznej inteligencji. Znacznie ważniejsze jest zachowanie procesów, które pióro, klawiatura albo dobrze zaprojektowana interakcja z modelem mogą uruchamiać: selekcji, syntezy, rekonstrukcji, wątplenia, redagowania i samokontroli. Narzędzie jest wtórne wobec architektury poznawczej działania.

Człowiek może pisać ręcznie i nie myśleć; może pisać na komputerze i przeprowadzać niezwykle wymagające rozumowanie; może też współpracować z AI w sposób bardziej intelektualnie aktywny niż ktoś, kto samodzielnie odtwarza cudze klisze. Granica nie przebiega więc pomiędzy analogiem a cyfrą ani nawet pomiędzy człowiekiem a maszyną. Przebiega pomiędzy aktywnym konstruowaniem znaczenia a biernym przyjmowaniem znaczenia przygotowanego przez inny system.

Przyjęcie tego rozróżnienia pozwala uczciwie przejść do historii maszyn językowych. Nie powstały one bowiem z jednego nagłego odkrycia ani z osobliwej ambicji współczesnych przedsiębiorstw technologicznych.

Dążenie do uczynienia znaczenia obliczalnym

Ich rodowód wiedzie przez kryptografię, logikę matematyczną, teorię informacji i automatykę, przez tłumaczenie maszynowe, pierwsze programy konwersacyjne oraz wieloletnie próby sprowadzenia języka do operacji wykonywalnych przez maszynę. Zanim komputer nauczył się układać przekonujące akapity, musiał najpierw zostać nauczony traktowania słowa jako czegoś, co można zakodować, policzyć, rozłożyć na elementy i przekształcić według reguł. Historia tego procesu pokazuje, że dzisiejsze modele językowe nie są nagłym wtargnięciem obcej inteligencji do kultury pisma. Stanowią końcowy, choć zapewne nie ostatni etap znacznie starszego marzenia: aby znaczenie uczynić obliczalnym. To historia maszyny językowej: od Enigmy do Transformera.

Dzisiejszy użytkownik modelu językowego może odnieść wrażenie, że maszyna zaczęła pisać niemal nagle. Jeszcze niedawno komputer był urządzeniem wykonującym polecenia, przeszukującym bazy danych i poprawiającym literówki; w ciągu kilku lat zaczął jednak tworzyć eseje, streszczać rozprawy, tłumaczyć języki, komponować argumenty i prowadzić dialog sprawiający wrażenie rozmowy z kompetentnym partnerem. Historycznie jest to złudzenie skali czasu. Generatywna sztuczna inteligencja nie pojawiła się jako nieoczekiwany gość kultury pisma, lecz jest rezultatem trwającego niemal stulecie programu intelektualnego. Jego zasadniczym założeniem było przekonanie, że przynajmniej część operacji kojarzonych z językiem, myśleniem i znaczeniem można przekształcić w działania wykonywalne przez urządzenie.

Zanim komputer nauczył się generować zdania, człowiek musiał najpierw nauczyć się patrzeć na język jak na strukturę podatną na formalizację. Jednym z najbardziej dramatycznych laboratoriów tego myślenia była kryptografia. W szyfrze język zostaje celowo odarty z bezpośredniej dostępności semantycznej. Odbiorca widzi ciąg znaków, lecz nie może odczytać znaczenia, dopóki nie zrekonstruuje procedury transformacji. Praca kryptologa polega zatem na operowaniu językiem bez konieczności rozpoczynania od jego sensu. Liczą się regularności, permutacje, kombinacje, rozkłady i powtarzalność.

W tym właśnie środowisku Marian Rejewski dokonał w 1932 roku przełomu w analizie wojskowej Enigmy, wykorzystując aparat matematyczny zamiast dominujących wcześniej metod czysto lingwistycznych. Później, wraz z Jerzym Różyckim i Henrykiem Zygalskim, rozwijał metody pozwalające odzyskiwać dzienne ustawienia maszyny; przed wojną Polacy przekazali wyniki swoich prac Brytyjczykom i Francuzom. Amerykańska NSA zalicza osiągnięcie Rejewskiego do kluczowych sukcesów kryptologii, podkreślając, że jego algebraiczne podejście otworzyło drogę do późniejszego alianckiego wykorzystania Enigmy. Nie należy jednak tworzyć z tej historii wygodnego mitu, według którego Enigma „prowadziła wprost” do współczesnych modeli językowych. Genealogie technologiczne są bardziej złożone.

Znaczenie kryptografii polega raczej na przesunięciu epistemologicznym: pokazała ona, że zachowanie systemu symbolicznego można analizować matematycznie nawet wtedy, gdy jego znaczenie pozostaje chwilowo niedostępne. Jest to intuicja niezwykle ważna dla późniejszego przetwarzania języka. W obu przypadkach struktura może zostać potraktowana jako problem obliczeniowy. Nie trzeba najpierw „rozumieć” wiadomości w taki sposób, w jaki rozumie ją człowiek, aby odkrywać regularności pozwalające ją transformować. Język zaczyna być widziany równocześnie jako nośnik ludzkiego sensu i jako obiekt formalny.

Do tego przesunięcia Claude Shannon dołożył w 1948 roku teorię informacji. Jego osiągnięcie jest szczególnie istotne, ponieważ Shannon świadomie oddzielił techniczny problem komunikacji od

semantycznej treści komunikatu. W klasycznym artykule „A Mathematical Theory of Communication” interesowało go to, jak wiadomość wybrana spośród wielu możliwych może zostać skutecznie zakodowana, przesłana przez zaszumiony kanał i odtworzona po drugiej stronie. Znaczenie przekazu uznał za nieistotne dla tego konkretnego problemu inżynierskiego. Informacja otrzymywała miarę matematyczną niezależną od tego, czy przesyłany ciąg oznacza deklarację miłości, prognozę pogody czy przypadkową sekwencję znaków. To rozdzielenie informacji i znaczenia okazało się jednym z najbardziej płodnych zabiegów abstrakcyjnych XX wieku, umożliwiając powstanie infrastruktury komunikacyjnej, bez której współczesny świat cyfrowy byłby niewyobrażalny.

Rozdział między funkcjonalnym rezultatem a podmiotowością

Wprowadziło to napięcie, które powraca w dzisiejszej debacie o sztucznej inteligencji. Człowiek odbiera komunikat przede wszystkim jako znaczenie; maszyna natomiast może skutecznie operować jedynie jego reprezentacją formalną. Znaczna część historii informatyki polegała na rozszerzaniu zakresu tych operacji: od bitów do znaków, od znaków do słów, od słów do zdań, a wreszcie od zdań do relacji kontekstowych. Za każdym razem granica między manipulowaniem reprezentacją a rozumieniem znaczenia przesuwała się tak dalece, że obserwator mógł odnieść wrażenie, iż maszyna wkroczyła w kolejną domenę ludzkich zdolności.

Alan Turing nadał temu problemowi formę filozoficzną. W artykule „Computing Machinery and Intelligence” z 1950 roku, opublikowanym w czasopiśmie „Mind”, postawił pytanie o to, czy maszyny mogą myśleć. Szybko jednak uznał spór o definicje za mało produktywny i zastąpił go „grą w naśladownictwo”. To posunięcie było genialne właśnie dlatego, że omijało metafizyczne wnętrze systemu. Skoro nie możemy bezpośrednio zajrzeć do cudzego umysłu, możemy przynajmniej oceniać jego zachowanie. Turing przewidywał nawet, że wraz z rozwojem technologii język mówiący o „myśleniu” maszyn zostanie kulturowo oswojony.

Naomi S. Baron słusznie zestawia tę orientację ze stanowiskiem Geoffreya Jeffersona. Ten nie chciał przyznać maszynie statusu porównywalnego z człowiekiem tylko dlatego, że potrafiłaby ona napisać sonet. Domagał się związku między utworem a przeżyciem oraz świadomości faktu, iż dzieło zostało stworzone. Jest to w istocie różnica między kryterium rezultatu a kryterium procesu wewnętrznego – w tym sporze zakorzeniona jest niemal cała dzisiejsza filozofia generatywnej AI.

Turing pyta: czy odpowiedź działa? Jefferson odpowiada: ale co znajduje się po drugiej stronie tej odpowiedzi? Pierwszy ustanawia kryterium funkcjonalne, drugi broni związku inteligencji z podmiotowością. Ironią historii jest fakt, że zanim spór ten osiągnął pełną dojrzałość filozoficzną, komputery zaczęły już produkować teksty imitujące najbardziej ludzkie formy wypowiedzi. Christopher Strachey uruchomił na Ferranti Mark I program generujący listy miłosne (datowany na początek lat pięćdziesiątych). Mechanizm był prosty: gotowe kategorie słów, szablony składniowe i losowy wybór elementów pozwalały tworzyć niezliczone warianty wyznań. System nie posiadał modelu miłości, osoby kochanej ani relacji interpersonalnej, a mimo to produkt przyjmował formę kulturowo rozpoznawalnego aktu intymnego. To jeden z najstarszych przykładów fundamentalnego paradoksu maszynowego pisania: forma intencjonalności może pojawić się bez podmiotu posiadającego intencję, którą ta forma zdaje się wyrażać.

Listy Stracheya są czymś więcej niż technologiczną ciekawostką. Pokazują one silną tendencję człowieka do uzupełniania brakującego wnętrza systemu. Widząc formę listu miłosnego, automatycznie uruchamiamy kulturową wiedzę o zakochaniu, pragnieniu, nadawcy i adresacie. Znaczenie nie musi być zatem wytworzone wyłącznie po stronie generatora – część zostaje dostarczona przez odbiorcę.

Maszyna tworzy strukturę wystarczającą do wywołania ludzkiej interpretacji, a człowiek wykonuje resztę pracy semantycznej. Właściwość ta stała się później jednym z fundamentów efektu ELIZY.

Symulacja funkcji nie jest tożsama z procesem rozumienia

W 1956 roku, podczas letniego projektu badawczego w Dartmouth, John McCarthy, Marvin Minsky, Nathaniel Rochester i Claude Shannon zgromadzili środowisko naukowe wokół programu, który później uznano za symboliczne narodziny sztucznej inteligencji jako odrębnej dziedziny. Sam termin „artificial intelligence” pojawił się już w propozycji projektu z 1955 roku. Założenie organizatorów było niezwykle ambitne: przyjęto hipotezę, że każdy aspekt uczenia lub każda cecha inteligencji może zostać opisana na tyle precyzyjnie, by można było skonstruować maszynę zdolną do jego symulacji. Kluczowe jest tutaj słowo „symulacja”. Nie zakłada ono bowiem, że maszyna będzie odtwarzać biologiczny mechanizm ludzkiego myślenia; wystarczy, że zdoła wykonywać funkcję, którą wcześniej uznawano za znamienne dla inteligencji. To rozróżnienie między symulacją a emulacją pozostaje fundamentalne. Samolot lata, lecz nie imituje anatomii ptaka. Kalkulator wykonuje rachunki, lecz nie musi rekonstruować ludzkiej arytmetyki mentalnej.

