

PIOTR ZAWOJSKI

Uniwersytet Śląski

<https://orcid.org/0000-0003-2291-0446>

Kulturowe i filozoficzne konteksty ekspansji sztucznej inteligencji

Cultural and Philosophical Contexts of the Expansion of Artificial Intelligence

Artykuł recenzowany / Peer-reviewed article

Abstract

The expansion of Artificial Intelligence, particularly in the wake of the presentation by OpenAI of the GPT model in November 2022, lived to see a number of elaborations and assorted interpretations, predominantly referring to innovative solutions making use of artificial neural networks and deep learning. The author of the presented text refers to such discussions but chiefly attempts to focus on the cultural and philosophical consequences of the development of new technologies associated with AI. Consideration involving relations between human and non-human intelligence, the reinterpretation of the Serres conception of the human parasite, reflection on threats that can appear within the domain of man's environment, epistemology, and the domination of algorithms – these are some of the broached problems. The perspective of cultural studies-philosophical examination making use of knowledge about the technological-computer science aspects of AI, constitutes for the author a methodological framework of sorts for the studied phenomena.

Keywords: Artificial Intelligence; non-human intelligences; Michel Serres; parasite; Byung-Chul Han

Abstrakt

Ekspansja sztucznej inteligencji, zwłaszcza po zaprezentowaniu przez OpenAI modelu GPT w listopadzie 2022 roku, doczekała się szeregu omówień i różnych interpretacji, przede wszystkim w odniesieniu do rozwiązań wykorzystujących sztuczne sieci neuronowe (*artificial neural network*) oraz głębokie uczenie (*deep learning*). Autor nawiązuje do tych dyskusji, ale stara się przede wszystkim skupić na kulturowych i filozoficznych konsekwencjach rozwoju nowych technologii związanych z AI. Namysł nad relacjami pomiędzy ludzką i nieludzką inteligencją, reinterpretacja Serresowskiej koncepcji pasożytnictwa, przemyślenie zagrożeń, jakie można zarysować w obszarze środowiska człowieka, epistemologii, dominacji algorytmów – to tylko niektóre z podejmowanych zagadnień. Perspektywa kulturoznawczo-filozoficzna, wykorzystująca wiedzę na temat technologiczno-informatycznych aspektów AI, stanowi dla autora rodzaj metodologicznej ramy dla rozpatrywanych zjawisk.

Słowa kluczowe: Artificial Intelligence; nieludzkie inteligencje; Michel Serres; pasożyty; Byung-Chul Han

O autorze

Piotr Zawojski – prof. dr hab. w Instytucie Nauk o Kulturze Uniwersytetu Śląskiego. Wykłada na ASP w Krakowie (Wydział Intermediów). Zajmuje się problematyką fotografii, filmu i kina, sztuki nowych mediów oraz cyberkultury i technokultury. Autor książek: *Elektroniczne obrazy* (2000), *Wielkie filmy przełomu wieków. Subiektywny przewodnik* (2007), *Cyberkultura. Syntopia sztuki, nauki i technologii* (2010, 2018), *Sztuka obrazu i obrazowania w epoce nowych mediów* (2012), *Technokultura i jej manifestacje artystyczne. Medialny świat hybryd i hybrydyzacji* (2016), *Ruchome obrazy zatrzymane w pamięci. Reminiscencje teoretyczne i krytyczne* (2017). Redaktor naukowy *Ku filozofii fotografii Viléma*

Flussera (2004, 2015). Redaktor tomów zbiorowych *Digitalne dotknięcia. Teoria w praktyce / Praktyka w teorii* (2010), *Bio-techno-logiczny świat. Bio art oraz sztuka technonaukowa w czasach posthumanizmu i transhumanizmu* (2015), *Klasyczne dzieła sztuki nowych mediów* (2015); twórca antologii *Teoria i estetyka fotografii cyfrowej* (2017). Członek Komitetu Nauk o Kulturze PAN.

Kulturowe i filozoficzne konteksty ekspansji sztucznej inteligencji

Ten esej poświęcony jest kulturowym konsekwencjom ekspansji sztucznej inteligencji, czego symbolicznym momentem przełomowym stało się zaprezentowanie przez OpenAI modelu GPT w listopadzie 2022 roku. Doczekało się to szeregu omówień i różnych interpretacji, przede wszystkim w odniesieniu do nowatorskich rozwiązań wykorzystujących sztuczne sieci neuronowe (*artificial neural networks*) i głębokie uczenie (*deep learning*). Nawiązując do tych dyskusji, ale staram się przede wszystkim skupić uwagę na szeregu kontekstów związanych z kulturowym wymiarem przemian generowanych przez AI oraz pytaniami z obszaru filozofii mediów. Interesuje mnie głównie problematyka relacji pomiędzy ludzką i nieludzką inteligencją w zakresie epistemologicznym, społecznym, a także technologicznym.

Perspektywa kulturoznawczo-filozoficznego oglądu opisywanych zjawisk nie zwalnia z dbałości o precyzję proponowanych interpretacji, równocześnie jednak uprawomocnia większą swobodę w formułowaniu własnych pomysłów. Te zaś nie układają się w rodzaj systematycznego wywodu, nie taki bowiem był mój cel; są raczej krążeniem wokół kilku kluczowych problemów takich jak dominacja algorytmów w technokulturze, pasożytnictwo systemów generatywnej sztucznej inteligencji, zagrożenia, ale i nadzieje związane z rozwojem AI, i wreszcie relacje człowieka z (inteligentnymi?) maszynami. Warto jednak tylko nieznacznie (wydawałoby się) cofnąć do przeszłości, by uświadomić sobie, jak nawet najodważniejsi w swoich eksperymentach myślowych autorzy mylili się w przewidywaniach dotyczących przyszłości sztucznej inteligencji. To nie tylko przestroga, ale też, paradoksalnie, zachęta do nieskrępowanego wymogami dyskursu naukowego podążania za własnymi intuicjami i hipotezami badawczymi.

