

Przenikanie umysłów czy iluzja komunikowania? Koncepcja oralności i piśmienności Waltera J. Onga wobec generatywnych modeli językowych (LLM)

Wojciech Furman

Dynamiczny rozwój generatywnych modeli językowych (*Large Language Models*) stał się jednym z wyzwań dla badań nad komunikowaniem. Narzędzia takie jak ChatGPT, Gemini czy Claude nie tylko automatyzują wytwarzanie tekstu, lecz także sprawiają wrażenie wchodzenia w obszary dotychczas zarezerwowane dla ludzkich umiejętności, a mianowicie streszczania, interpretowania, argumentowania, a nawet twórczości. W związku z tym pojawia się pytanie postawione w tytule artykułu: czy mamy do czynienia z nową formą dialogu i „przenikaniem umysłów”¹, czy tylko z iluzją komunikowania się człowieka z narzędziem, działającym bez jakichkolwiek intencji?

Pomocą w poszukiwaniu odpowiedzi na powyższe pytanie może być koncepcja oralności i piśmienności Waltera J. Onga², która pozwala spojrzeć z innego punktu widzenia na problem współpracy człowieka z modelami językowymi, podjęty przez Aleksandrę Przegalińską i Tamillę Triantoro³. Ong wskazał że pismo, traktowane jako technologia komunikacyjna, nie jest neutralnym narzędziem, lecz głęboko przekształca ludzkie sposoby myślenia, organizowania wiedzy i działania językiem. Choć jego prace powstawały w epoce mediów analogowych, proponowana przez Onga rama analityczna okazuje się zaskakująco użyteczna.

Struktura artykułu opiera się na trzech węzłowych problemach: (1) sposobach konstruowania tekstu; (2) technologii, która jest wsparciem i rozszerzeniem umysłu; (3) słowie stanowiącym działanie. Część teoretyczna została uzupełniona refleksją empiryczną w formie autoetnografii autora, opisującej skrótoowo proces tworzenia tekstu naukowego z użyciem LLM. Takie rozwiązanie pozwala uchwycić realne funkcjonowanie modelu językowego w praktyce badawczej.

Sposób konstruowania tekstu: standardowa formuła – wzorzec – prawdopodobieństwo

Jednym z istotnych punktów wspólnych między Ongowską analizą oralności pierwotnej a opisem działania współczesnych generatywnych modeli językowych jest sposób powstawania tekstu. W obu przypadkach tworzenie tekstu nie jest całkowicie dowolne, lecz polega na rekonfigurowaniu istniejących zasobów: formuł, wzorców, struktur i powtarzalnych elementów językowych.

¹ To sformułowanie znalazło się w tytule książki A. Przegalińskiej i T. Triantoro, *Przenikanie umysłów. Potencjał twórczy współpracy z AI*, AI Books, Warszawa 2024.

² W.J. Ong, *Oralność i piśmienność. Słowo poddane technologii*, tłum. J. Japola, wyd. 2, Wydawnictwa Uniwersytetu Warszawskiego, Warszawa 2011.

³ A. Przegalińska, T. Triantoro, dz. cyt.

Ong: Oralność pierwotna i gotowe formuły

Opowieści oralne są budowane z gotowych tematów, schematów fabularnych, zwrotów, formuł i typów ludzkich. Przykładem może być dzielny rycerz, piękna księżniczka lub zły czarownik. Wygłaszane rytmicznie formuły są pełne powtórzeń, epitetów, przeciwieństw i przysłów. Taka konstrukcja pozwala na łatwiejsze zapamiętywanie i odtwarzanie opowieści w społeczeństwie, które nie zna pisma.

Bard nie jest autorem opowieści we współczesnym znaczeniu tego słowa. On nie wymyśla treści, lecz je rekonstruuje, zszywa, aktualizuje, aby dopasować swą opowieść do sytuacji i oczekiwań słuchaczy, przestrzegając przy tym rytmicznych rygorów metrum. Można powiedzieć, że bard rekonfiguruje istniejące zasoby kulturowe.