Możliwe jest zatem, że maszyna będzie osiągać coraz większą sprawność językową bez reprodukcji tego, co człowiek nazywa doświadczeniem, świadomością albo rozumieniem. Błąd zaczyna się wtedy, gdy podobieństwo funkcji automatycznie uznajemy za dowód podobieństwa procesu. Współczesne modele generatywne czynią ten błąd szczególnie kuszącym, gdyż język jest podstawowym medium, za pomocą którego ludzie przypisują sobie nawzajem istnienie umysłu. Wczesne tłumaczenie maszynowe dobitnie ukazało przepaść między entuzjazmem technologicznym a trudnością uchwycenia znaczenia. Demonstracja Georgetown-IBM z 1954 roku, wykorzystująca komputer IBM 701 do przekładu rosyjskiego na angielski, wywołała ogromne zainteresowanie.

Z dzisiejszej perspektywy system ten był jednak prymitywny: operował na około 250 elementach słownikowych i sześciu regułach, a prezentowane zdania pochodziły ze starannie ograniczonego zbioru. Mimo to demonstracja wykreowała oczekiwanie, że pełna automatyzacja tłumaczenia jest tuż za rogiem. To pierwsza z wielu lekcji historii AI: system działający znakomicie w kontrolowanym otoczeniu może stwarzać złudzenie rozwiązania problemu ogólnego. Język jest szczególnie bezlitosny wobec takich złudzeń, ponieważ jego znaczenie zależy od kontekstu, wiedzy o świecie, pragmatyki, kultury, ironii, intencji i niezliczonych założeń niewypowiedzianych w samym zdaniu.

Raport ALPAC, opublikowany w 1966 roku przez National Academy of Sciences-National Research Council, chłodno ocenił ówczesny stan automatycznego tłumaczenia, co przyczyniło się do wygaszenia entuzjazmu wobec szybkiego rozwiązania problemu. Pierwsza fala marzeń o maszynie językowej zderzyła się z tym, co można nazwać oporem semantycznym świata.

Iluzja rozumienia w ewolucji systemów językowych

W 1966 roku Joseph Weizenbaum opublikował opis programu ELIZA. System był technicznie skromny, a jego najślynniejszy skrypt imitował psychoterapeutę, przekształcając fragmenty wypowiedzi użytkownika w pytania i odpowiedzi podtrzymujące rozmowę. ELIZA nie dysponowała współczesnym modelem językowym ani rozległą reprezentacją świata, a mimo to użytkownicy przypisywali jej zrozumienie, zainteresowanie i psychologiczną głębię. Sam Weizenbaum był zaskoczony łatwością, z jaką ludzie emocjonalnie angażowali się w relację z programem. Jego eksperyment pozostaje jedną z najważniejszych lekcji humanistyki cyfrowej: aby człowiek przypisał maszynie umysł, maszyna nie musi go posiadać; czasami wystarczy, że prawidłowo uruchomi nasze mechanizmy społeczne.

Publikacja programu w "Communications of the ACM" stała się punktem odniesienia dla całej późniejszej historii konwersacyjnej AI. Od tego momentu rozwój NLP można postrzegać jako

kolejne próby ograniczania przestrzeni, w której język wymyka się formalizacji. Systemy regułowe dążyły do jawnego zapisu wiedzy językowej. Metody statystyczne przesunęły ciężar z ręcznie tworzonych reguł na prawdopodobieństwa wyuczone z korpusów. Sieci neuronowe umożliwiły konstruowanie coraz bogatszych reprezentacji, a uczenie głębokie zwiększyło skalę i złożoność tych procesów.

Każdy etap oznaczał stopniowe odejście od idei, że programista musi precyzyjnie instruować maszynę o strukturze języka. System miał zamiast tego samodzielnie wydobywać regularności z przykładów. Przełom z 2017 roku i artykuł "Attention Is All You Need" należy rozumieć właśnie w tym kontekście. Transformer zaproponowany przez Vaswaniego i współautorów zrezygnował z dominującej wcześniej konieczności rekurencyjnego przetwarzania sekwencji na rzecz architektury opartej na mechanizmach attention. Umożliwiło to znacznie bardziej równoległe trenowanie i niezwykle skuteczne modelowanie zależności między elementami sekwencji. Pierwotny artykuł demonstrował przewagę architektury przede wszystkim w tłumaczeniu maszynowym, a nie istnienie nowej cyfrowej świadomości. Dopiero późniejsze zwiększanie skali modeli, danych i mocy obliczeniowej ujawniło, jak daleko można rozciągnąć tę architekturę.

Warto tutaj precyzyjnie rozdzielić dwa twierdzenia, które w dyskusji publicznej często zlewają się w jedno. Transformer potrafi modelować kontekst słowa z niezwykle skuteczną, lecz nie wynika z tego automatycznie, że rozumie go w takim sensie, w jakim człowiek rozumie sytuację społeczną, własne doświadczenie czy konsekwencje wypowiedzi. Kiedy model różnicuje znaczenie wyrazu "bank" występującego w otoczeniu "kredytu" od tego samego słowa w sąsiedztwie "rzeki", wykazuje imponującą kompetencję kontekstową. Jednak kompetencja operacyjna nie rozstrzyga sporu Jeffersona; nie mówi nam, czy istnieje ktoś, kto wie, czym jest bank, rzeka, kredyt albo mokra stopa zanurzona w wodzie.

W tym świetle metafora Naomi Baron o "maszynce do mielenia języka" nabiera właściwej mocy. Nie powinna być rozumiana jako próba ośmieszenia osiągnięć inżynierii – przeciwnie: maszynka jest niezwykle dobra. Potrafi przekształcać olbrzymie zbiory danych w reprezentacje umożliwiające generowanie odpowiedzi, których płynność jeszcze niedawno wydawała się domeną science fiction.

Problem pojawia się dopiero wtedy, gdy jakość produktu końcowego przesłania pytanie o naturę procesu, który go wytworzył. Przeciwwstawienie "gładkiej osłonki" nieprzejrzystemu wnętrzu systemu służy właśnie temu ostrzeżeniu: fenomenologiczna perfekcja tekstu może uczynić jego pochodzenie niewidzialnym. Droga od Enigmy do Transformera nie jest prostą historią maszyny, która stopniowo "nauczyła się myśleć". To historia coraz bardziej wyrafinowanego odkrywania, jak wiele zachowań językowych można uzyskać bez konieczności rozstrzygnięcia, czym właściwie jest ludzkie rozumienie. Kryptolog pokazał, że symbole można matematycznie rekonstruować. Shannon

dowiodł, że informację można formalizować bez odwoływania się do semantyki. Turing zaproponował badanie inteligencji poprzez zachowanie. Strachey wykazał, że nawet intymną formę językową można generować kombinatorycznie. ELIZA ujawniła gotowość człowieka do przypisywania maszynie podmiotowości.

Statystyczne NLP pokazało potęgę regularności zawartych w korpusach, a Transformer zwiększył skalę ich wykorzystania. Każde z tych osiągnięć przesunęło granicę tego, co wcześniej wydawało się wymagać udziału człowieka. Żadne z nich nie rozwiązało jednak pytania, od którego zaczęła się nasza analiza: co znaczy być autorem? Aby odpowiedzieć, trzeba przyjrzeć się chwili, w której eksperymentalna zdolność generowania języka przestała być laboratoryjną ciekawostką, a stała się ekonomiczną możliwością masowej produkcji tekstu. W tym miejscu na scenę ponownie wejdzie Roald Dahl ze swoim automatycznym gramatyzatorem.

Ekonomiczna logika automatyzacji treści i koniec niedoboru twórczego

Opowiadanie to okazuje się mniej fantazją o maszynie niż przenikliwą teorią rynku, na którym twórczość może zostać odłączona od autora, zwielokrotniona i sprzedana jak każdy inny produkt przemysłowy. A gdyby maszyny naprawdę potrafiły pisać? Dahl, Grammatizator i ekonomia przemysłowej twórczości.

Kiedy w 1953 roku Roald Dahl opublikował „The Great Automatic Grammatizator”, największą osobliwością przedstawionego świata nie była sama maszyna zdolna do produkowania opowiadań. Literatura fantastyczna знаła już automaty, cudowne urządzenia i techniczne protezy człowieka.

Dahl dostrzegł problem subtelniejszy: co stanie się z kulturą, gdy pisanie przestanie być dobrem ograniczonym czasem, cierpliwością i zdolnościami pojedynczego autora. Jego Adolph Knipe rozumie jak przedsiębiorca, a nie romantyczny wynalazca. Skoro język ma strukturę, fabuły wykazują powtarzalność, a rynek literacki nagradza rozpoznawalne formy, można potraktować opowiadanie jak produkt przemysłowy. Wystarczy przełożyć proces twórczy na procedurę i osiągnąć przewagę, której ludzki pisarz nie zneutralizuje talentem – człowiek napisze jedną książkę w czasie, w którym maszyna wyprodukuje setki.

Dahl antycypował zatem nie tyle dzisiejszą architekturę modeli językowych, co ich ekonomiczną logikę. Fikcyjny Grammatizator nie przypomina współczesnego Transformera pod względem technicznym, ale niezwykle celnie opisuje konsekwencję automatyzacji wytwarzania dóbr symbolicznych: przejście od niedoboru zdolności produkcyjnej do nadmiaru produktu. Przez

większość historii pisma zdanie stawało się tanie dopiero wtedy, gdy ktoś poniósł koszt jego wymyślenia. Druk obniżył koszt kopiowania, lecz nie tworzenia pierwotnego tekstu. Internet zbliżył koszt dystrybucji do zera, ale proces wytwarzania materiału wciąż pozostawał względnie drogi. Generatywna sztuczna inteligencja uderza właśnie w to ostatnie ograniczenie.

Automatyzuje ona nie tylko kopiowanie i przesyłanie, lecz coraz większą część samej produkcji językowej. Jest to zmiana porównywalna nie z wynalezieniem lepszego pióra, lecz z mechanizacją manufaktury. Maszyna przemysłowa nie musi pracować dokładnie tak jak człowiek, by radykalnie zmienić ekonomiczną wartość jego umiejętności; wystarczy, że wytwarza produkt spełniający wymagania odbiorcy szybciej, taniej lub w większej skali. To właśnie kryterium ekonomiczne jest dla analizy generatywnej AI ważniejsze niż filozoficzny spór o to, czy maszyna „naprawdę” tworzy. Rynek nie musi rozstrzygać dylematów Turinga i Jeffersona.

Konsument kupi tekst dlatego, że jest wystarczająco dobry, nie pytając, czy system podczas pisania cokolwiek doświadczał. Jeśli funkcjonalna substytucja osiągnie odpowiedni poziom, brak świadomości maszyny nie uchroni człowieka przed konkurencją. Tutaj ujawnia się ekonomiczna asymetria między autorem biologicznym a systemem generatywnym. Twórczość człowieka podlega ograniczeniom czasu; każdy tekst konkuruje z innymi możliwymi sposobami wykorzystania życia autora. Godzina spędzona nad jednym akapitem nie może zostać równocześnie poświęcona napisaniu dziesięciu innych. Ekonomista nazwałby to kosztem alternatywnym.