Rozpocznę od krótkiego omówienia przypadku radykalnej zmiany w podejściu do kwestii sztucznej inteligencji, którą uosabia postawa Douglasa Hofstadtera. Jego książkę *Gödel, Escher, Bach: an Etemal Golden Braid* opublikowaną w 1979 roku¹ czytałem dawno temu, właściwie zapomniałem, że końcowe jej rozdziały zostały poświęcone sztucznej inteligencji. Przypomniła mi o tym, oraz przede wszystkim o poglądach i prognozach tam zaprezentowanych, Hannah Fry w swojej publikacji zatytułowanej *Hello World. Jak być człowiekiem w epoce maszyn*. Zrobiła to niejako pośrednio, odwołując się do pewnego wydarzenia z 1997 roku, którego zresztą Hofstadter był organizatorem. W University of Oregon odbył się krótki koncert, w trakcie którego zostały zaprezentowane tylko trzy utwory wykonane przez pianistę Winifreda Knera: pierwszym była mało znana kompozycja Jana Sebastiana Bacha, drugim kompozycja Steva Larsona (profesora muzyki na uniwersytecie w Oregonie), trzecim zaś kompozycja stworzona przez algorytm napisany przez Davida Cope'a eksperymentującego z tak zwaną inteligencją muzyczną (*Experiments in Musical Intelligence*, EMI). Publiczność zapytana o to, który z zaprezentowanych utworów jest dziełem Bacha, w zdecydowanej większości uznała, że

to ten napisany przez algorytm². Amerykański myśliciel był nie tylko zdumiony, ale i zaniepokojony. Mówił tak:

Jedynę pocieszenie, jakie mogę w tym momencie odczuwać, pochodzi ze świadomości, że EMI samo w sobie nie generuje stylu. Polega to na naśladowaniu wcześniejszych kompozytorów. Ale to wciąż nie jest duże pocieszenie. W jakim stopniu muzyka składa się z „riffów”, jak mówią ludzie jazzu? Jeśli tak jest, oznaczałoby to, ku mojemu całkowitemu zdruzgotaniu, że muzyka reprezentuje sobą znacznie mniej, niż kiedykolwiek myślałem³.

A przecież w *GEB* (tak określa się najsławniejsze dzieło Hofstadtera) autor z dużym przekonaniem twierdził, że tylko ucieleśnione doświadczenie może być gwarantem i podstawą tworzenia muzyki, która wyrasta z emocjonalnego stosunku kompozytora (ludzkiego) do świata, czerpiącego z własnych doświadczeń i przeżyć.

Co zatem się stało, że raptem kilkanaście lat później, niejako na własne życzenie, jego przekonanie o wyjątkowości ludzkich zdolności twórczych uległo tak istotnej przemianie? Co prawda także w tym samym roku (1997) Deep Blue stworzony przez IBM pokonał genialnego szachowego mistrza Gariego Kasparowa w stosunku 3½: 2½, ale czy świadomość tego faktu oraz wysłuchanie „zimitowanej” przez program komputerowy w stylu Bacha kompozycji, a także jej ocena przez słuchaczy, było tak dramatycznym przeżyciem, które zachwiało wiarę autora *GEB* w wyjątkowość człowieka?

W roku 2018 Hofstadter opublikował krótki tekst poświęcony „płytkości” tłumaczeń dokonywanych przez translator Google'a⁴; dowodził w nim, że jeśli chodzi o proste teksty, a także ultraszybkie przetwarzanie tekstu, program ten spełnia swoje zadania, natomiast nie jest w stanie nawet zbliżyć się do pracy ludzkiego tłumacza, który nie działa mechanicznie, ale rozumie to, co tłumaczy. To powoduje, że nasycy tłumaczony przez siebie tekst zarówno kontekstualnymi odniesieniami wynikającymi ze świadomości zanurzenia każdego tekstu w kulturze, ale i z przekonania, iż nie chodzi o doskonałość wiernego

tłumaczenia, ale przede wszystkim o oddanie „ducha” (tak to nazwę), a nie dosłowne przełożenie tekstu, zwłaszcza jeśli chodzi o twórczość literacką⁵.

Problem maszynowego tłumaczenia przez wiele lat był jednym z istotnych zagadnień rozpatrywanych w kontekście wydajności, ale też jakości pracy systemów sztucznej inteligencji, opinia Hofstadtera z przywoływanego artykułu była dosyć powszechna. Jak jednak miało się wkrótce okazać (raptem kilka lat po opublikowaniu tego artykułu) amerykański intelektualista, podobnie jak w przypadku prognoz dotyczących muzyki czy też szachów, zmienił zasadniczo swoje zdanie, co odbiło się szerokim echem w rozmaitych środowiskach, na co zwrócił uwagę choćby David Brooks⁶. A że uczynił to w „The New York Times”, to sprawa stała się głośna.

Przypadek Hofstadtera traktuję jako zobrazowanie pewnego trendu w postrzeganiu możliwości (i niemożliwości, tych jednak jest coraz mniej) sztucznej inteligencji w całkowitym przeobrażeniu współczesnego świata. Jeśli nawet tak kompetentni, a przy tym sceptyczni wobec myślenia o niej jako o fundamentalnym *game changerze* autorzy nie mają już dziś wątpliwości, że w istocie jest ona właśnie tym czymś, co radykalnie zmienia naszą rzeczywistość, to rodzaj symetriizmu, w którym zastanawiamy się nad pozytywnymi i negatywnymi tego zjawiska, musi ustąpić na rzecz realnej oceny długofalowych konsekwencji zmian kulturowych. Te zaś są nieuniknione i dokonują się na naszych oczach w sposób wykładniczy. Ową zmianę myślenia Hofstadtera można potraktować symbolicznie, jest ona znamienna, bowiem osadzona została w głębokiej znajomości nie tylko tego, czym była i jest historia tworzenia myślących maszyn, ale dotyczy przede wszystkim fundamentalnych zagadnień relacji zachodzących pomiędzy człowiekiem a technologią, czy też technologiami, to zaś zmusza nas do powrotu do podstawowych pytań o istotę człowieczeństwa.

Dokonana przez Hofstadtera rewizja poglądów wywołana została przez fakt, że wszystko to, w co wierzył przez omalże pół wieku, czyli perspektywa powstania niesamowicie wydajnych systemów sztucznej inteligencji, oddalona jest w czasie na przynajmniej kilkadziesiąt lat – dziś jest naszą rzeczywistością. Po raz kolejny się pomylił, co dobitnie pokazały symboliczne porażki człowieka w konfrontacji z „myślącymi maszynami” (Garii Kasparow, Lee Sedol – przegrywający z AlphaGo w 2016 roku). Znacznie bardziej traumatycznym doświadczeniem może jednak być świadomość „ja”, że już wkrótce ludzie mogą być „przyćmieni” przez systemy sztucznej inteligencji. To dziś wydaje się kwestią zasadniczą, dekonstruuje bowiem fundamenty przekonania nie tylko o wyjątkowości człowieka w złożonym systemie nieorganicznych, organicznych i technologicznych bytów, ale też możliwość zachowania *status quo* antropocentrycznego układu zależności. Stąd bierze się przysłgnięcie Hofstadtera zrodzone ze świadomości, że jego (nasze) zdolności w porównaniu z systemami obliczeniowymi są miliardy razy mniej wydajne. To sytuacja bezprecedensowa, absolutny przełom

w historii ludzkości, który symbolicznie porównać można z wynalezieniem ognia czy też koła.