Formuły oralne są silnie zakorzenione w pamięci zbiorowej i bliskie doświadczeniom słuchaczy. Odwołują się do znanych sytuacji, napięć i konfliktów, a cała opowieść jest wytwarzana w bezpośredniej relacji barda z jego audytorium.

LLM: Wzorce, analogie i statystyczna aproksymacja

Generatywne modele językowe są trenowane na ogromnych korpusach tekstów. Te teksty zostają przeszukane w celu znalezienia regularności, analogii i wzorców. Na tej podstawie model generuje kolejne teksty, przewidując najbardziej prawdopodobną kolejność słów.

Jak pisał Andrzej Dragan, modele składają się z sieci neuronowych, których kluczową kompetencją jest dostrzeganie analogii oraz kompresja danych⁴. Modele nie przechowują wszystkich danych treningowych, lecz tylko ich esencję, czyli zbiory reguł i zależności, które pozwalają na uogólnienia i odpowiadanie na nowe, choć podobne problemy. W tym sensie tekst generowany przez model nie jest prostą recytacją, która tylko odtwarza znaną treść, lecz jest mieszaniną zapamiętanych fragmentów oraz formuł tworzonych przez analogię⁵.

Dobór słów w tekście tworzonym przez model językowy ma charakter probabilistyczny i jest pozbawiony odniesień do rzeczywistego świata. Model nie ma odniesienia semantycznego do świata pozajęzykowego, a jedynie operuje na poziomie relacji między znakami. Skutkiem tego jest podatność stosowanych aktualnie wersji modeli na tzw. halucynacje, czyli tworzenie formalnie spójnych, lecz fałszywych treści.

Podobieństwa i różnice

Na poziomie formalnym podobieństwo między bardem oralnym a modelem językowym jest uderzające. W obu przypadkach przekaz powstaje poprzez rekombinację gotowych elementów, którymi są schematy, formuły, słowa lub tokeny (token jest to zbitka liter, stanowiąca najmniejszą jednostkę analizowaną przez LLM⁶). Co najwyżej można dodać, że z upływem czasu owe podstawowe elementy przekazu stają się coraz mniejsze.

⁴ A. Dragan, *Quo vAldis*, Wydawnictwo Otwarte, Kraków 2025, s. 156.

⁵ A. Dragan, tamże, s. 157, 214.

⁶ A. Dragan, tamże, s. 226.

Różnice są jednak zasadnicze. Bard działa w konkretnej sytuacji społecznej i pozostaje w bezpośrednim kontakcie ze swymi odbiorcami. Zna albo na bieżąco poznaje oczekiwania słuchaczy, reaguje na wyrażaną przez nich aprobatę lub niezadowolenie. Wyraża swój stosunek do rzeczywistości i w tę rzeczywistość ingeruje. Natomiast model językowy nie definiuje sytuacji, nie ma doświadczenia ani pamięci biograficznej. Działa na poziomie struktur językowych, bez związku z rzeczywistością.

Pomiędzy bardem a modelem językowym można usytuować twórcę tekstu pisanego. On także nie ma bezpośredniego kontaktu z odbiorcami i może ich sobie tylko wyobrażać, lecz jego twórczość pozostaje powiązana z rzeczywistym światem oraz osobistym doświadczeniem.

Technologia jako wsparcie i rozszerzenie umysłu

Jednym z istotnych wątków w koncepcji Onga jest teza, że pismo nie tylko pośredniczy w przekazywaniu treści, lecz także przekształca procesy ludzkiego myślenia. Pismo nie jest zewnętrznym dodatkiem, lecz czymś w rodzaju rusztowania, pozwalającego na organizowanie poznawania, konceptualizowanie doświadczeń oraz wyrażanie myśli. Takie spojrzenie na pismo pozwala traktować modele językowe jako kolejną technologię rozszerzającą ludzki umysł.