Przesunięcie wartości z produkcji tekstu na autoryzację i uwagę

Model generatywny działa inaczej. Po poniesieniu ogromnych kosztów stałych związanych z infrastrukturą, treningiem, energią, danymi i sprzętem, koszt wygenerowania kolejnej jednostki tekstu staje się relatywnie niewielki. To klasyczna cecha gospodarki cyfrowej: wysokie bariery wejścia przy bardzo niskich kosztach krańcowych reprodukcji lub świadczenia usługi. W kulturze pisma prowadzi to do gwałtownego wzrostu podaży wypowiedzi, jednak sama obfitość nie oznacza automatycznie przyrostu wiedzy. Należy tu oddzielić produkcję znaków od produkcji wartości informacyjnej. Jeśli tysiąc wariantów tekstu można stworzyć niemal tak samo łatwo jak jeden, rzadkim zasobem przestaje być tekst jako taki.

Zasobami stają się uwaga, selekcja, wiarygodność, reputacja, oryginalne dane, doświadczenie oraz zdolność rozpoznania, który z tysiąca wariantów zasługuje na zaufanie. Gospodarka generatywna może zatem prowadzić do paradoksu nadmiaru: im więcej treści jesteśmy zdolni wytworzyć, tym

bardziej kosztowne staje się ustalenie, które z nich mają znaczenie. Dawniej problemem autora było pytanie: „jak stworzyć tekst?”. Dziś coraz częściej problemem odbiorcy staje się kwestia: „dlaczego miałbym czytać właśnie ten?”. Dahlowska maszyna trafia zatem w samo serce ekonomii uwagi. Przemysłowa produkcja słowa nie znosi konkurencji, lecz zmienia jej przedmiot. Gdy koszt napisania kolejnego artykułu spada, rośnie względna wartość jego dystrybucji, pozycjonowania, marki, rekomendacji i dostępu do odbiorcy. Twórca może zostać uwolniony od części pracy kompozycyjnej, by jednocześnie znaleźć się w środowisku, w którym walka o zauważenie staje się jeszcze brutalniejsza. Paradoksalnie demokratyzacja produkcji może więc wzmacniać koncentrację dystrybucji. Kiedy każdy może wytworzyć profesjonalnie brzmiący komunikat, przewagę zdobywa ten, kto posiada sieć odbiorców, dane o ich zachowaniu, platformę rekomendacyjną albo zdolność kupowania zasięgu.

Demokracja pióra nie musi oznaczać demokracji uwagi. W opowiadaniu Dahla szczególnie przewidujący jest motyw wykupywania nazwisk pisarzy. Maszyna może produkować, lecz rynek nadal potrzebuje autoryzującego znaku człowieka. Nazwisko pełni funkcję reputacyjnego certyfikatu: mówi czytelnikowi, z jakim typem doświadczenia, estetyki albo jakości ma do czynienia.

Knipe nie musi więc niszczyć figury autora; może ją przejąć i wykorzystać jako markę nakładaną na produkcję, której autor faktycznie nie wykonał. To wizja subtelniejsza niż proste stwierdzenie, że „roboty zabiorą pracę pisarzom”. Autor może pozostać na okładce nawet wtedy, gdy jego udział w wytwarzaniu tekstu zostanie radykalnie ograniczony. Jego nazwisko staje się aktywem, a niekoniecznie opisem procesu twórczego. Współczesna gospodarka zna ten mechanizm od dawna. Wielkie przedsiębiorstwa publikują komunikaty podpisane nazwiskiem prezesa, choć tekst przygotował dział komunikacji. Polityk wygłasza przemówienie stworzone przez speechwritera. Celebryta wydaje autobiografię powstałą przy znacznym udziale ghostwritera. Profesor może kierować zespołem badawczym, nie wykonując osobiście każdego obliczenia ani eksperymentu. Sztuczna inteligencja nie tworzy więc od zera rozdziału między nazwiskiem a fizycznym procesem pisania.

Radykalizuje ona istniejącą praktykę, zwiększając skalę i zmniejszając koszt pośrednictwa. Ghostwriter jest drugim człowiekiem, któremu trzeba zapłacić, przekazać kontekst i poświęcić czas. Model jest infrastrukturą zdolną do natychmiastowej iteracji. To rozróżnienie prowadzi do ważnej korekty moralnej. Nie każda delegacja tekstu jest oszustwem, tak samo jak nie każde użycie maszyny jest utratą autorstwa. Dyrektor podpisujący profesjonalnie przygotowany raport może rzeczywiście odpowiadać za jego treść, choć nie napisał każdego zdania. Badacz pozostaje autorem artykułu przygotowanego zespołowo, jeżeli wkład, kontrola i odpowiedzialność spełniają normy

danej wspólnoty. Analogicznie człowiek korzystający z AI może zachować istotną część autorstwa, jeżeli to on określa problem, źródła, strukturę, kryteria prawdziwości i ostateczną treść. Pytanie o autorstwo nie daje się więc sprowadzić do ruchu palców po klawiaturze.

Pułapka produktywności i inflacja kompetencji pozornej

Kwestia ta dotyczy struktury kontroli nad procesem. Możliwość automatyzacji rodzi pokusę, którą trafnie uchwycił Dahl: skoro odbiorca widzi przede wszystkim rezultat, a ten można uzyskać taniej, system ekonomiczny zacznie premiować tych, którzy najszybciej przeniosą koszt z człowieka na maszynę. Presja konkurencyjna może działać niezależnie od indywidualnych preferencji twórców. Dziennikarka może cenić ręczne konstruowanie tekstu, copywriter własny styl, a prawnik samodzielne budowanie argumentacji, lecz jeśli konkurenci wykonują tę samą usługę trzy razy szybciej dzięki automatyzacji, pozostanie przy dawnym procesie staje się decyzją ekonomicznie kosztowną. Wydajność przestaje być prywatnym wyborem i staje się normą organizacyjną.

To zjawisko można opisać jako "pułapkę produktywności". Narzędzie początkowo obiecuje oszczędzić pracownikowi czas, lecz po jego upowszechnieniu organizacja niekoniecznie pozwala ten czas zachować. Zamiast napisać pięć tekstów i skończyć wcześniej, pracownik ma wykonać dziesięć.

Indywidualny zysk produktywności zostaje przekształcony w nowy standard wydajności. Historia technologii pracy zna ten mechanizm doskonale: automatyzacja zmniejsza nakład potrzebny do wykonania jednostki zadania, a następnie konkurencja zwiększa oczekiwaną liczbę tych jednostek. To, co miało być ulgą, staje się wymaganiem.

Generatywne pisanie może zatem prowadzić nie tylko do świata, w którym ludzie piszą mniej, lecz również do takiego, w którym organizacje publikują nieporównanie więcej. Badania z pierwszej fazy wdrażania generatywnej AI pokazują, dlaczego ta presja jest tak silna. W eksperymentach obejmujących zadania zawodowe osoby korzystające z modeli generatywnych pracowały szybciej, a przeciętna jakość produktu rosła; korzyści były szczególnie widoczne u uczestników, którzy bez wsparcia narzędzia osiągalni słabsze wyniki. Nie oznacza to oczywiście, że każdy rodzaj pisania zostanie równie łatwo zautomatyzowany. Pokazuje jednak mechanizm ekonomiczny o ogromnym znaczeniu: AI może kompresować różnice produktywności między wykonawcami w zadaniach, dla których dobrą odpowiedź da się wygenerować na podstawie istniejących wzorców. W takim środowisku techniczna poprawność tekstu traci wartość jako sygnał wyjątkowej kompetencji. Prowadzi to do drugiej konsekwencji, której Dahl nie mógł przewidzieć w pełnej skali: inflacji kompetencji pozornej.

Jeżeli profesjonalnie brzmiący list motywacyjny, raport, oferta, komentarz ekspercki czy analiza mogą zostać wygenerowane przez osoby o bardzo różnym poziomie rzeczywistych umiejętności, styl tekstu staje się słabszym sygnałem jakości autora. Ekonomia sygnałów uczy, że informacja posiada wartość wtedy, gdy trudno ją wiarygodnie imitować bez posiadania cechy, którą ma sygnalizować. Dyplom, portfolio czy trudny egzamin pełnią funkcję sygnału właśnie dlatego, że ich zdobycie jest kosztowne. Jeżeli poprawna proza przestaje być kosztowna, maleje jej wartość diagnostyczna. Nie oznacza to, że społeczeństwo przestanie cenić umiejętność pisania. Może wydarzyć się coś odwrotnego: wartość przesunie się od samej zdolności wytworzenia gładkiej wypowiedzi ku zdolności jej oceny, obrony, weryfikacji i twórczego różnicowania. Jeśli każdy może uzyskać poprawną średnią, większego znaczenia nabiera odchylenie od tej średniej.

Ekonomiczny paradoks AI: premiowanie unikalności przy zagrożeniu dla źródeł wiedzy

Ekonomicznie wartościowy autor będzie tym, który dostarczy perspektywę trudną do skopiowania: unikalne doświadczenie, dostęp do danych, oryginalny smak, rozpoznawalny głos, ryzyko intelektualne lub reputację zdobytą poza samym tekstem. W ten sposób AI może jednocześnie uśredniać ogromną część rynku i zwiększać premię za autentycznie odróżnialne autorstwo. Nowsze modele ekonomiczne analizujące konkurencję między treścią ludzką a generowaną wskazują właśnie na taką możliwość. Wraz ze wzrostem jakości AI rośnie substytucyjność jej produktów względem treści tworzonych przez ludzi, a tym samym presja konkurencyjna na twórców. Jedną z odpowiedzi jest silniejsza dyferencjacja: człowiek musi tworzyć coś, czego syntetyczna przeciętność nie dostarcza równie łatwo. Nie jest więc przesądzone, że obecność AI zawsze zmniejszy różnorodność kultury.

W pewnych warunkach może ona zmusić twórców do odejścia od środkowego, najbardziej przewidywalnego segmentu rynku. Problem polega jednak na tym, że nie każdy autor dysponuje zasobami pozwalającymi przetrwać taką transformację. Rynek może stać się bardziej różnorodny na krańcach, lecz jednocześnie bardziej brutalny dla średniej zawodowej.

Opowieść Dahla można czytać jako teorię twórczej klasy średniej. Największe zagrożenie nie musi dotyczyć wyłącznie najwybitniejszych pisarzy ani całkowicie rutynowych producentów tekstu. Najbardziej narażone mogą być zawody pośrednie, których wartość wynikała dotąd z umiejętności tworzenia kompetentnej, poprawnej i przewidywalnej prozy: opisów produktów, komunikatów korporacyjnych, streszczeń, standardowych analiz, materiałów marketingowych, podstawowej dokumentacji, raportów czy tłumaczeń użytkowych.