Bagatelizowanie tych faktów, rodzaj stoickiego spokoju albo wyparcia (to dwa modele podejścia do AI) są reakcjami obronnymi. Nie prowadzą do jakichkolwiek racjonalnych działań w zakresie przemyslenia i instytucjonalnego określenia zasad tworzenia, funkcjonowania i regulacji (prawnych, etycznych, politycznych), które należałoby wdrażać wraz z pojawianiem się coraz bardziej złożonych systemów autonomizujących własne działanie. Sprowadzanie ich wyłącznie do roli programów polegających na zwykłym „autouzupełnianiu tekstu” to wyraz niezrozumienia istoty ich działania, które obecnie zbliża się do rozumienia słów i koncepcji, a nawet, w pewnym sensie, zdolne jest do odczuwania emocji, a to jeden z powracających motywów rozprawiania o AI.

Geoffrey Hinton, najważniejsza postać w historii myślenia o tym, by komputery mogły uczyć się uczenia, dostrzega trzy główne obszary zagrożeń związanych z AI: niekontrolowane generowanie *fake newsów*, wyeliminowanie pewnych rodzajów pracy oraz, być może najważniejsze, zagrożenie egzystencjalne. To ostatnie wiąże się z powstaniem

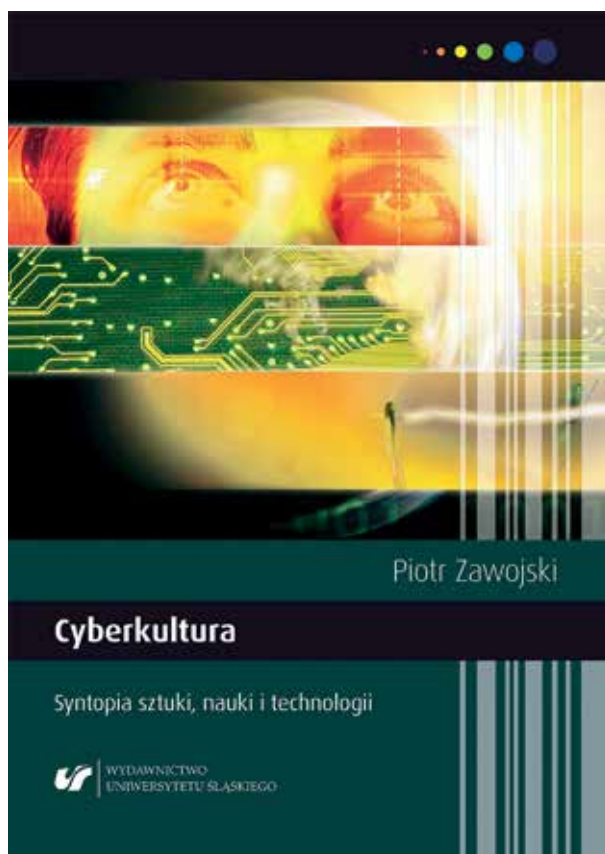
cyfrowej inteligencji, którą rozwijamy. Może [ona] stanowić lepszą, skuteczniejszą formę inteligencji niż mózgi biologiczne. Zawsze myślałem, że AI czy uczenie głębokie może jedynie próbować naśladować działanie naszego mózgu, ale nie może mu dorównać: naszym celem było coraz dokładniejsze upodobnienie maszyn do nas samych. W ciągu ostatnich kilku miesięcy zmieniałem swoje stanowisko. I dziś uważam, że możemy już opracować coś, co będzie znacznie wydajniejsze niż ludzki mózg⁷.

Choćby z tego powodu, że systemy cyfrowe można kopiować, a mózgu biologicznego (póki co) nie.

Nie ulega jednak wątpliwości, że tego typu dylematy lokują się w centrum najważniejszych dziś debat dotyczących nie tylko kwestii sztucznej inteligencji, ale i przyszłości relacji człowieka z nowymi technologiami oraz po prostu przyszłości człowieka. Dwa obszary wydają się tutaj kluczowe, a przy tym należy rozpatrywać je łącznie jako decydujący punkt zwrotny w historii ludzkości. To sztuczna inteligencja oraz biologia syntetyczna rozwijające się w nieokiełznany sposób, co zresztą zwłaszcza w środowiskach szeroko rozumianego high-techu przyjmowane jest z entuzjazmem i programowym wręcz optymizmem, a zarazem zastanawiającym i konsekwentnym niedostrzeganiem zagrożeń. Można to tłumaczyć w sposób najprostszy poprzez konieczność skutecznego zabezpieczania interesów największych graczy w obszarze rynku nowych technologii komputerowych, informatycznych i biologicznych czy też biotechnologicznych. Logika kapitału polega przecież na maksymalizacji zysku, a nie na podcinaniu gałęzi, na której się siedzi.

Ta „nadchodząca fala”, jak ją określa Mustafa Suleyman⁸, wydaje się nie do powstrzymania. W swojej niezwykle przenikliwej analizie „sztucznej inteligencji, władzy i najważniejszego dylematu ludzkości w XXI wieku” (tak brzmi podtytuł jego książki) przedstawia podstawowe problemy i zadania, z jakimi musimy się zmierzyć. Wspomniany już nadmierny optymizm to rewers „awersji do pesymizmu”, wedle określenia Suleymana; podejście to sprawia, że elity nie dostrzegają lub odrzucają narracje, które uznają za przesadnie negatywne, to niejako odnowienie sporu technofobów i technofili, tyle że w nowych okolicznościach technologicznych. Podstawowym zadaniem, jakie powinniśmy realizować, jest zdecydowane „ujęcie w ryzyko technologii”, które wydaje się karkołomne lub wręcz niemożliwe do zrealizowania. Koncepcja „powstrzymywania” polegająca na zdolności do „monitorowania, ograniczania, kontrolowania, a potencjalnie nawet wycofywania z użytku technologii”⁹. I choć, paradoksalnie, rozdział pierwszy książki zatytułowany jest *Powstrzymywanie technologii jest niemożliwe*, to w jego zakończeniu autor pisze, że „na pierwszy rzut oka powstrzymywanie technologii nie wydaje się możliwe. A jednak dla dobra nas wszystkich musi być możliwe”¹⁰. Ten apel wzywający do wstrzymania prac nad sztuczną inteligencją na pół roku – podpisany między innymi przez Elona Muska, Steve’a Wozniaka, Yuvala Noaha Harariego i tysiące innych osób – nie spotkał się z powszechną aprobatą i został zignorowany przez firmy rozwijające prace nad generatywnymi systemami AI, nie doczekał się też jakichkolwiek reakcji ze strony rządów. Sygnatariusze apelu, podkreślając rosnące zagrożenia i ryzyko, odwołali się do głośnych *Zasad AI z Asilomar* (z roku 2017)¹¹, w których sformułowano dwadzieścia trzy reguły obejmujące trzy obszary: badania, etykę i wartości oraz długoterminowe bezpieczeństwo. Prace nad AI powinny być zaplanowane i kontrolowane, obecnie nic takiego nie ma miejsca.