Ong: Pismo jako rusztowanie poznawcze

Pismo pozwala na oddzielenie wypowiedzi od sytuacji mówienia, a także na utrwalenie przekazu oraz wielokrotne do niego powracanie. Dzięki temu stają się możliwe takie działania, które byłyby niezwykle trudne lub wręcz niemożliwe w kulturze oralnej. Posługując się tekstem pisany ludzie mogą precyzyjnie wskazywać na porządek wynikania, abstrahować, klasyfikować, interpretować oraz dokonywać krytycznej analizy. Możliwe staje się porównywanie stanowisk, śledzenie argumentacji i korygowanie błędów. W ten sposób pismo sprzyja rozwojowi nauki, sprzyja także powstawaniu metajęzyka, czyli języka refleksji nad językiem. Krańcową manifestacją oderwania języka od kontekstu są współczesne słowniki⁷.

Ong dostrzegł ambiwalentny charakter pisma. Przy wszystkich swoich zaletach pismo częściowo zwalnia człowieka z potrzeby pamiętania oraz wprowadza dystans między autorem wypowiedzi a anonimowym, pojedynczym czytelnikiem. Zapisany tekst daje wprawdzie czytelnikowi szansę na lekturę i interpretację, lecz to od czytelnika zależy, czy i w jaki sposób ową szansę wykorzysta oraz jakie wnioski z niej wyciągnie. Przypadki zbiorowego wysłuchiwanie zapisanych wcześniej tekstów, takie jak akcja „Narodowe Czytanie”, zdarzają się przecież wyjątkowo.

LLM: Model językowy jako pomocnik w myśleniu

Model językowy potrafi przejąć od człowieka część operacji poznawczych. Model przegląda ogromne zasoby tekstów, sporządza streszczenia, porządkuje informacje i nadaje im strukturę, a także proponuje warianty rozwiązywania proble-

⁷ R. Harris, *Racjonalność a umysł piśmienny*, tłum. M. Rakoczy, Wydawnictwa Uniwersytetu Warszawskiego, Warszawa 2014, s. 252.

mów. Działając w ten sposób model pomaga człowiekowi zmniejszać kłopoty, jakie wynikają z potrzeby orientowania się w nadmiarze informacji.

Model nie działa jednak autonomicznie. To człowiek formułuje i kalibruje prompt, wyznacza kierunek poszukiwań, ocenia trafność odpowiedzi i decyduje o przyjęciu, modyfikowaniu lub odrzuceniu propozycji wypracowanych przez model. Model wnosi do tego procesu moc obliczeniową, zdolność do statystycznej aproksymacji oraz generowania licznych wariantów. A człowiek korzysta ze swej intuicji, emocji oraz doświadczeń.

W tym sensie model jest pomocnikiem, lecz nie jest kreatywnym partnerem w myśleniu. Relacja człowiek-LLM ma charakter asymetryczny. LLM jest bardzo sprawnym i niezwykle wydajnym narzędziem, które jednak nie dąży do poznawania ani działania.

Podobieństwa i różnice

Zarówno pismo jak i modele językowe spełniają funkcje pomocnicze wobec ludzkiego umysłu. Ułatwiają gromadzenie i porządkowanie wiedzy, rozszerzają zakres operacji poznawczych i podnoszą efektywność pracy umysłowej. W obu przypadkach technologia wpływa korzystnie na tempo, skalę i styl myślenia człowieka.

Różnica polega na stopniu autonomii i na stwarzaniu pozoru podmiotowości. Model językowy nie jest aktywnym podmiotem ani partnerem w dialogu, nie wyraża intencji ani nie ponosi odpowiedzialności za skutki wyprodukowanych przez siebie wypowiedzi. Funkcjonowanie modelu nie znosi zatem fundamentalnej różnicy między działającym podmiotem a jego narzędziem.