To właśnie tam różnica między „bardzo dobrym automatem” a „dostatecznie dobrym człowiekiem” najłatwiej daje się przeliczyć na czas i pieniądze. Jednocześnie gospodarka generatywna nie może całkowicie odłączyć się od ludzkiej produkcji znaczeń bez popadnięcia w osobliwy paradoks. Modele uczą się na świecie wytworzonym w ogromnej części przez ludzi. Dziennikarstwo wymaga, aby ktoś udał się na miejsce wydarzenia, zdobył dokument, porozmawiał ze świadkiem i ustalił, że wydarzyło się coś nowego. Nauka wymaga przeprowadzenia eksperymentu, pomiaru, obserwacji lub stworzenia teorii. Literatura czerpie z doświadczeń przeżywanych przez ludzi.

Jeżeli system generatywny potrafi tanio przetwarzać istniejącą wiedzę, nie oznacza to, że potrafi równie tanio wytwarzać pierwotne fakty, doświadczenia i odkrycia, z których ta wiedza się składa. W ekonomii informacji pojawia się zatem problem źródła: jeśli automatyczna synteza podcina ekonomiczne podstawy kosztownego tworzenia informacji pierwotnej, kto sfinansuje jej dalsze powstawanie?

To napięcie jest szczególnie widoczne w dziennikarstwie. Można wyobrazić sobie model generujący ogromną liczbę artykułów na podstawie dostępnych wiadomości. Jeżeli jednak nikt nie poniesie kosztu pierwotnego reporterstwa, system zacznie coraz sprawniej przetwarzać coraz starszy zasób. Syntetyczna nadbudowa potrzebuje ludzkiego podłoża. Podobny problem dotyczy nauki, sztuki, danych i wiedzy eksperckiej.

Maszyna może obniżyć koszt przekształcania istniejących zasobów, ale jej sukces ekonomiczny może jednocześnie osłabić bodźce do tworzenia zasobów nowych. Jest to jeden z głębszych paradoksów gospodarki AI: technologia korzystająca z olbrzymiego dziedzictwa ludzkiego autorstwa może zmieniać rynek w sposób utrudniający finansowanie kolejnych pokoleń tego dziedzictwa.

Funkcjonalne rozwarstwienie i dekonstrukcja figury autora

Wizja Dahla składa się z trzech poziomów. Pierwszy dotyczy technologii: język można częściowo zautomatyzować. Drugi odnosi się do rynku: automatyzacja zmienia koszty, skalę i strukturę konkurencji. Trzeci ma charakter instytucjonalny: gdy produkcja tekstu staje się tania, musimy na nowo zdefiniować to, co właściwie chcemy wynagradzać. Czy będzie to liczba słów? Rezultat? Oryginalny pomysł? Wiarygodność? Doświadczenie? Ryzyko? Odpowiedzialność? Dostarczenie nowej informacji? A może samo nazwisko?

W kulturze przedgeneratywnej elementy te tworzyły jeden pakiet zwany „autorem”. Automatyzacja rozbija ten pakiet na części i wystawia każdą z nich na osobną wycenę. Pytanie Grammatizatora nie brzmi zatem: czy maszyna potrafi napisać opowiadanie równie dobrze jak człowiek? Bardziej niepokojące jest pytanie o to, co stanie się z instytucją autorstwa, gdy rynek przestanie potrzebować człowieka do wytworzenia wszystkich elementów tekstu, lecz nadal będzie wymagał jego nazwiska, odpowiedzialności, doświadczenia lub reputacji. Możliwe, że przyszłość nie przyniesie prostego zaniku autora, lecz jego funkcjonalne rozwarstwienie. Jedni będą dostarczać doświadczeń i danych, inni nadawać kierunek, jeszcze inni reprezentować reputację, podczas gdy maszyny zajmą się formą językową i skalą.

Autor przestaje być samotnym źródłem tekstu, a staje się węzłem w systemie produkcyjnym: projektantem, kuratorem, redaktorem, selekjonerem, weryfikatorem lub gwarantem. W wariantcie optymistycznym oznacza to uwolnienie człowieka od językowej pracy niskiego poziomu i przesunięcie go ku zadaniom wymagającym sądu, wyobraźni i odpowiedzialności. W wariantcie pesymistycznym człowiek staje się jedynie biologiczną pieczęcią przyłożoną do syntetycznego produktu, ponieważ rynek nadal potrzebuje kogoś, kto podpisze się pod treścią i weźmie za nią odpowiedzialność.

Roald Dahl opisał tę możliwość na długo przed powstaniem modeli fundamentowych. Jego satyra pozostaje trafna, gdyż nie opiera się na prognozie konkretnej technologii, lecz na znajomości ludzkich instytucji. Jeśli pojawia się maszyna zdolna wytwarzać dobro szybciej i taniej, przedsiębiorca zwiększy skalę, rynek obniży wartość tego, co stało się obfite, a dotychczasowi producenci zostaną zmuszeni do renegocjowania swojej pozycji.

W przypadku generatywnej AI dobrem tym jest jednak coś niezwykłego: język – podstawowa infrastruktura ludzkiej wiedzy, reputacji, prawa, polityki i kultury. Aby zrozumieć zasięg tej przemysłowej transformacji, należy przyjrzeć się samemu procesowi generowania. Metafora "maszynki do mielenia języka" jest efektowna, lecz dopiero jej techniczna dekonstrukcja wyjaśni, dlaczego współczesne modele tworzą tak przekonujące rezultaty, mimo że spór o ich rozumienie pozostaje nierozstrzygnięty.

Zajrzyjmy zatem pod błyszczącą powierzchnię tekstu: do tokenów, reprezentacji, prawdopodobieństw, mechanizmu uwagi i ogromnych zbiorów danych, z których nowoczesna maszyna rekonstruuje statystyczną architekturę naszych sposobów mówienia. Maszynka do mielenia języka prowokuje do analizy tego, co kryje się pod eleganckim produktem. Naomi S. Baron wykorzystuje ją jako ostrzeżenie przed sytuacją, w której odbiorca widzi jedynie płynny, gramatycznie uporządkowany tekst i na tej podstawie zakłada istnienie kompetentnego autora. Jest to jednak metafora, a nie techniczny opis modelu językowego. Dosłowne potraktowanie tego obrazu

mogłoby prowadzić do błędnego przekonania, że system generatywny działa jak wyszukiwarka fragmentów zdań, która wycina kawałki istniejących tekstów i skleja z nich nową wypowiedź. Mechanizm jest znacznie ciekawszy.

W procesie treningu sieć uczy się wielowymiarowych regularności obecnych w języku, a następnie wykorzystuje je do konstruowania kolejnych elementów wypowiedzi. To właśnie zdolność generalizacji, a nie zwykłe kopiowanie, pozwala modelom odpowiadać na pytania sformułowane w sposób, który nie musiał wcześniej wystąpić w danych treningowych. Pierwszą niespodzianką dla intuicji humanisty jest fakt, że model zazwyczaj nie zaczyna od „słów” w takim sensie, w jakim doświadcza ich człowiek.

Matematyczna reprezentacja tekstu nie jest rozumieniem

Tekst zostaje rozłożony na tokeny – jednostki, które mogą odpowiadać całym słowom, ich fragmentom, znakom lub innym często występującym sekwencjom. Tokenizacja wydaje się jedynie technicznym etapem poprzedzającym właściwe modelowanie, lecz coraz wyraźniej widać, że wpływa ona na efektywność, wielojęzyczność, długość kontekstu oraz sposób reprezentowania struktur językowych. Badania z 2025 i 2026 roku wskazują wręcz, że wybór tokenizera należy traktować jako element architektury modelu, a nie neutralną czynność przygotowania danych. Podczas gdy człowiek widzi zdanie jako strukturę semantyczną, system na pierwszym poziomie otrzymuje sekwencję identyfikatorów, którą dopiero przekształca w reprezentacje matematyczne.

Każdemu tokenowi przypisywana jest reprezentacja wektorowa. Nie jest to jednak prosty słownik, w którym jednej jednostce odpowiada gotowa definicja. Wektor stanowi punkt w wielowymiarowej przestrzeni, którego położenie zostało wyuczone poprzez analizę ogromnej liczby zależności w danych. Elementy występujące w podobnych środowiskach językowych uzyskują reprezentacje pozwalające modelowi rozpoznawać między nimi podobieństwa i różnice.

Nie oznacza to jednak, że gdzieś wewnątrz sieci znajduje się elegancka kartoteka zatytułowana „znaczenie demokracji”, „pojęcie śmierci” albo „uczucie zazdrości”. Wiedza jest rozproszona pomiędzy ogromną liczbą parametrów i aktywacji. Już na tym poziomie metafora biblioteki zaczyna zawodzić. Model nie przechowuje tekstów tak, jak biblioteka przechowuje książki; uczy się raczej struktury pozwalającej przewidywać i rekonstruować regularności obecne w języku.

Mechanizm attention, fundament architektury Transformera, jeszcze bardziej komplikuje ten obraz. W klasycznych modelach sekwencyjnych informacje musiały przepływać przez kolejne stany

zgodnie z następstwem elementów. Transformer zaproponowany w 2017 roku umożliwił bezpośrednie modelowanie zależności pomiędzy elementami sekwencji poprzez mechanizmy uwagi, rezygnując w pierwotnej architekturze z rekurencji i konwolucji. Autorzy wykazali jego przewagę w zadaniach tłumaczenia oraz znacznie większą możliwość równoległego trenowania.

Osiągnięcie to nie polegało jednak na wszczępieniu maszynie ludzkiej uwagi psychologicznej. „Attention” jest nazwą operacji matematycznej, która pozwala systemowi dynamicznie ważyć relacje między elementami reprezentacji. Gdy w zdaniu występuje słowo wieloznaczne, człowiek szybko rozstrzyga jego sens, wykorzystując kontekst. Transformer również uwzględnia zależności między tym słowem a innymi elementami wypowiedzi, tworząc reprezentację dostosowaną do konkretnego otoczenia. Przykładowo: sąsiedztwo słów „kredyt” i „pieniądze” kieruje system w inną stronę niż obecność „rzeki” i „ryby” przy analizie słowa „bank”.

Należy jednak zachować ostrożność przy używaniu słowa „rozumie”. Technicznie można stwierdzić, że model skutecznie reprezentuje i wykorzystuje zależności kontekstowe. Filozoficznie zaś twierdzenie, że system rozumie bank tak jak człowiek stojący przed okienkiem kredytowym albo na brzegu Wisły, wymagałoby dodatkowej teorii znaczenia. Właśnie tę różnicę podkreślili Emily Bender i Alexander Koller, przeciwstawiając językową formę znaczeniu zakotwiczonemu w intencji komunikacyjnej i świecie. Ich argument nie neguje osiągnięć modeli językowych; przeciwnie – uznaje ich imponującą skuteczność, ale ostrzega przed nieuprawnionym przejściem od stwierdzenia, że „system wykonuje zadanie wrażliwe na znaczenie”, do wniosku, że „system posiada znaczenie w sensie analogicznym do człowieka”. To rozróżnienie jest kluczowe. Termometr reaguje na temperaturę, lecz nie przypisujemy mu doświadczenia gorąca. Aparat fotograficzny rozróżnia parametry światła, nie przeżywając koloru.