Spotkanie ponad stu badaczy, filozofów, ekonomistów, przedsiębiorców – zorganizowane przez Future of Life Institute w centrum konferencyjnym Asilomar w Pacific Grove w Kalifornii – było niewątpliwie istotnym wydarzeniem zwracającym uwagę opinii publicznej na zagrożenia płynące ze strony bezprzykładnego rozwoju nowych technologii, w tym przede wszystkim AI. Można tylko wyrazić zdziwienie, że w tym gronie nie było twórców kultury ani artystów, co jest zastanawiającym brakiem. Czy jednak te szlachetne przesłanki i cele w jakiś konkretny i wymierny sposób przyczyniły się do zmiany myślenia i regulacji dotyczących koniecznych ograniczeń, jakie powinny być zastosowane wobec tych, którzy tworzą nowe technologie informatyczne? Niestety, jak często bywa w przypadku tego typu słusznych inicjatyw środowisk intelektualnych i naukowych – w gruncie rzeczy nie spowodowało to żadnych praktycznych działań ani ze strony korporacji myślących przede wszystkim o monetyzacji swoich wytworów, ani ze strony rządów, które powinny narzucać ramy prawne i etyczne. Ten rozjazd w ocenie zarówno obecnej



Piotr Zawojski, *Cyberkultura. Syntopia sztuki, nauki i technologii*, Wydawnictwo Uniwersytetu Śląskiego, Katowice 2018.

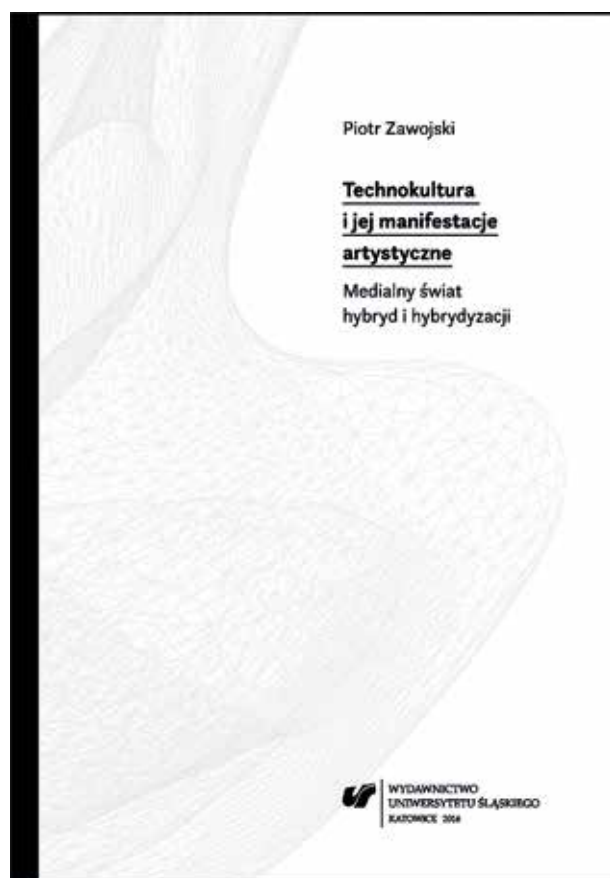
sytuacji, jak i w kreśleniu przyszłościowych scenariuszy rozwoju relacji człowieka oraz sztucznej inteligencji – który wyraża się w hurraoptymizmie przedstawicieli big techów, a więc „bandy czworga” (GAFA) albo „wielkiej piątki” (GAFAM), czyli Google’a (obecnie Alphabet), Amazona, Facebooka (obecnie Meta), Apple’a oraz Microsoftu, krytycznie oceniających dokonujące się zmiany naukowców oraz aktywistów – wydaje się znaczący.

W odróżnieniu od wspomnianego wcześniej i krytycznie ocenianego przez Suleymana nadmiernego optymizmu, w trakcie spotkania w Asilomar pojawił się głos Erika Brynjolfssona, który postulował „świadomy optymizm”, czyli przekonanie, że w omawianej materii wydarzą się dobre rzeczy, ale by tak się stało, „konieczne jest wypracowanie pozytywnej wizji przyszłości”¹² wbrew dominującym tendencjom do ujęć dystopijnych. Taki punkt widzenia jest radykalnie odmienny niż ten prezentowany przez Suleymana.

Pytanie o to, czy algorytmy tylko kompilują dane, którymi zostały nakarmione, są doskonałymi imitatorami, czy też są w stanie stworzyć coś oryginalnego i nowatorskiego, przewija się nieustannie w dyskusjach na temat specyfiki działań systemów generatywnych, co ma szczególne znaczenie w obszarze *AI art* oraz istotnych kwestii związanych z rozpatrywaniem pytań o kreatywność tego typu „narzędzi”. Używam cudzysłowu, bowiem pytanie o to, czy są to tylko narzędzia, czy też systemy zdolne do autonomicznej

kreatywności, pozostaje otwarte. Stanowisko Andrzeja Dragana jest w tym kontekście warte przemyślenia, uważa on bowiem, że maszyny mogą być twórcze, że nie są tylko narzędziami (takimi jak w przeszłości był na przykład aparat fotograficzny czy kamera filmowa).

Wspominam o aparacie, bowiem Dragan poza tym, że naukowo zajmuje się fizyką teoretyczną, jest także artystą, który wypracował własny, oryginalny styl w fotografii portretowej. W jednym z wywiadów podaje przykład odnoszący się do tej twórczości. Otóż w 2005 roku wykonał portret Davida Lyncha trzymającego kurę; dodam, że to niezwykle frapująca fotografia, chociaż samo to określenie może wydawać się nie do końca adekwatne, bowiem finalny efekt został osiągnięty poprzez wykorzystanie cyfrowego instrumentarium postprodukcyjnego, może zatem należałoby mówić po prostu o obrazie cyfrowym, to jednak tylko marginalny przypis. Ważne jest, że zdjęcie bez trudu można znaleźć w sieci, może to uczynić zarówno człowiek, jak i tym bardziej system(y) karmiące się wielkimi zbiorami danych w celu udoskonalania własnej wydajności. Nie ma innych wersji tego obrazu, jest on niepowtarzalny, nie ma mowy o jakiejś serii, wariacjach, kolejnych wersjach, co przecież jest naginą dziś praktyką w obszarze sztuki funkcjonującej w epoce cyfrowej wytwórczości, a co mogło spowodować, że *output* w postaci obrazu musiałby pokrywać się jeden do jednego z materiałem źródłowym.



Piotr Zawojski, *Technokultura i jej manifestacje artystyczne. Medialny świat hybryd i hybrydyzacji*, Wydawnictwo Uniwersytetu Śląskiego, Katowice 2016.