Pozór prowadzenia dialogu może jednak sprawiać, że sam człowiek zaczyna traktować model językowy jako partnera, czyli że nadaje mu status uczestnika rozmowy. Czasami nawet zwierza mu się ze swych kłopotów i traktuje model jak niezawodnego, życzliwego i zawsze dostępnego powiernika oraz fachowego doradcę. Opisy takich sytuacji można znaleźć w prasie⁸.

Słowo jako działanie

Jednym z istotnych wymiarów różnicy między ludzkim komunikowaniem a funkcjonowaniem modeli językowych jest potraktowanie słowa jako działania. Według Onga w kulturach oralnych mowa nie jest jedynie przekazywaniem wiadomości, lecz jest przede wszystkim wydarzeniem społecznym. Bard coś czyni, zmienia relacje i wywołuje reakcje. To spostrzeżenie Onga nawiązuje do teorii aktów mowy Austina⁹ i potwierdza tę teorię.

Ong: Mowa jako czyn, wydarzenie i agon

W ujęciu Onga słowo mówione jest zawsze wydarzeniem. Powstaje w określonym czasie i miejscu, jest kierowane do konkretnych słuchaczy. Opowiadanie, przemowa czy pieśń nie są neutralnym przekazem wiadomości, lecz sposobem oddziaływania – próbą wywołania emocji, zaproszeniem do uzyskania odpowiedzi,

⁸ Pisze o tym na przykład J. Cieśla, *Chat mój brat*, „Polityka” 2025, nr 51/52, s. 36-38.

⁹ J.L. Austin, *Mówienie i poznawanie. Rozprawy i wykłady filozoficzne*, tłum. B. Chwedeńczuk, WN PWN, Warszawa 1993.

wciąganiem słuchaczy w atmosferę wyrażającą określony porządek wartości i znaczeń. Należy dodać, że takie ujęcie procesu komunikowania oralnego przez Onga kładzie nacisk na tylko jeden człon diady, jaką wskazał John Fiske. Wedle Fiskego komunikowanie należy traktować jako przekazywanie wiadomości lub jako tworzenie i wymianę znaczeń¹⁰.

Ong zwraca też uwagę na agonistyczny charakter komunikowania oralnego. Mowa jest często pojedynkiem na słowa, jest zmaganiem się, w którym ujawniają się konflikty i napięcia. Retoryka oralna dopuszcza przesadę, opieranie się na kontrastach, opozycjach i dobitnych ocenach, ponieważ jej celem jest nie tyle obiektywne przedstawienie faktów, co skuteczne oddziaływanie na odbiorców. Perswazja retoryczna jest przekonywaniem zmierzającym do wyboru wartości, opinii, mniemań¹¹. Kluczowe znaczenie ma także antycypowane sprzężenie zwrotne, gdyż mówca dostosowuje swoją wypowiedź do reakcji słuchaczy, do ich aprobaty lub sprzeciwu.

Mowa jest zawsze relacyjna i performatywna. Sens mowy wyłania się w interakcji mówcy i słuchaczy, nie może być oddzielony od intencji mówiącego ani od społecznych konsekwencji jego wypowiedzi.

LLM: Generowanie wypowiedzi bez intencji działania

Model językowy produkuje wypowiedzi, odpowiadając w ten sposób na prompt użytkownika. Nie działa w świecie społecznym, nie wyraża intencji, nie dąży do zamierzonych skutków ani nie antycypuje reakcji użytkownika. Tworzone teksty są rezultatem operacji obliczeniowych, a nie aktami działania.

Model nie bierze udziału w konfliktach ani potyczkach językowych. Może wprawdzie generować teksty o charakterze perswazyjnym, polemicznym lub nawet emocjonalnym, lecz czyni to wyłącznie na poziomie formalnym, symulując postępowanie znane z danych treningowych.