Skala parametrów generuje kompetencje bez intencjonalnego rozumienia

Model językowy potrafi operować niezwykle bogatymi regularnościami semantycznymi, nie musząc rozstrzygać, czy posiada fenomenologiczne lub intencjonalne rozumienie treści, o których mówi. Modele generatywne typu autoregresyjnego wprowadzają do tego procesu dodatkowy mechanizm: w najprostszym ujęciu uczą się one przewidywać następny token na podstawie poprzedniego kontekstu. Pozornie brzmi to banalnie. Jeśli jednak zadanie to jest realizowane na gigantycznej liczbie przykładów przez sieć o ogromnej liczbie parametrów, system musi przyswoić regularności składniowe, semantyczne, stylistyczne i pragmatyczne, które pomagają przewidzieć dalszy ciąg. Skala zaczyna rodzić kompetencje, których nie zapisano w modelu jako ręcznie skonstruowanych

reguł. GPT-3 z 2020 roku, dysponujący 175 miliardami parametrów, demonstrował zdolność wykonywania wielu zadań na podstawie instrukcji i niewielkiej liczby przykładów podanych w kontekście, bez konieczności klasycznego dostrajania parametrów do każdego z nich. To właśnie takie zjawiska gwałtownie przesunęły nasze wyobrażenia o granicach modeli językowych. Sformułowanie „przewiduje następny token” jest jednak równie podatne na ideologiczne nadużycie, co stwierdzenie, że model „rozumie język”. Pierwsze może służyć do bagatelizowania technologii, drugie zaś do jej antropomorfizacji. Oba opisy bywają niewystarczające.

Człowieka również można opisać na różnych poziomach abstrakcji: można powiedzieć, że porusza mięśniami krtani, wypowiada zdanie, argumentuje lub próbuje kogoś przekonać. Żaden z tych opisów sam w sobie nie wyczerpuje zjawiska. Podobnie fakt, że model generatywny optymalizuje przewidywanie elementów sekwencji, nie oznacza, że jego wewnętrzne reprezentacje są banalne. Trening może prowadzić do powstania złożonych struktur pozwalających na realizację zadań, których projektanci nie zaprogramowali wprost. Nie wolno więc krytykować modeli poprzez udawanie, że ich sukces jest pozorem. To właśnie realność tego sukcesu czyni problem autorstwa interesującym. Pojęcie "stochastycznej papugi", wprowadzone do debaty przez Emily Bender, Timnit Gebru, Angelinę McMillan-Major i Margaret Mitchell, należy odczytywać w tym samym duchu. Nie jest to twierdzenie, że model losowo powtarza zasłyszane zdania. Autorzy zwracali uwagę na znacznie szerszy problem: ogromne modele uczące się z masowych zbiorów danych mogą produkować niezwykle przekonujące formy językowe, podczas gdy dane wejściowe niosą ze sobą nierówności społeczne, stereotypy i dominację określonych języków czy grup, a koszt środowiskowy i finansowy skalowania nie jest neutralny. Pytanie brzmiało zatem nie tylko o to, „czy model działa?”, ale również: „co zostaje utrwalone, wykluczone i ukryte w procesie doprowadzania do tego, że działa?”.

W tym miejscu powraca zasada GIGO – garbage in, garbage out. Jest ona potrzebna, lecz wymaga korekty. W klasycznym systemie deterministycznym wadliwe dane prowadzą wprost do wadliwych wyników. W modelu uczącym się relacja ta jest bardziej złożona. Sieć potrafi generalizować, uśredniać i wykrywać wzorce, a procedury filtrowania, dostrajania oraz uczenia na preferencjach mogą znacząco modyfikować ostateczne zachowanie modelu.

Kultura jako surowiec i ryzyko konfabulacji

"Śmieci na wejściu" nie przechodzą przez system jak kamień przez rurę. Stanowią element statystycznego środowiska, z którego model wywodzi regularności. Zanieczyszczone dane

pozostawiają jednak strukturalny ślad: utrwalają stereotypy, błędne korelacje oraz nadreprezentację jednych perspektyw kosztem innych.

Pytanie o dane treningowe ma zatem wymiar epistemologiczny i polityczny. Internet nie jest neutralną próbką ludzkiej wiedzy, lecz archiwum tego, co określone społeczeństwa miały możliwość, motywację i władzę opublikować oraz cyfrowo utrwalić. Obok encyklopedii i teorii naukowej mieści się w nim propaganda, pornografia, marketing, błędy, trolling czy spam; dyskusje ekspertów przeplatają się z wypowiedziami osób całkowicie niekompetentnych. Niektóre języki dysponują miliardami cyfrowych słów, inne jedynie szczątkową reprezentacją. Podobnie jest z grupami społecznymi – jedne pozostawiły po sobie rozległe archiwa, inne były historycznie wypychane poza instytucje wytwarzania tekstu.

Model uczący się z takich zasobów nie spogląda na ludzkość z zewnątrz; widzi raczej odbicie tej jej części, która znalazła się w danych. Sam termin "dane" bywa tu zwodniczy. Fragment powieści był dla pisarza dziełem, wpis na forum elementem rozmowy, artykuł naukowy rezultatem badań, a fotografia dokumentem życia. Z perspektywy uczenia maszynowego wszystkie te wytwory stają się przykładami służącymi do wyprowadzania regularności. Następuje radykalna zmiana ontologiczna: kultura staje się zbiorem treningowym. To, co człowiek stworzył jako komunikat o konkretnym celu, dla maszyny jest materiałem do optymalizacji funkcji celu. Stąd bierze się jeden z kluczowych konfliktów współczesnej gospodarki generatywnej: kultura może być jednocześnie ekspresją osób oraz surowcem infrastruktury obliczeniowej.

Warto przy tym odróżnić wiedzę zakodowaną w parametrach od klasycznej bazy danych. Model nie przechowuje twierdzeń w uporządkowanej tabeli z polami takimi jak "prawda", "źródło" czy "data aktualizacji". Ta różnica jest kluczowa dla zrozumienia problemu halucynacji. Jeśli system ma wygenerować ciąg odpowiadający statystycznie i pragmatycznie kontekstowi, może stworzyć wypowiedź idealnie pasującą do formy oczekiwanej przez użytkownika, mimo że jej szczegóły są fałszywe. W tradycyjnej bazie wiedzy brak rekordu skutkowałby komunikatem "brak danych".

Model generatywny został jednak skonstruowany właśnie po to, aby generować. Zdolność do płynnego uzupełniania brakującej struktury jest źródłem jego imponującej użyteczności, ale i mechanizmem prowadzącym do konfabulacji. Wyjaśnia to paradoks trudny dla ludzkiej intuicji: model może być jednocześnie znakomity w objaśnianiu skomplikowanego zagadnienia i zdolny do wymyślenia nieistniejącego źródła, nazwiska lub faktu.

Asymetria między jakością retoryczną a gwarancją epistemiczną

Nie jest to koniecznie „kłamstwo”, gdyż kłamstwo w klasycznym sensie wymaga rozpoznania prawdy oraz intencji wprowadzenia odbiorcy w błąd. Bardziej adekwatne jest pojęcie zawodności epistemicznej generacji: mechanizm odpowiadający za płynność nie jest sam w sobie gwarantem prawdziwości. Współczesne systemy próbują niwelować tę różnicę poprzez dodatkowy trening, wyszukiwanie informacji, narzędzia, zewnętrzne bazy wiedzy oraz procedury weryfikacyjne, lecz problem ten pozostaje kluczowy dla teorii autorstwa. Autor ludzki również może się mylić; oczekujemy jednak, że jest on w stanie wskazać źródło swojego twierdzenia i uzasadnić, dlaczego uznał je za prawdziwe.

Tu pojawia się kwestia odpowiedzialności epistemicznej. Człowiek piszący artykuł naukowy nie powinien być oceniany wyłącznie przez pryzmat prawdopodobieństwa językowego swoich zdań. Musi on odróżniać informację potwierdzoną od hipotezy, fakt od interpretacji, źródło pierwotne od wtórnego oraz korelację od przyczynowości. Model może wspierać wszystkie te operacje, ale końcowa wiarygodność tekstu wymaga procedury, która nie sprowadza się do generowania najbardziej przekonującej kontynuacji. Im poważniejsze konsekwencje wypowiedzi, tym bardziej konieczne staje się rozdzielenie funkcji generatora od funkcji arbitra prawdy.

Paradoksalnie rosnąca sprawność modeli może ten problem pogłębiać. Słaby system sygnalizuje własne ograniczenia niezgrabnością – tekst pełen błędów składniowych budzi czujność czytelnika. Dobry model usuwa właśnie te sygnały ostrzegawcze. Potrafi przedstawić fałsz w formie, którą kulturowo kojarzymy z kompetencją: spokojnie, systematycznie, terminologicznie poprawnie i z logicznie brzmiącymi przejściami. Baronowska „lśniaca osłonka” jest więc metaforą asymetrii między jakością retoryczną a gwarancją epistemiczną. Im piękniejsza osłonka, tym większa potrzeba wiedzy o tym, co znajduje się pod nią.

Nie można jednak zakończyć tej analizy stwierdzeniem, że skoro model nie posiada ludzkiego doświadczenia, jego reprezentacje są poznawczo bezwartościowe. Byłoby to odwrócenie tego samego antropomorficznego błędu. W języku zakodowana jest ogromna ilość wiedzy o świecie właśnie dlatego, że ludzie używali go przez stulecia do jego opisywania. Statystyczne regularności języka nie są czystą powierzchnią. Słowa występują w określonych związkach, ponieważ pewne rzeczy rzeczywiście współwystępują, praktyki społeczne mają określoną strukturę, a pojęcia należą do wspólnych domen. Modelowanie języka może więc prowadzić pośrednio do reprezentacji odzwierciedlających realne regularności świata. Otwarte pozostaje pytanie, czy tak powstałą strukturę należy nazywać „rozumieniem”, czy raczej wyrafinowaną kompetencją reprezentacyjną.

Warto zatem odrzucić zarówno mistycyzm technologiczny, jak i jego lustrzane odbicie. Model nie staje się człowiekiem tylko dlatego, że stworzył piękny akapit. Z drugiej strony człowiek nie dowodzi jego banalności, powtarzając, że system „tylko przewiduje tokeny”. Słowo „tylko” wykonuje tu więcej pracy retorycznej niż naukowej.