Kiedy Dragan zwrócił się do sztucznej inteligencji, by ta (na podstawie jego prostego promptu) stworzyła obraz „Davida Lyncha trzymającego kurę”, okazało się, że „nie zdarzyło się, aby AI stworzyła obraz, który wykorzystywał choć jeden element tego oryginalnego zdjęcia. Czyli nie kopiowała, lecz stworzyła swoje własne, oryginalne dzieło”¹³. To dla niego przekonujący dowód na to, że maszyny mogą być twórcze, co rzecz jasna jest ciekawym, choć nie do końca przekonującym argumentem. Warto wziąć go jednak pod uwagę.

Inną rzeczą jest jednak wyrażone *expressis verbis* przekonanie, że to pierwszy moment w historii, w którym tworzymy czy też używamy technologii, których istoty działania nie rozumiemy. Pozwalam sobie na wątplenie co do zasadności tego typu rozumowania, idąc zresztą za wskazaniem samego Andrzeja Dragana, który z przekonaniem twierdzi, podążając za sugestią Bertranda Russella (ja też ją podzielam), że podstawową doktryną badawczą i poznawczą powinno być wątplenie¹⁴. Zatem wątpię, że to AI jest tym pierwszym przypadkiem, który dowodzi naszego niezrozumienia tego, jak działają pewne wytworzone przez nas narzędzia. Dragan ma przy tym wielu sprzymierzeńców, którzy rozumują w podobny sposób, trudno im odmówić kompetencji i głębokiej znajomości dyskutowanych zagadnień. Na przykład Jon Kleinberg i Sendhil Mullainathan twierdzą, że

Kolejne generacje algorytmów są coraz inteligentniejsze, ale zarazem stają się coraz bardziej niezrozumiałe. Tymczasem, aby poradzić sobie z maszynami, które myślą, musimy zrozumieć sposób, w jaki to robią. Chyba po raz pierwszy w historii stworzyliśmy maszyny, których działania nie pojmujemy¹⁵.

Łatwo przeoczyć ową niezrozumiałość, tym bardziej że projektujemy maszyny, by myślały tak jak my, i być może tak w istocie jest, ale jednocześnie trzeba założyć, że te systemy wcale nie myślą (jeśli „myślą”, a nie tylko analizują i przeliczają dane) tak jak my. Posługują się one inną logiką, chociaż wytrenowaną dosłownie i w przenośni na nas; pasożytują na naszych danych po to, by zwiększać wydajność własnego działania. Nie do końca jednak przekonujące wydaje się stwierdzenie, że po raz pierwszy stworzyliśmy „rzeczy” (algorytmy, programy, software), których działania nie jesteśmy w stanie zrozumieć, jak ma to miejsce w przypadku maszynowego uczenia głębokiego czy też sztucznych sieci neuronowych. Wielokrotnie powracający w dysputach na temat nowych technologii wątek *black box*, czyli systemów, które przetwarzując surowe dane wejściowe, produkują wartościowe dane wyjściowe, a przy tym ten proces „przemiany” jest całkowicie „nieprzezroczysty” – po prostu nie wiemy, jak to się dzieje – był opisywany zarówno w odniesieniu do konkretnych problemów (na przykład działania ludzkiego mózgu), choć zapewne częściej używano tego określenia w wymiarze metaforycznym.

Wracam do zagadnienia pasożytnictwa systemów karmiących się (czy też karmionych) danymi w celu tworzenia różnych modeli generatywnych przetwarzających dane. GPT oznacza „wstępnie przeszkolony transformator generatywny” (*generative pre-trained transformer*) – to fundament tego typu programów, ale też jeden z kluczowych problemów implikujących pytania o prawa autorskie, ochronę danych osobowych czy idee wspólnej inteligencji sieciowej opartej na zasadzie dzielenia się jednostkowym wkładem do kolektywnej inteligencji. Jej infrastruktura miałaby wynikać z bezinteresownego dzielenia się wiedzą i umiejętnościami, co było jednym z mitów założycielskich wczesnej cyberkultury.

Jak w tym kontekście można byłoby wykorzystać metaforyczne, ale i dosłowne zastosowanie koncepcji pasożytów w odniesieniu do systemów generatywnej sztucznej inteligencji? Już w pionierskim okresie rozwoju sztuki sieciowej (*net art*) pojawiły się intuicje, by rozpatrywać ją w odniesieniu do koncepcji pasożytnictwa. Erik Hobijn i Andreas Broeckman¹⁶ w interesujący sposób starali się pokazać, że wiele z pionierskich projektów *net art*owych opierało się na swoistym pasożytnictwie, które jednak nie było konotowane pejoratywnie. Te wczesne (z początku lat 90. ubiegłego stulecia) rozpoznania dotyczące działań opierających się na zasobach magazynowanych w sieci, na długo przed teoretycznymi opisami baz danych (*big data*), takich jak chociażby analityka kulturowa spod znaku Lva Manovicha¹⁷ – kumulowane i rozwijane w efekcie doprowadziły do wykrystalizowania się dataizmu jako nowej formy ideologii albo „religii”. Termin ten został użyty po raz pierwszy przez Davida Brooksa w 2013 roku, ale zapewne to Yuval Noah Harari w największym stopniu przyczynił się do rozpropagowania tego pojęcia w obszarze refleksji historyczno-filozoficznej¹⁸. Natomiast w obszarze działań artystycznych można byłoby przywołać takich twórców konceptualizujących zagadnienie sztuki danych (*art of data* albo *data art*) jak Albert-László Barabási¹⁹ czy też Ryoji Ikeda, o którego inspirujących dziełach pisałem w innym miejscu²⁰.

Sztuczna inteligencja oparta jest obecnie na wszechstronnym przetwarzaniu danych zgromadzonych w gigantycznych zasobach sieciowych. Logika procesów, w efekcie których możemy używając promptów, tworzyć (?) obrazy, teksty, muzykę czy wideo, opiera się na nieskończonym przeliczaniu. To poruszanie się w obszarze reżimu informacyjnego opartego na dominacji algorytmów, które zmierzają w stronę totalitarnego zawłaszczania naszych poczynań, wyborów (jeśli możemy o czymś takim jeszcze mówić), decyzji – ma wszelkie znamiona zmechanizowanych czynności²¹. Filozoficzna interpretacja takich procesów, odchodząca od języka informacyjnego tłumaczącego te zjawiska na poziomie technologicznym, okazuje się niezwykle istotnym wymiarem refleksji dotyczącej bazowego wymiaru sztucznej inteligencji, czyli zbiorów danych, bez których nie mogłaby ona się rozwijać. Jak pisze Byung-Chul Han:

Wraz z totalitaryzmem żegnamy rzeczywistość daną nam przez pięć zmysłów. Za tym co dane totalitaryzm konstruuje aktualną rzeczywistość, która domaga się szóstego zmysłu. Dataizm natomiast nie potrzebuje szóstego zmysłu. Nie wykracza poza immanencję tego, co dane, czyli danych. Dosłowne znaczenie łacińskiego *datum*, wywodzącego się od *dare* (dawać), oznacza to, co dane. Dataizm nie tworzy innej rzeczywistości poza danymi, jest totalitaryzmem bez ideologii²².