Interakcja człowieka z LLM, nawet jeśli stwarza wrażenie dialogu, pozostaje asymetryczna. To użytkownik nadaje sens wypowiedziom modelu językowego, interpretuje je jako znaczące i decyduje o ewentualnym użyciu tych wypowiedzi jako narzędzi pomagających w konstruowaniu przekazu. Możliwe tryby współpracy człowieka z LLM wedle Przeglasińskiej i Triantoro to: praca równoległa (człowiek i model działają niezależnie), praca kontrybutywna (model wspiera człowieka w określonym zakresie), praca zespołowa (model działa jako członek zespołu ludzi), praca hybrydowa (model tworzy pomysły, a człowiek nadaje kierunek, wnosi zrozumienie kontekstu i uwzględnia kwestie etyczne), praca synergiczna (ciągły wzajemny wpływ człowieka i modelu)¹².

Podobieństwa i różnice

Zarówno teoria oralności Onga, jak i teoria aktów mowy Austina, wskazują na związek mówienia i działania. Komunikujący się człowiek jest działającym podmiotem, korzystającym z różnych mediów i różnych narzędzi. Wśród tych narzędzi

¹⁰ J. Fiske, *Wprowadzenie do badań nad komunikowaniem*, tłum. A. Gierczak, Astrum, Wrocław 1999, s. 16.

¹¹ M. Korolko, *Sztuka retoryki. Przewodnik encyklopedyczny*, Wiedza Powszechna, Warszawa 1990, s. 35.

¹² A. Przeglasińska, T. Triantoro, dz. cyt., s. 22-26.

są modele językowe, spełniające urozmaicone funkcje – od czysto instrumentalnych po quasi-dialogowe.

Choć interakcja człowieka z generatywnym modelem językowym przyjmuje formę dialogu, nie oznacza to, że mamy do czynienia z rzeczywistym komunikowaniem. Wymiana przekazów zakłada bowiem aktywność podmiotów wyrażających intencje, odpowiedzialnych i usytuowanych w społeczeństwie¹³. Natomiast model językowy nie chce niczego osiągnąć ani nie ponosi konsekwencji swoich wypowiedzi. Model nie jest rozmówcą, lecz swoistym medium – technologicznym interfejsem, reagującym na bodźce językowe zgodnie z rozkładami prawdopodobieństwa.

Część empiryczna, czyli jak powstał ten artykuł

Ta część artykułu jest próbą opisu działania autora we współpracy z generatywnym modelem językowym ChatGPT. Była to najprostsza wersja tego modelu, dostępna bezpłatnie na przełomie lat 2025/2026.

Podstawowy materiał empiryczny stanowiły fiszki sporządzone przez autora podczas lektury trzech kluczowych, wspominanych już książek, autorstwa Waltera J. Onga, Aleksandry Przeglasińskiej i Tamilli Triantoro oraz Andrzeja Dragana. Każda z fiszek liczyła od 12 do 30 tysięcy znaków ze spacjami. Fiszki powstały wyłącznie na użytek autora, miały charakter selektywny i interpretacyjny. Nie były streszczeniem dzieła, lecz zapisem fragmentów uznanych za istotne, uzupełnionych o nasuujące się w trakcie lektury komentarze, skojarzenia, uwagi i pytania badawcze.

Podstawą dalszej pracy z modelem nie były cyfrowe wersje wspomnianych książek ani ich pełne papierowe teksty, lecz jedynie autorskie notatki. Zapewniało to autorskie odczytanie każdego dzieła

Na kolejnym etapie fiszki zostały skopiowane do interfejsu modelu językowego, któremu wydano polecenie: „Na podstawie poniższych notatek zaproponuj tytuł, węzłowe problemy i układ artykułu naukowego”. Odpowiedź modelu dotycząca tytułu miała postać umieszczoną w ramce.