Istotne pytanie brzmi: jakie struktury trzeba wykształcić, aby skutecznie przewidywać język na poziomie osiąganym przez współczesne systemy i które z nich funkcjonalnie odpowiadają zdolnościom dotąd kojarzonym z rozumowaniem, pamięcią, abstrahowaniem i generalizacją. „Maszynka do mielenia języka” powinna być więc rozumiana jako urządzenie epistemicznie podwójne. Z jednej strony ujawnia potęgę statystycznego przetwarzania kultury: z olbrzymiej przestrzeni przykładów można nauczyć się struktur pozwalających tworzyć nowe, zadziwiająco kompetentne wypowiedzi. Z drugiej strony zaciera pochodzenie i status tych treści. Czytelnik widzi produkt, nie widząc zbioru danych, procesu treningu, parametrów, filtrów, kryteriów optymalizacji czy procedur post-treningu. W klasycznym autorstwie pytaliśmy o biografię człowieka; w autorstwie algorytmicznym potrzebujemy czegoś na kształt biografii systemu. Dopiero z tej perspektywy można zrozumieć kolejną fazę historii.

Maszyny przestały być jedynie laboratoriami do badania języka. Zaczęły pisać wiadomości agencyjne, raporty finansowe, prognozy, książki, tłumaczenia, dokumenty prawne, reklamy i teksty literackie. „Maszynka” wyszła z laboratorium i wkroczyła w obszary zawodów, w których słowo było dotąd podstawowym świadectwem ludzkiej kompetencji. Wraz z nią pojawiło się pytanie znacznie bardziej praktyczne niż spór o naturę świadomości: skoro system potrafi wytwarzać tekst, który wykonuje pracę, do czego właściwie potrzebny jest jeszcze zawodowy autor?

Przejście od automatycznego raportu do zawodu post-edytora zmienia charakter całego sporu o autorstwo. Dopóki maszyna generuje eksperymentalny list miłosny albo literacką osobliwość, można traktować ją jako ciekawostkę z historii informatyki. Gdy jednak jej produkt zaczyna zastępować raport finansowy, wiadomość prasową, tłumaczenie, fragment pisma procesowego albo streszczenie literatury naukowej, pytanie „czy maszyna naprawdę rozumie?” zostaje uzupełnione przez pytanie znacznie bardziej brutalne: czy jej tekst wykonuje pracę, za którą dotąd płacono człowiekowi? Rynek zawodowy nie jest laboratorium filozofii umysłu. Jeżeli wytwór spełnia standard funkcjonalny, zostaje włączony w proces organizacyjny nawet wtedy, gdy problem jego intencjonalności pozostaje całkowicie nierozstrzygnięty.

Automatyzacja pisania przekroczyła tym samym granicę pomiędzy demonstracją technologiczną a zmianą instytucjonalną. Źródło przedstawia ten proces jako przejście „od listów miłosnych Stracheya do Beta Writer”, po którym maszyny wkraczają do dziennikarstwa, prawa i tłumaczeń, a

człowiek może zostać przesunięty z roli autora do roli post-edytora. Jest to trafna oś interpretacyjna, wymagająca jednak jednego ważnego doprecyzowania.

Automatyzacja dekomponuje zawód na zadania

Historia automatyzacji pisania zawodowego nie polega na nagłym pojawieniu się maszyny zdolnej zastąpić cały zawód. Znacznie częściej dochodzi do dekompozycji zawodu na zadania. Automatyzowane są te fragmenty pracy, które charakteryzują się regularnym wejściem, powtarzalną strukturą i mierzalnym rezultatem. Człowiek pozostaje tam, gdzie niezbędne są niepełne informacje, odpowiedzialność, negocjowanie znaczeń, ocena konsekwencji, pozyskiwanie nowych danych lub przekraczanie schematu. Technologia nie tyle "pożera zawód", ile najpierw rozcina go na części. Widać to wyjątkowo wyraźnie w dziennikarstwie finansowym.

Associated Press już w połowie poprzedniej dekady wykorzystała automatyzację do przygotowywania krótkich informacji o wynikach kwartalnych amerykańskich spółek. Wcześniej redakcja tworzyła ręcznie około trzystu takich materiałów na kwartał; dzięki rozwiązaniu Automated Insights i danym Zacks Investment Research skala ta wzrosła do kilku tysięcy tekstów. AP przedstawiało ten cel nie jako zastąpienie dziennikarstwa, lecz przesunięcie reporterów od mechanicznego przepisywania danych ku analizie, rozmowom ze źródłami i poszukiwaniu tematów wymagających rzeczywistego warsztatu. Przypadek ten jest modelowy, gdyż ukazuje dwa możliwe odczytania tej samej technologii. W perspektywie optymistycznej maszyna uwalnia dziennikarza od rutyny. Z perspektywy organizacyjnie bardziej sceptycznej udowadnia jednocześnie, że znaczna część pracy, za którą wcześniej płacono człowiekowi, mogła zostać opisana jako powtarzalna procedura. Podobny eksperyment przeprowadził "The Washington Post", wdrażając w 2016 roku system Heliograf.

Początkowo generował on krótkie aktualizacje dotyczące igrzysk olimpijskich, później posłużył do relacjonowania niemal pięciuset wyścigów wyborczych. Istotny jest jednak szczegół, który w uproszczonej historii "robotów piszących wiadomości" łatwo przeoczyć. Redakcja przedstawiała Heliografa jako system hybrydowy: automat aktualizował dane i tworzył fragmenty tekstu, a dziennikarze mogli dodawać analizę, kontekst, reportaż lub całkowicie nadpisywać materiał generowany maszynowo. Już ten wczesny przykład ukazuje architekturę, która stała się później podstawowym modelem współpracy człowieka z generatywną AI: system szybko wykonuje to, co powtarzalne, a człowiek zachowuje odpowiedzialność za to, czego nie można sprowadzić do aktualizacji szablonu. Nie jest przypadkiem, że automatyzacja najwcześniej osiągnęła wysoką skuteczność tam, gdzie język posiada wyraźną strukturę. Wynik meczu, notowanie giełdowe czy

raport pogodowy mają dane wejściowe o stabilnej formie; tekst można w nich zbudować według przewidywalnej zależności między liczbami a zdaniami. W tym kontekście warto przywołać wcześniejsze systemy automatycznych streszczeń oraz specjalistyczne zastosowania, w których copywriter zaczyna przypominać "Copy Directora" zarządzającego procesem generowania, a nie rzemieślnika konstruującego każdą frazę. Ekonomicznie jest to istotne rozróżnienie.

Automatyzacja nie musi zastępować wysokiej kreatywności, aby zmienić rynek pracy. Wystarczy, że przejmie warstwę tekstów "wystarczająco standardowych", z których utrzymywała się znaczna część profesjonalnych autorów. Przypadek Beta Writer pokazuje podobną logikę w wydawnictwie naukowym, a zarazem przestrzega przed przesadzaniem w opisach osiągnięć maszyny. W 2019 roku Springer Nature opublikował "Lithium-Ion Batteries: A Machine-Generated Summary of Current Research", określając publikację jako pierwszy prototyp książki wygenerowanej maszynowo. System opracowany we współpracy z Goethe University Frankfurt grupował recenzowane artykuły, organizował je w rozdziały i tworzył automatyczne streszczenia oraz elementy orientacyjne. Nie był to zatem autonomiczny badacz dokonujący nowych eksperymentów czy tworzący teorię; był to zaawansowany mechanizm syntezy istniejącego korpusu.

Maszyna potrafiła przekształcić istniejącą literaturę w użyteczny artefakt wydawniczy, lecz źródłowa wiedza naukowa została wcześniej wytworzona przez ludzi. Springer Nature poszedł później w stronę modelu jeszcze wyraźniej hybrydowego, łącząc automatyczne przeglądy literatury z ludzkim komentarzem naukowym. Rozwiązanie to jest interesujące nie jako kompromis między "człowiekiem" a "robotem", lecz jako rozdzielenie dwóch odmiennych rodzajów wartości epistemicznej.

System jest niezwykle efektywny w przeszukiwaniu, grupowaniu i kondensowaniu dużych zbiorów dokumentów. Badacz powinien natomiast odpowiadać za znaczenie syntetyzowanego materiału, ocenę jego jakości, dostrzeżenie sporów, ograniczeń metodologicznych i wskazanie tego, co w danym polu naprawdę zasługuje na uwagę. Automat potrafi pomóc zobaczyć las. Nie wynika z tego jednak, że wie, które drzewo jest chore ani dlaczego warto je badać. Jeszcze bardziej radykalny eksperyment przeprowadził Philip M. Parker z INSEAD, którego opatentowana metoda automatycznego tworzenia i marketingu publikacji pozwalała generować książki oraz raporty na podstawie ustrukturyzowanych informacji. W 2007 roku INSEAD informował o około dwustu tysiącach wygenerowanych tytułów. W rozmowie z 2026 roku Parker stwierdził, że liczba ta przekroczyła 1,6 miliona.

Automatyzacja procesu a paradoks drabiny kompetencyjnej

Skala tego zjawiska robi wrażenie, ale to właśnie ona, a nie literacka jakość poszczególnych książek, jest tu kluczowa. Parker pokazał, że po właściwym sformalizowaniu procesu wydawniczego produkcja tytułu może zostać potraktowana jak generowanie wariantu z przestrzeni danych. W tej logice milionowa bibliografia nie świadczy o milionie lat doświadczenia autorskiego, lecz o milionie iteracji systemu. Tu ponownie powraca problem sygnałów. Przez stulecia liczba książek napisanych przez jedną osobę była traktowana jako – choć niedoskonały – wskaźnik ogromnego nakładu czasu. Automatyzacja zrywa ten związek. Liczba tekstów przestaje informować o czasie poznawczym potrzebnym do ich wytworzenia. W świecie publikacji maszynowej można stworzyć obszerną bibliotekę bez odpowiadającej jej liczby lat pracy autora. To samo dzieje się dziś na mniejszą skalę w marketingu, komunikacji korporacyjnej i mediach społecznościowych.

Portfolio tekstowe nie musi już sygnalizować analogicznego portfolio doświadczenia. Tym większego znaczenia nabiera pytanie o metodę powstawania dokumentu. Prawo jest tu szczególnie interesującym laboratorium transformacji, ponieważ język stanowi w nim jednocześnie narzędzie pracy i instrument wywoływania skutków instytucjonalnych. Pismo procesowe nie jest tylko tekstem – może uruchomić postępowanie, przyznać lub podważyć roszczenie, zinterpretować normę i wpłynąć na rozkład praw oraz obowiązków. Dlatego automatyzacja prawna obnaża ograniczenia prostego kryterium „dobrze brzmi”. Narzędzia takie jak LegalMation deklarują automatyczne przygotowywanie odpowiedzi procesowych, materiałów discovery czy streszczeń na podstawie wpływających pism i reguł danej jurysdykcji – przy założeniu nadzoru prawników. Właśnie te ostatnie słowa są najważniejsze: „przy nadzorze prawników”. Odpowiedzialność w prawie nie daje się bowiem delegować równie łatwo, co wygenerowanie akapitu.