Rozprawianie o pasożytnictwie w kontekście sztucznej inteligencji, a zwłaszcza modeli głębokiego uczenia i sztucznych sieci neuronowych, najczęściej sprowadzane było do negatywnej oceny tego typu praktyk. Już wcześniej zresztą takie interpretacje pojawiały się w odniesieniu do generatywnych programów komputerowych naśladujących sztukę tworzoną przez ludzi. Jon McCormack (i jego współpracownicy) przywołuje w tej kwestii poglądy filozofa Anthony’ego O’Hary²³, który w roku 1995 wskazywał, że wykorzystanie programów komputerowych do kopiowania określonych technik artystycznych czy konwencji w gruncie rzeczy jest praktyką pasożytniczą, a efekty ich działania są problematyczne, jeśli chodzi o ocenę wartości artystycznych i estetycznych.

Twierdzimy, że modele głębokiego uczenia wyszkolone na dużych zbiorach danych, czyli sztuki stworzonej przez człowieka – takich jak systemy TtI [Text-to-Image – dop. P.Z.] – pasożytują na ludzkiej sztuce i kreatywności, ponieważ czerpią i utrzymują się ze sztuki ludzkiej kosztem współczesnej sztuki tworzonej przez ludzi. W biologii, aby związek między gatunkami miał charakter pasożytniczy, pasożyt musi odnosić korzyści kosztem żywiciela. Różni się to od innych relacji takich jak symbioza czy mutualizm, w których obie strony mogą coś zyskać. Bezpośrednia zależność modeli TtI wyszkolonych na zbiorach danych nie wymaga dalszych wyjaśnień, więc aby uzasadnić nasze twierdzenie o pasożytnictwie, musimy wykazać wpływ na sztukę ludzką. Jeśli zatem sztuka ludzka jest umniejszana przez TtI, wówczas tę relację można uznać za pasożytniczą²⁴.

Tyle że rozumienie pasożytnictwa, biorąc pod uwagę reinterpretację tego zjawiska zaproponowaną przez Michela Serresa²⁵, może być poszerzone o bardziej skomplikowane relacje oraz obopólną korzyść, którą można zastosować do opisu niejednoznacznych stosunków człowieka i sztucznej inteligencji. „AI jako pasożyt”, w tradycyjnym rozumieniu pasożytnictwa, to zdecydowanie uproszczona formuła, która zakłada, że tylko jeden z dwóch „organizmów” czerpie korzyści ze współżycia, natomiast drugi ponosi wyłącznie szkody.

Wykorzystując rozważania Serresa, moglibyśmy powiedzieć, że technopasożyty funkcjonujące w obszarze systemów głębokiego uczenia i sztucznych sieci neuronowych przyczyniają się do, po pierwsze, rozjaśnienia

epistemologicznych uproszczeń wynikających z bezkrytycznego przyjmowania standardowych twierdzeń dotyczących tych zagadnień, po drugie zaś, skierowują naszą uwagę ku możliwości nowej, odmiennej od zastanych standardów interpretacji skomplikowanych relacji zachodzących pomiędzy ludzkimi i nieludzkimi inteligencjami. Mówiąc krótko: metaforycznie ujęte pasożytnictwo systemów czerpiących z baz danych to sposób na „użyłizację” (w rozumieniu „odzyskiwania”) doświadczeń ludzkich dostarczycieli danych, ale też możliwość dostrzeżenia swojego „wkładu” w proces nieustannej wymiany danych pomiędzy „gościem” i „gospodarzem”. Według Serresa pasożyt „nie jest złodziejem ani nikczemnikiem: żywiciel stwarza pasożytowi, wprost albo w sposób pośredni, zachęcające warunki”²⁶. To one determinują możliwości fortunnego rozwoju i kooperacji ludzkich i nieludzkich inteligencji. Bez namysłu nad tymi zagadnieniami przełamanie impasu symetrycznego rozprawiania o nadziejach i zagrożeniach związanych z rozwijaniem filozofii i technologii sztucznej inteligencji będzie tylko nieustannym poszukiwaniem jakiegoś dobrze umocowanego, choć raczej niemożliwego do jednoznacznego określenia punktu oparcia dla spekulacji (bo przecież nie teorii) dotyczących nie tylko wpływu AI na nas (ludzi), ale i na świat (ludzki i pozaludzki).

Wyłaniający się nowy porządek świata zdominowanego przez cyfrowe narzędzia to nie jest wizja przyszłości – to nasza rzeczywistość. Można się spierać, czy zaproponowana przez Maxa Tegmarka, współtwórcę wspomnianego wcześniej Future of Life Institute, koncepcja „życia 3.0” jeszcze nie istnieje na Ziemi, jak pisał jeszcze kilka lat temu przed rewolucyjnymi zmianami będącymi konsekwencją eksplozji generatywnej sztucznej inteligencji, czy właśnie jesteśmy świadkami jego narodzin. Szwedzko-amerykański fizyk i kosmolog, ale też błyskotliwy myśliciel, wyróżnił trzy stadia życia: życie 1.0 (biologicznie proste), życie 2.0 (tworzące kulturę) oraz życie 3.0 (technologiczne). To trzecie stadium, w którym „życie projektuje swoje organy i oprogramowanie”, cechuje „zdolność do wykonania każdego zadania kognitywnego przynajmniej równie dobrze jak człowiek”²⁷, co miałyby charakteryzować Ogólną Sztuczną Inteligencję (AGI).

Nie możemy zapominać, że koszty związane z tworzeniem AI są ogromne. To koszty psychologiczne, społeczne, kulturowe, epistemologiczne, ale też te związane z klimatem i ekologią. Najczęściej postrzegamy sztuczną inteligencję jako coś abstrakcyjnego i niematerialnego, by jednak mogła ona funkcjonować i się rozwijać, potrzebne są olbrzymie ilości energii pozyskiwanej z węgla, ropy, litu (to „biała ropa”, a przy tym „krwawy minerał”) służących do wytwarzania baterii w smartfonach, laptopach i tabletach. Zużywa się przy tym niesłychane ilości wody, eksploatuje tanią siłę roboczą pracującą często w nieludzkich (*nomen omen*) warunkach. To współczesny proletariatus, czy raczej klasa niewolnicza w takich krajach jak Iran, Chiny, Boliwia czy Mongolia, która niewiele ma wspólnego

z konsumtariatem państw Zachodu czy uprzywilejowanej globalnej Północy. Mówiąc krótko: sztuczna inteligencja jest ucieleśniona i na wskroś materialna, i żeby mogła zaistnieć, musiała rozwinąć się na wielką skalę „działalność wydobywcza”. Wnikliwie i przekonująco analizuje te procesy Kate Crawford²⁸. Kluczowym pojęciem okazuje się ekstrakcja, czyli wydobywanie, zarówno materialnych (paliwa kopalne), jak i niematerialnych (dane) surowców; bez nich rozwijanie sztucznej inteligencji byłoby niemożliwe. Jak tłumaczy badaczka:

W tej książce twierdzą, że sztuczna inteligencja nie jest ani sztuczna, ani inteligentna. Sztuczna inteligencja jest ucieleśniona i materialna, powstaje z naturalnych zasobów, paliwa, ludzkiej pracy, infrastruktury, logiki, historii i klasyfikacji. Systemy AI nie są autonomiczne, racjonalne ani zdolne do rozpoznawania czegokolwiek bez kosztownego obliczeniowo trenowania na wielkich zbiorach danych lub z określonymi z góry zasadami i nagrodami²⁹.