Propozycja tematu artykułu

Wariant główny (najbardziej „nośny” teoretycznie)

„Od oralności wtórnej do współdzielonej sprawczości. Sztuczna inteligencja jako nowa technologia słowa i komunikowania”

Warianty alternatywne

„AI jako technologia słowa: oralność, piśmienność i komunikowanie w epoce generatywnych modeli językowych”

„Komunikowanie jako antycypowane sprzężenie zwrotne: Walter J. Ong i współpraca człowieka z AI”

„Przenikanie umysłów czy iluzja dialogu? Ongowska teoria komunikowania wobec generatywnej sztucznej inteligencji”

Pierwszą reakcją na taką odpowiedź było poczucie dezorientacji i dystansu. Każda wersja tytułu rozmijała się z intuicją badawczą autora, jaka wynikała z lektury trzech wspomnianych wyżej książek. Ten moment zniechęcenia, lub nawet od-

¹³ R. Burkart, *Kommunikationswissenschaft. Grundlagen und Problemfelder*, Böhlau Verlag, Wien 2002, S. 25-29.

rzucenia, okazał się jednak poznawczo istotny. Zamiast spodziewanego przyspieszenia pracy pojawiła się konieczność gruntownego namysłu nad tym, dlaczego propozycje modelu odbiegają od rysujących się niewyraźnie i początkowo trudno uchwytnych zamysłów autora. W ten sposób następowało stopniowe precyzowanie własnego stanowiska oraz artykułowanie różnic między myśleniem autora a niekoniecznie trafnymi propozycjami modelu.

Tytuł figurujący w ramce na pierwszym miejscu i określony jako najlepszy stanowi zawężenie koncepcji Onga do oralności wtórnej, czyli komunikowania za pośrednictwem analogowego telefonu, radia i telewizji, natomiast pomija najbardziej istotne spostrzeżenia Onga dotyczące oralności pierwotnej. Tytuł drugi, podobnie jak pierwszy, określa AI jako technologię słowa, co jest bardzo ogólnym i raczej dziwacznym sformułowaniem, skoro od dawna posługujemy się bardziej trafnym terminem „medium”. Ponadto takie określenie AI nie dostrzega możliwości tworzenia obrazów i dźwięków. Tytuł trzeci zawęża problematykę zamierzonego artykułu do tylko jednej węzłowej kwestii. Dopiero ostatnia wersja tytułu kwalifikowała się do dalszych rozważań i po wniesieniu poprawek okazała się możliwa do przyjęcia.

Powyższe postępowanie zostało powtórzone wielokrotnie w ciągu kilkunastu tygodni, przy czym każda kolejna interakcja z modelem opierała się na coraz bardziej precyzyjnie formułowanych wskazówkach. Model proponował liczne warianty, które podlegały krytycznej ocenie i stanowiły podstawę do dalszego myślenia. Różnica polegała na tym, że autor zwykle potrzebował kilku dni, a model – około jednej minuty. Ponadto model dysponował treścią fiszek i swoimi zasobami pamięci, a autor miał pod ręką papierowe oryginały książek.

Najpierw rozważany był tytuł, a następnie węzłowe problemy i plan artykułu. Kiedy plan został uznany za satysfakcjonujący, analogiczna procedura została zastosowana do kolejnych punktów artykułu, zawsze na podstawie przyjętego, a następnie korygowanego i uzupełnianego planu. Tekst wygenerowany przez model stanowił roboczy szkic, wymagający zmian i uzupełnień. LLM podpowiadał, lecz nie zawsze trafnie. Jego podpowiedzi wymagały wielokrotnej oceny i korekty.

Współpraca z modelem spowodowała zmianę dotychczasowego trybu pracy autora. Dotychczas po lewej stronie ekranu zwykle figurował plan artykułu lub rozdziału, na bieżąco korygowany i uzupełniany o pojawiające się pomysły, a po prawej stronie znajdował się powstający tekst. Teraz ekran został podzielony na trzy części. Po lewej stronie znalazły się prompty autora i odpowiedzi modelu, na środku plan podlegający korektom i uzupełnieniom, a po prawej stronie – powstający tekst artykułu.

Przedstawiony powyżej skrótowo sposób powstawania artykułu potwierdza, że model językowy może pełnić funkcję rusztowania poznawczego oraz generatora wariantów tekstu, lecz nie przejmuje odpowiedzialności za sens. Model wspierał myślenie, lecz pozostawał narzędziem, a nie podmiotem komunikowania.