Można zautomatyzować wyszukiwanie wzorca, stworzenie pierwszej wersji dokumentu, porównanie klauzul czy streszczenie materiału. Ktoś jednak musi rozstrzygnąć, czy argument jest dopuszczalny, czy źródło istnieje, czy precedens pozostaje aktualny i czy nie pominięto stanu faktycznego, którego model nie rozpoznał. Musi też ocenić, czy podpisanie pisma jest zgodne z interesem klienta i obowiązkami zawodowymi. Generatywna AI przesuwą więc prawnika z pozycji autora pierwszej wersji ku roli walidatora, stratega i gwaranta. Nie jest oczywiste, że rola ta jest mniej wymagająca; jest natomiast inna. Dane rynkowe pokazują, że przejście to nie jest już marginalnym eksperymentem.

W badaniu Thomson Reuters z 2025 roku odnotowano wyraźny wzrost wykorzystania generatywnej AI w organizacjach prawniczych względem roku poprzedniego; raport wskazuje

również, że większość profesjonalistów spodziewa się istotnego miejsca tej technologii w przyszłych przepływach pracy. Najczęściej wymieniane zastosowania dotyczą właśnie czynności tekstowych: researchu, streszczania czy przygotowywania pierwszych wersji dokumentów – etapów poprzedzających końcową decyzję eksperta. To kluczowe, ponieważ przemiana zawodu dokonuje się nie wtedy, gdy ostatni prawnik zostanie zastąpiony, lecz znacznie wcześniej: gdy zmienia się struktura tego, za co klient płaci i czego junior musi nauczyć się samodzielnie. Tutaj pojawia się problem, który można nazwać paradoksem drabiny kompetencyjnej. Zawody eksperckie tradycyjnie uczą poprzez wykonywanie prostszych zadań. Młody prawnik przegląda dokumenty, pisze pierwsze wersje pism i wyszukuje orzeczenia, stopniowo budując intuicję niezbędną do rozwiązywania problemów strategicznych. Młoda dziennikarka przepisuje dane i tworzy krótkie depesze, ucząc się weryfikacji źródeł, zanim otrzyma wielomiesięczne śledztwo. Tłumacz zaczyna od mniej skomplikowanych materiałów, zanim przejmie teksty wymagające szczególnej odpowiedzialności.

Jeżeli automatyzacja usuwa z rynku zadania juniorskie, pojawia się pytanie: w jaki sposób człowiek zdobędzie doświadczenie potrzebne do nadzorowania maszyny, skoro właśnie te zadania, na których dawniej budował warsztat, wykonuje teraz automat? Nie jest to argument przeciw automatyzacji, lecz przeciw naiwnemu przekonaniu, że wszystkie oszczędzone godziny są czystym zyskiem. Instytucja może dziś dysponować seniorem wyposażonym w niezwykle wydajne narzędzia, ale za dziesięć lat będzie potrzebowała nowego seniora. Jeśli ścieżka prowadząca do jego kompetencji została zautomatyzowana, konieczne jest stworzenie alternatywnego systemu kształcenia. W przeciwnym razie organizacja może popaść w osobliwą zależność: coraz mniej osób potrafi wykonać zadanie samodzielnie, choć coraz więcej potrafi zatwierdzić wynik maszyny. Dopóki system działa sprawnie, różnica pozostaje niewidoczna; staje się krytyczna dopiero wtedy, gdy trzeba rozpoznać subtelny błąd.

Tłumaczenie profesjonalne obrazuje tę zmianę jeszcze wyraźniej, ponieważ proces post-editingu maszynowego stał się odrębną praktyką zawodową na długo przed eksplozją generatywnej AI. Tłumacz otrzymuje wersję wytworzoną automatycznie, a następnie poprawia błędy i dostosowuje terminologię oraz styl. Notatka źródłowa interpretuje ten model krytycznie, wskazując ryzyko „translatorese”, spłaszczenia języka i przekształcenia tłumacza w post-edytora. Warto jednak uniknąć zbyt łatwego wartościowania.

Badania nad post-editingiem pokazują, że narzędzia te mogą zmniejszać określone rodzaje wysiłku poznawczego i poprawiać produktywność, choć efekty nie są jednokierunkowe i zależą od jakości wejścia oraz środowiska pracy. Przykładowo badanie profesjonalnych tłumaczy z 2024 roku analizujące post-editing z wykorzystaniem dodatkowej syntezy mowy wykazało obniżenie pewnych

miar wysiłku poznawczego, ale równocześnie pogorszenie produktywności w tej konkretnej konfiguracji. Człowiek i automat nie układają się więc w prostą linię: „więcej AI = mniej pracy”.

Wartość tekstu zależy od funkcji i procesu

Kluczowa jest jednak zmiana jakości samej uwagi. Tłumacz tworzący tekst od podstaw musi konstruować rozwiązanie; post-edytor zaczyna od propozycji, która wyznacza punkt odniesienia. Psychologia decyzji opisuje tu efekt zakotwiczenia: pierwsza przedstawiona wartość wpływa na kolejne sądy, nawet jeśli wiemy, że jest ona jedynie propozycją. Analogicznie, wersja wygenerowana automatycznie może narzucać sposób widzenia zdania. Pojawia się nowe zadanie poznawcze: nie wymyślenie odpowiedzi, lecz rozpoznanie, kiedy gotowy tekst jest wystarczająco dobry, a kiedy jego elegancja maskuje błąd. To kompetencja niezwykle potrzebna, choć odmienna od umiejętności wytwarzania tekstu od zera.

Pojęcia „post-edytora” nie należy więc traktować wyłącznie jako obelgi wobec zdegradowanego autora. Można wyobrazić sobie wysoką formę post-edycji, przypominającą pracę redaktora naukowego, sędziego jakości czy architekta argumentacji. Problem pojawia się wtedy, gdy człowiek staje się post-edytorem bez czasu, wiedzy lub uprawnień do rzeczywistej oceny. Jeśli organizacja oczekuje jedynie szybkiego zatwierdzania treści z systemu, nie mamy do czynienia z rozszerzeniem profesjonalizmu, lecz z automatyzacją pierwszego rzędu i rytualnym nadzorem drugiego rzędu. Człowiek pozostaje formalnie odpowiedzialny, choć coraz mniejsza część procesu znajduje się pod jego realną kontrolą. To szczególnie niebezpieczna konfiguracja, gdyż pozwala instytucji korzystać z wydajności maszyny, zachowując jednocześnie człowieka jako bufor odpowiedzialności. Algorytm generuje, pracownik podpisuje. Gdy wszystko działa, sukces przypisuje się automatyzacji.

Gdy jednak występuje błąd, pytanie brzmi: dlaczego człowiek go nie zauważył? Powstaje asymetria charakterystyczna dla wielu systemów automatycznych: maszyna przejmuje sprawstwo operacyjne, ale człowiek zachowuje odpowiedzialność normatywną. W lotnictwie, medycynie, finansach i administracji problem ten jest znany od dawna; generatywne pisanie przenosi go teraz w domenę języka.

Przyszłości zawodów pisarskich nie należy więc opisywać jako pojedynku człowieka z maszyną. Trafniejszy jest model rekonfiguracji łańcucha wartości. Maszyna przejmuje pewne etapy: wyszukiwanie, streszczanie, tworzenie wersji roboczych, wariantowanie, formatowanie czy personalizację. Człowiek przesuwa się ku definiowaniu problemu, zdobywaniu informacji pierwotnych, kontroli źródeł, interpretacji, strategii i nadawaniu znaczenia. W teorii brzmi to jak awans: maszyny wykonują rutynę, ludzie zajmują się tym, co twórcze. Historia organizacji

przestrzega jednak, że taki rezultat nie pojawia się automatycznie. Pracownik może zostać pozbawiony rutyny po to, by zyskać czas na myślenie, ale może też zostać zmuszony do obsługi dziesięciokrotnie większej liczby dokumentów. „Pokusa wydajności” jest tu kategorią ekonomiczną w równym stopniu co psychologiczną. Nacisk na wolumen może oddalać człowieka od procesu twórczego i spychać go w stronę funkcji operatora lub post-edytora. Nie powinniśmy jednak zakładać, że każda taka zmiana jest degradacją.

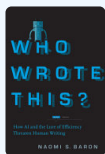
Jeżeli automat tworzy rutynowe zestawienie wyników giełdowych, a dziennikarka dzięki temu ma czas ujawnić oszustwo, automatyzacja zwiększa zakres ludzkiego dziennikarstwa. Jeśli prawnik przestaje ręcznie kopiować standardowe elementy odpowiedzi procesowej i poświęca czas na strategię sprawy, jego kompetencja zostaje wzmocniona. Ale jeżeli ten sam system sprawia, że reporter przestaje weryfikować źródła, prawnik traci umiejętność konstruowania argumentu, a tłumacz jedynie wygładza syntetyczną prozę według normy wydajności – narzędzie zaczyna zmieniać samą treść profesjonalizmu. Z tej analizy wynika ogólna zasada: automatyzacja tekstu jest bezpieczna tam, gdzie wartość leży w poprawnym rezultacie; staje się problematyczna tam, gdzie wartość wynika także z procesu dochodzenia do tego rezultatu. Wiadomość o wyniku meczu może być znakomita, choć nikt jej nie „przeżył”.

Opinia prawna, tekst naukowy, reportaż śledczy czy przekład literacki mają inną strukturę wartości, ponieważ proces selekcji, interpretacji i odpowiedzialności stanowi część ich jakości. To prowadzi do rozróżnienia, które musi znaleźć się w centrum niniejszego wywodu. Literatura, pamiętnik, tekst naukowy, akt prawny, oferta reklamowa czy instrukcja techniczna – wszystkie należą do świata pisma, lecz nie z tego powodu są wartościowe. W jednych tekstach liczy się przede wszystkim funkcja. W innych źródło, biografia, doświadczenie, oryginalność lub osobiste ryzyko wypowiedzi są częścią samego znaczenia. Dlatego pytanie „czy maszyna może pisać?” jest zbyt ogólne. Maszyna już pisze. Pytaniem znacznie trudniejszym pozostaje: które rodzaje tekstu mogą utracić ludzkiego autora bez utraty tego, co czyniło je wartościowymi?

W świecie, w którym syntetyczna przeciętność staje się darmowa i powszechna, paradoksalnie może wzrosnąć wartość tego, co nieobliczalne: unikalnego głosu, ryzyka intelektualnego i osobistego świadectwa. Prawdziwym wyzwaniem przyszłości nie będzie zatem walka z maszyną o prawo do pisania, lecz umiejętność odmawiania sobie wygody w imię zachowania zdolności do myślenia. Czy w pogoni za bezszwową efektywnością nie usuniemy właśnie tego poznawczego tarcia, które czyni nas ludźmi?