Konsekwencje tych procesów układają się w siatkę trzech głównych zagrożeń dla współczesnego świata – to kwestie środowiska, epistemologii i demokratycznych wyborów, połączonych zresztą szeregiem zależności. Sztuczną inteligencję, jako rodzaj nowego medium w szerokim sensie tego słowa, należałoby ulokować w tym obszarze, o którym pisał niegdyś Jussi Parikka³⁰. Zwracał on uwagę na to, że McLuhanowska koncepcja mediów jako ekstensji ludzkich zmysłów musi być zastąpiona myśleniem w innych kategoriach – media stają się ekstensją Ziemi, to z niej bowiem, czy też jej eksploatacji, wywodzą się nowe technologie. By mogły się one rozwijać, by mogły przybierać materialną postać, ale też funkcjonować na poziomie niematerialnych algorytmów, potrzebny jest „przemysł wydobywczy”, czyli czerpanie z zasobów geologicznych Ziemi. To zaś rodzi szereg skutków, których najczęściej nie dostrzegają, albo starają się je ukrywać, wielcy gracze przemysłu informatycznego. Nie przez przypadek Nvidia właśnie stała się najwyżej wycenianą firmą na świecie, pozostawiając w tyle dotychczasowych czempionów takich jak Microsoft, Google czy Apple. Dlaczego? Bo jest potentatem w produkcji procesorów graficznych, bez których niemożliwe jest progres firm funkcjonujących w domenie sztucznej inteligencji. To Nvidia – w osobie ekscentrycznego prezesa Jensa Hena Huang – w istocie sprawuje obecnie władzę nad światem sztucznej inteligencji, bo to ona może dyktować warunki i wyznaczać priorytetowych odbiorców swojego produktu.

A przecież ponad tymi kwestiami, wywodzącymi się w gruncie rzeczy z rezerwuaru spraw ekonomicznych i biznesowych, związanych z finansowym i gospodarczym wymiarem rozwoju AI, warto zadawać pytania fundamentalne, dotyczące kulturowych i filozoficznych konsekwencji zmian, jakie dokonują się tu i teraz, ale też dzięki naszej świadomej i nieświadomej partycypacji. W końcu to

każdy z nas jest dawcą danych, to każdy z nas karmi ten wszytkożerny system absorbujący i przetwarzający każdy cyfrowy ślad, który pozostawiamy w sieci(ach) internetowych.

Nie można zatrzymywać się na poziomie technologii i infrastruktury – należy też uwzględniać daleko idące konsekwencje ontologiczne i epistemologiczne, tak jak to czyni Byung-Chul Han. Odnajduję w nim spadkobiercę i kontynuatora zarazem tego sposobu rozprawiania o nowych technologiach, które zapoczątkował Vilém Flusser; on sam zresztą przyznaje się do tych inspiracji, nieraz cytując czeskiego myśliciela. W jego tekstach nieustannie natrafiamy na zdania, które można potraktować jako celne i precyzyjne, choć jednocześnie wieloznaczne i czasem omalże poetyckie wersy. Brzmiały one jak filozoficzna aforystyka. Odwołuję się w tym miejscu do jego książki zatytułowanej *Nie(do)rzeczy*, w której w syntetyczny sposób kreśli współczesną trajektorię odchodzenia od rzeczy materialnych do nie(do)rzeczy funkcjonujących jako niematerialne byty w przestrzeni cyfrowej. To znaczy, że „naszą obsesją nie są już rzeczy, ale informacje i dane”³¹.

Proponując takie pojęcia jak „infomania” czy też „infomaty” (to rzeczy nieustannie przetwarzające informacje i przeniknięte informacją; dziś głównym infomatem jest rzecz jasna smartfon), posługuje się też pojęciem „infor-ga” przejętym z propozycji teoretycznych Luciano Floridiego. „Infor-gi”, czyli „organizmy informacyjne”, to ludzie „czwartej rewolucji przemysłowej”, czasu infokracji, w którym nowe technologie w sposób organiczny splatają świat fizyczny z cyfrowym i biologicznym.

Dziś nie ma wątpliwości, że sztuczne inteligencje odgrywają i w coraz większym stopniu będą odgrywać kluczową rolę. Dodam tylko, że wedle Floridiego transhumanistyczne wizje „biotechnologicznych transformacji” człowieka, jego genetycznych modyfikacji, nie są w tym zakresie najważniejsze; o wiele bardziej znaczący jest fakt nasycania nas informacjami. Jak pisze:

Zaczęliśmy traktować siebie jako infor-gów, nie poprzez biotechnologiczne modyfikacje naszych ciał, ale bardziej poważnie i realistycznie poprzez radykalną transformację naszego środowiska i działających w nim agentów³².

Czy za takiego agenta można uznać sztuczną inteligencję? Jak najbardziej, a przy tym w wielorakim sensie tego słowa, bo nie ma jednoznacznej jego definicji. To, co nasuwa się w pierwszej kolejności, to takie cechy jak autonomiczność, komunikatywność, percepcja, ale też zdolność do wykorzystywania wiedzy, używania abstrakcji i symboli, adaptacji w celu osiągnięcia określonych celów, uczenia i komunikowania się w języku naturalnym, przeprowadzania operacji w czasie rzeczywistym, by przywołać tylko kilka najważniejszych atrybutów agenta³³. Sztuczne inteligencje mogą być uznane za metaagenty działające w pewnych środowiskach, zdolne do komunikowania się

i podejmowania autonomicznych decyzji. To w gruncie rzeczy infor-gi przypominające infor-gi ludzkie.