Wnioski

Artykuł nie został oparty na przeciwstawieniu oralności, piśmienności i komunikowania cyfrowego z użyciem modeli językowych, lecz stanowi próbę spojrzenia na trzy węzłowe, przekrojowe problemy: sposób konstruowania tekstu, technologii jako wsparcie umysłu oraz słowo rozumiane jako działanie.

Odpowiedź na pytanie postawione w tytule artykułu nie jest jednoznaczna. Koncepcja oralności i piśmienności Onga skonfrontowana z działaniem generatywnego modelu językowego nie prowadzi ani do prostego potwierdzenia tezy o przenikaniu umysłów, ani do zdecydowanego uznania interakcji człowieka z LLM za iluzję komunikowania. Modele rzeczywiście rekonfigurują istniejące zasoby kulturowe, rozpoznają wzorce i analogie oraz generują nowe wypowiedzi w sposób, który przypomina praktyki międzyludzkie. Jednocześnie brak im kluczowych cech działania językowego w sensie wskazanym przez Onga, a mianowicie intencji, rozumienia sytuacji, zakorzenienia w relacjach społecznych oraz odpowiedzialności za skutki wypowiedzi.

Z takiej perspektywy modele językowe należy uznać za zaawansowane narzędzia poznawcze, wspierające i na swój sposób rozszerzające działanie ludzkiego umysłu, czyli funkcjonujące tak, jak pismo w analizach Onga. Modele potrafią przyspieszać operacje poznawcze i wskazywać warianty rozwiązań, lecz ich propozycje wymagają rozpatrywania i korygowania przez człowieka.

Empiryczna część artykułu potwierdziła tę diagnozę. Opis pracy z LLM wskazuje, że jest to bardzo użyteczne narzędzie, pomocne w porządkowaniu materiału, tworzeniu propozycji i testowaniu alternatywnych ujęć. Szkice formułowane przez model były następnie oceniane i odrzucane bądź korygowane przez autora. Gdyby sięgnąć do wspomnianej wyżej propozycji Przeglasińskiej i Triantoro, taki tryb współpracy należałoby określić jako kontrybutywny bądź hybrydowy.

Ong dostrzegł, że przejście od oralności do piśmienności nie jest ruchem jednokierunkowym, lecz towarzyszą mu powroty na wyższym poziomie, określone przezeń jako oralność wtórna. Na tę kwestię zwracało uwagę wielu badaczy, na przykład zdaniem Maryli Hopfinger zapis cyfrowy stanowi wspólny mianownik dla tekstów wszelkiego rodzaju¹⁴. Również tłumacz dzieła Onga na język polski Józef Japola przytoczył z aprobatą opinię, iż język symboliczny i język analityczny funkcjonują jako tandem¹⁵. Oba te języki uzupełniają się wzajemnie i podnoszą efektywność oraz fortunność komunikowania – to pierwsze rozumiane jako przekazywanie i odczytywanie wiadomości, a to drugie jako tworzenie i rozumienie znaczeń. W tym sensie generatywne modele językowe mogą jawić się jako kolejny odcinek spirali, w której późniejsze technologie nie znoszą poprzednich, lecz wchodzą z nimi w dynamiczne sprzężenie.

¹⁴ M. Hopfinger, *Literatura i media. Po 1989 roku*, Oficyna Naukowa, Warszawa 2010, s. 72.

¹⁵ J. Japola, *Od tłumacza*, [w:] W.J. Ong, dz. cyt., s. 24.