Inspiracje

[1]



"Who Wrote This" Naomi S Baron

ISBN: 9781503643574

Naomi S. Baron (ur. 1946) jest wybitną językoznawczynią i emerytowaną profesorką językoznawstwa na Uniwersytecie Amerykańskim w Waszyngtonie. Uzyskała tytuł doktora na Uniwersytecie Stanforda i piastowała prestiżowe stanowiska, w tym stypendystkę Fundacji Guggenheima, stypendystkę Fulbrighta i profesor wizytującą w Centrum Studiów Zaawansowanych Nauk Behawioralnych Stanforda. Jej bogata kariera akademicka koncentruje się na styku języka i technologii, a w szczególności na badaniu, w jaki sposób media cyfrowe, komunikacja mobilna i sztuczna inteligencja przekształcają ludzkie procesy czytania, pisania i procesy poznawcze. Jako płodna autorka, opublikowała liczne książki zgłębiające te tematy, w tym „Words Onscreen” i „How We Read Now”. Jej prace są powszechnie uznawane za krytyczną analizę wpływu technologii zorientowanych na wydajność na umiejętność czytania i pisania, autorstwo i zachowanie ludzkiej ekspresji twórczej w erze cyfrowej.

Naomi S. Baron (born 1946) is a prominent linguist and Professor Emerita of Linguistics at American University in Washington, D.C. She earned her Ph.D. from Stanford University and has held prestigious positions, including Guggenheim Fellow, Fulbright Fellow, and Visiting Scholar at the Stanford Center for Advanced Study in the Behavioral Sciences. Her extensive academic career has focused on the intersection of language and technology, specifically examining how digital media, mobile communication, and artificial intelligence reshape human reading, writing, and cognitive processes. A prolific author, she has published numerous books exploring these themes, including 'Words Onscreen' and 'How We Read Now.' Her work is widely recognized for its critical analysis of how efficiency-driven technologies impact literacy, authorship, and the preservation of human creative expression in the digital age.

American University

Literatura uzupełniająca

Książki

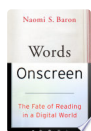
Autorzy przywoływani w artykule:



[1] "How We Read Now"

Naomi S. Baron

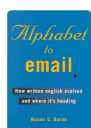
2021 | Oxford University Press | ISBN: 9780190084103



[2] "Words Onscreen"

Naomi S. Baron

2015 | Oxford University Press, USA | ISBN: 9780199315765



[3] "Alphabet to Email"

Naomi S. Baron

2000 | Psychology Press | ISBN: 9780415186858



[4] "Reader Bot"

Naomi S. Baron

2026 | ISBN: 9781503643949



[5] "Speech, Writing, and Sign"

Naomi S. Baron

1981

Literatura pokrewna (Google Scholar):

[6] "Data and its (dis)contents: A survey of dataset development and use in machine learning research"

Emily Bender



[7] "Computer Power and Human Reason: From Judgment to Calculation"

Joseph Weizenbaum

ISBN: 9780140179118

[8] "On the Dangers of Stochastic Parrots: Can Language Models Be Too Big? 🦜"

Emily Bender

[9] "Computation: Finite and Infinite Machines"

Marvin Minsky

[10] "SLIP"

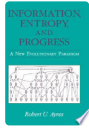
Joseph Weizenbaum

[11] "ELIZA"

Joseph Weizenbaum

[12] "Zygalski sheets"

Henryk Zygalski



[13] "information entropy"

Claude Shannon

Springer Science & Business Media | ISBN: 9780883189115



[14] "The mathematical theory of communication"

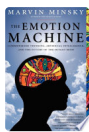
Claude Shannon

[15] "Programming a Computer for Playing Chess"

Claude Shannon

[16] "M.U.C. Love Letter Generator"

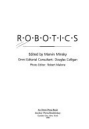
Christopher Strachey



[17] "The Emotion Machine"

Marvin Minsky

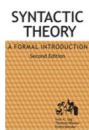
Simon and Schuster | ISBN: 9781416579304



[18] "Robotics"

Marvin Minsky

Anchor Books



[19] "Syntactic Theory: A Formal Introduction"

Emily Bender

Center for the Study of Language and Information Publica Tion

[20] "Distributed Artificial Intelligence Research Institute"

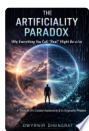
Timnit Gebru

[21] "Black in AI"

Timnit Gebru

[22] "Botipedia"

Philip M. Parker



[23] "THE ARTIFICIALITY PARADOX"

Dwyrnir Diningrat

Independently Published

Artykuły naukowe

Autorzy przywoływani:

[1] "Empathy and experience in HCI"

Peter Wright, John McCarthy

2008 | DOI: 10.1145/1357054.1357156

[Google Scholar](#)

[2] "AI Skills for Business Competency Framework"

The Alan Turing Institute, Innovate UK, Innovation and Technology (DSIT) Department for Science, Digital Catapult, Science and Technology Facilities Council

2026 | Zenodo (CERN European Organization for Nuclear Research) | DOI: 10.5281/zenodo.18892559

[Google Scholar](#)

[3] "Text Processing"

C. Gregg, M. Sahami, J. Clarke, A. Turing

2020 | Raku Recipes | DOI: 10.1007/978-1-4842-6258-0_10

[Google Scholar](#)

[4] "MATHEMATICAL SOLUTION OF THE ENIGMA CIPHER"

Marian Rejewski

1982 | Cryptologia | DOI: 10.1080/0161-118291856731

[Google Scholar](#)

[5] "Mathematical solution of Enigma cipher"

Marian Rejewski

1989 | Artech House, Inc. eBooks

[Google Scholar](#)

[6] "A Mathematical Theory of Communication"

Claude E. Shannon

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

[Google Scholar](#)

[7] "The Mathematical Theory of Communication"

Claude E. Shannon, Warren Weaver, Norbert Wiener

1950 | Physics Today | DOI: 10.1063/1.3067010

[Google Scholar](#)

[8] "Communication Theory of Secrecy Systems"

Claude E. Shannon

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

[Google Scholar](#)

[9] "Toward a mathematical semantics for computer languages"

Dana Scott, C. Strachey

1971 | Medical Entomology and Zoology

[Google Scholar](#)

[10] "A Theory of Programming Language Semantics"

Robert Milne, C. Strachey

1977 | CERN Document Server (European Organization for Nuclear Research)

[Google Scholar](#)

[11] "Fundamental Concepts in Programming Languages"

C. Strachey

2000 | LISP and Symbolic Computation | DOI: 10.1023/a:1010000313106

[Google Scholar](#)

[12] "Continuations: A Mathematical Semantics for Handling Full Jumps"

C. Strachey, Christopher P. Wadsworth

2000 | LISP and Symbolic Computation | DOI: 10.1023/a:1010026413531

[Google Scholar](#)

[13] "Some Philosophical Problems from the Standpoint of Artificial Intelligence"

John McCarthy, Patrick J. Hayes

1981 | Elsevier eBooks | DOI: 10.1016/b978-0-934613-03-3.50033-7

[Google Scholar](#)

[14] "Technology as experience"

John McCarthy, Peter Wright

2004 | interactions | DOI: 10.1145/1015530.1015549

[Google Scholar](#)

[15] "(1969) Marvin Minsky and Seymour Papert, Perceptrons, Cambridge, MA: MIT Press, Introduction, pp. 1-20, and p. 73 (figure 5.1)"

Marvin Minsky, Seymour Papert

1988 | Neurocomputing, Volume 1 | DOI: 10.7551/mitpress/4943.003.0015

[Google Scholar](#)

[16] "ELIZA—a computer program for the study of natural language communication between man and machine"

Joseph Weizenbaum

1966 | Communications of the ACM | DOI: 10.1145/365153.365168

[Google Scholar](#)

[17] "Computer Power and Human Reason: From Judgement to Calculation."

John F. Crowther, Joseph Weizenbaum

1977 | Contemporary Sociology A Journal of Reviews | DOI: 10.2307/2062802

[Google Scholar](#)

[18] "Attention Is All You Need"

Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones

2025 | DOI: 10.65215/2q58a426

[Google Scholar](#)

[19] "Music Transformer: Generating Music with Long-Term Structure"

Cheng-Zhi Anna Huang, Ashish Vaswani, Jakob Uszkoreit, Noam Shazeer, Ian Simon

2019 | International Conference on Learning Representations

[Google Scholar](#)

[20] "Self-Attention with Relative Position Representations"

Peter Shaw, Jakob Uszkoreit, Ashish Vaswani

2018 | DOI: 10.18653/v1/n18-2074

[Google Scholar](#)

[21] "Decoding with Large-Scale Neural Language Models Improves Translation"

Ashish Vaswani, Yinggong Zhao, Victoria Fossum, David Chiang

2013 | DOI: 10.18653/v1/d13-1140

[Google Scholar](#)

[22] "The Value and (Re)Use of Photograph Collections Globally"

Brenda Bernier, Penley Knipe

2025 | Conservation of Photographs | DOI: 10.4324/9781003096108-2

[Google Scholar](#)

[23] "Climbing towards NLU: On Meaning, Form, and Understanding in the Age of Data"

Emily M. Bender, Alexander Koller

2020 | Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics | DOI: 10.18653/v1/2020.acl-main.463

[Google Scholar](#)

[24] "Data Statements for Natural Language Processing: Toward Mitigating System Bias and Enabling Better Science"

Emily M. Bender, Batya Friedman

2018 | Transactions of the Association for Computational Linguistics | DOI: 10.1162/tacl_a_00041

[Google Scholar](#)

[25] "Data and its (dis)contents: A survey of dataset development and use in machine learning research"

Amandalynne Paullada, Inioluwa Deborah Raji, Emily M. Bender, Emily Denton, Alex Hanna

2020 | arXiv (Cornell University) | DOI: 10.1016/j.patter.2021.100336

[Google Scholar](#)

[26] "On Achieving and Evaluating Language-Independence in NLP"

Emily M. Bender

2011 | Linguistic Issues in Language Technology | DOI: 10.33011/lilt.v6i.1239

[Google Scholar](#)

[27] "Review of Gender, Sexuality and Meaning: Linguistic Practice and Politics by Sally McConnel-Ginet"

Veronika Koller,

2012 | Lodz Papers in Pragmatics | DOI: 10.1515/lpp-2012-0016

[Google Scholar](#)

Artykuły pokrewne (Google Scholar):

[28] "Always On"

Naomi S. Baron

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

[Google Scholar](#)

[29] "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](#)

[30] "See you Online"

Naomi S. Baron

2004 | Journal of Language and Social Psychology | DOI: 10.1177/0261927x04269585

[Google Scholar](#)

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

Naomi S. Baron

2015

[Google Scholar](#)

[32] "Text Messaging and IM"

Rich Ling, Naomi S. Baron

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

[Google Scholar](#)

[33] "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](#)

[34] "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](#)

[35] "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](#)

[36] "Alphabet to Email"

Naomi S. Baron

2002 | DOI: 10.4324/9780203194317

[Google Scholar](#)

[37] "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: 3a244e0e-168a-42e7-823e-4497296cd8c3