Bibliografia

- Austin J.L., *Mówienie i poznawanie. Rozprawy i wykłady filozoficzne*, tłum. B. Chwedeńczuk, WN PWN, Warszawa 1993.
- Burkart R., *Kommunikationswissenschaft. Grundlagen und Problemfelder*, Böhlau Verlag, Wien 2002.
- Cieśla J., *Chat mój brat*, „Polityka” 2025, nr 51/52, s. 36-38.
- Dragan A., *Quo vAldis*, Wydawnictwo Otwarte, Kraków 2025.
- Fiske J., *Wprowadzenie do badań nad komunikowaniem*, tłum. A. Gierczak, Astrum, Wrocław 1999.
- Harris R., *Racjonalność a umysł piśmienny*, tłum. M. Rakoczy, Wydawnictwa Uniwersytetu Warszawskiego, Warszawa 2014.
- Hopfinger M., *Literatura i media. Po 1989 roku*, Oficyna Naukowa, Warszawa 2010.
- Japola J., *Od tłumacza*, [w:] W.J. Ong, *Oralność i piśmienność. Słowo poddane technologii*, wyd. 2, Wydawnictwa Uniwersytetu Warszawskiego, Warszawa 2011, s. 17-26.
- Korolko M., *Sztuka retoryki. Przewodnik encyklopedyczny*, Wiedza Powszechna, Warszawa 1990.
- Ong W.J., *Oralność i piśmienność. Słowo poddane technologii*, tłum. J. Japola, wyd. 2, Wydawnictwa Uniwersytetu Warszawskiego, Warszawa 2011.
- Przegalińska A., Triantoro T., *Przenikanie umysłów. Potencjał twórczy współpracy z AI*, AI Books, Warszawa 2024.

Streszczenie

Artykuł stanowi próbę skonfrontowania koncepcji Onga ze sposobem działania generatywnych modeli językowych. Autor proponuje odejście od przeciwstawienia oralności, piśmienności i sztucznej inteligencji na rzecz analizy trzech węzłowych problemów dotyczących komunikowania: sposobu konstruowania przekazu, technologii jako wsparcia umysłu oraz słowa jako działania. W części teoretycznej odniesiono ustalenia Onga dotyczące oralności pierwotnej do generowania tekstu przez modele językowe, opartego na rozpoznawaniu wzorców i statystycznej aproksymacji. Szczególną uwagę poświęcono różnicy między ludzką mową jako intencjonalnym działaniem a generowaniem wypowiedzi przez modele językowe, pozbawione własnych celów, doświadczenia i biograficznej pamięci. Część empiryczna ma charakter autoetnograficzny i rekonstruuje proces powstawania artykułu jako formę kontrolowanej współpracy autora z modelem językowym. Analiza ta pokazuje, że model może pełnić funkcję poznawczego rusztowania, zwiększać wariantywność rozwiązań i tempo pracy, lecz nie zastępuje interpretacji, krytyki ani odpowiedzialności autora. Wnioski wskazują, że relacja człowiek-model językowy nie polega ani na „przenikaniu umysłów”, ani na rzeczywistym dialogu, lecz ma specyficzną formę zapośredniczonego komunikowania.

Słowa kluczowe: Walter J. Ong, generatywne modele językowe LLM, relacja człowiek-LLM.

Abstract

The article attempts to confront Ong's concept with the functioning of generative language models. The author proposes moving away from the opposition of orality, literacy, and artificial intelligence towards an analysis of three central problems concerning communication: the way of constructing the message, technology as a support for the mind, and the word as action. The theoretical part relates Ong's findings on primal orality to the generation of text by language models based on pattern recognition and statistical approximation. Special attention was paid to the difference between human speech as an intentional action and the generation of speech by language models, which lack their own goals, experience, and biographical memory. The empirical part is of an autoethnographic nature and reconstructs the process of creating the article as a form of controlled collaboration between the author and the language model. This analysis shows that the model can function as a cognitive framework, increase the variability of solutions and the speed of work, but it does not replace the interpretation, criticism or responsibility of the author. The results indicate that the human-language model relationship is neither a "converging mind" nor a real dialogue, but rather a specific form of mediated communication.

Key words: Walter J. Ong, generative language models LLM, human-LLM relationship.