Kreśląc filozoficzną różnicę pomiędzy inteligencją ludzką i nieludzką, Han zwraca uwagę na to, że myślenie jest procesem analogowym, zaś „sztuczna inteligencja nie myśli, bo nigdy nie jest poza sobą”; ludzkie myślenie fundowane jest w oparciu o „fenomenologię nastroju” (odwołuje się tutaj do Heideggerowskich konotacji, wywodzących myślenie z *pathos*, czyli przekonania o absolutnej pewności poznania). Czy prowadząc konwersację z chatem GPT mamy jakąkolwiek pewność, że nie jesteśmy świadkami jego halucynacji/konfabulacji? „*Pathos* jest początkiem myślenia. Sztuczna inteligencja jest *apathic*, czyli bez *pathos*, bez *passion*. Ona *racjonalna*”³⁴. To oczywiste, że istnieje różnica pomiędzy rachowaniem a myśleniem, ale jednocześnie kilka lat, które upłynęły od napisania tych słów w epoce przed pojawieniem się generatywnych sztucznych inteligencji, zmusza nas do ponownego przemyślenia tak jednoznacznie formułowanych sądów. Han w efektowny sposób puentykuje swój esej: „Sztuczna inteligencja nie jest zdolna do myślenia, ponieważ nie może być *faire l'idiot*”³⁵. Jest zbyt inteligentna, by być idiotą”³⁶.

Jeśli zgodzimy się, że myślenie ludzkie (i nieludzkie) to coś więcej aniżeli tylko obliczanie i rozwiązywanie problemów, ale też skłonność do tego, by halucynować czy też konfabulować, to właśnie teraz, po raptem kilku tylko latach rozwoju nowych form sztucznych inteligencji, ten argument wydaje się nieaktualny. Sztuczna inteligencja, tak jak ludzie, wykazuje dobitnie, że „może być idiotą”. Czyżby to właśnie miało być potwierdzeniem jej upodabniania się do człowieka?

Przypisy

- ¹ Zob. Douglas R. Hofstadter, *Gödel, Escher, Bach: an Eternal Golden Braid*, Basic Books, New York 1979.
- ² Hannah Fry, *Hello World. Jak być człowiekiem w epoce maszyn*, przeł. Sebastian Musielak, Wydawnictwo Literackie, Kraków 2019, s. 18–19.
- ³ George Johnson, *Undiscovered Bach? No, a Computer Wrote It*, „The New York Times”, 11.11.1997; <https://www.nytimes.com/1997/11/11/science/undiscovered-bach-no-a-computer-wrote-it.html>, dostęp: 25.06.2025.
- ⁴ Zob. Douglas Hofstadter, *The Shallowness of Google Translate*, „The Atlantic”, 30.01.2018; <https://www.theatlantic.com/technology/archive/2018/01/the-shallowness-of-google-translate/551570/>, dostęp: 25.06.2025.
- ⁵ Douglas R. Hofstadter, *Le Ton beau de Marot: In Praise of the Music of Language*, BasicBooks, New York 1997.
- ⁶ David Brooks, *Human Beings Are Soon Going to Be Eclipsed*, „The New York Times”, 13.07.2023; www.nytimes.com/2023/07/13/opinion/ai-chatgpt-consciousness-hofstadter.html, dostęp: 25.06.2025.
- ⁷ Manuel G. Pascual, *Laureat Nagrody Turinga: AI to nasza przyszłość. O ile ją poskromimy*, przeł. Łukasz Grzymisławski, „Gazeta Wyborcza”, 12.05.2024; <https://wyborcza.pl/7,179012,29745711,laureat-nagrody-turinga-ai-to-nasza-przyszlosc-o-ile-ja-poskromimy.html>, dostęp: 25.06.2025.

- Zob. też: Joshua Rothman, *Dlaczego ojciec chrzestny AI obawia się własnego dzieła*, przeł. Katarzyna Mojowska, „Pismo. Magazyn opinii” 2024, nr 5, s. 54–68.
- 8 Mustafa Suleyman, Michael Bhaskar, *Nadchodząca fala. Sztuczna inteligencja, władza i najważniejszy dylemat ludzkości w XXI wieku*, przeł. Justyn Hunia, Szczeliny, Kraków 2024. Omówienie i komentarz tej publikacji prezentuje Piotr Cieśliński, *Co się śni AI*, „Magazyn Gazety Wyborczej”, 8–9.06.2024, s. 8–11.
 - 9 Mustafa Suleyman, *Nadchodząca fala...*, dz. cyt., s. 14.
 - 10 Tamże, s. 40.
 - 11 Zob. *Asilomar AI Principles*, <https://futureoflife.org/open-letter/ai-principles/>, dostęp: 25.06.2025. Polskie tłumaczenie tego dokumentu można znaleźć w: Max Tegmark, *Życie 3.0. Człowiek w erze sztucznej inteligencji*, przeł. Tomasz Krzysztos, Prószyński i S-ka, Warszawa 2019, s. 422–425. Zob. też: Maciej Garbowski, *A Critical Analysis of the Asilomar AI Principles*, „Zeszyty Naukowe Politechniki Śląskiej. Seria Organizacja i Zarządzanie” 2018, z. 115, s. 45–55.
 - 12 Max Tegmark, *Życie 3.0...*, dz. cyt., s. 4, 28.
 - 13 Piotr Cieśliński, *Prof. Andrzej Dragan: Po raz pierwszy w historii nauczylimy się robić rzeczy, których nie rozumiemy*, „Gazeta Wyborcza”, 7.05.2024; <https://wyborcza.pl/7,75400,30547020,prof-andrzej-dragan-ai-jest-teraz-na-poziomie-osmiolatka.html?token=CCDHMNYfayUZtZ-4bSJpYbZqx4c68-v-Q32eOWUjZWA>, dostęp: 25.06.2025.
 - 14 Andrzej Dragan, *Kwantechizm 2.0, czyli klatka na ludzi*, Wydawnictwo Otwarte, Kraków 2022, s. 13.
 - 15 Jon Kleinberg, Sendhil Mullainathan, *Konstruujemy je, ale ich nie rozumiemy*, przeł. Kurhaus Publishing, GETIT, K. Kubala, [w:] *A ty, co sądzisz o myślących maszynach? Wizje przyszłości wybitnych umysłów ery sztucznej inteligencji*, red. John Brockman, Wydawnictwo Naukowe PWN, Warszawa 2020, s. 79.
 - 16 Eric Hobijn, Andreas Broeckmann, *Techno-Parasites: Bringing the Machinic Unconscious to Life*, <https://v2.nl/articles/techno-parasites>, dostęp: 25.06.2025; Andreas Broeckmann, *Sztuka sieci, maszyny i pasożyty*, [w:] *Media Art Biennale. WRO 97*, red. Piotr Krajewski, Violetta Kutlubasis-Krajewska, Open Studio, Wrocław 1997, s. 25–35.
 - 17 Lev Manovich, *Cultural Analytics*, The MIT Press, Cambridge, MA–London 2020.
 - 18 Zob. Yuval Noah Harari, *Homo deus. Krótka historia jutra*, przeł. Michał Romanek, Wydawnictwo Literackie, Kraków 2018, s. 467–504.
 - 19 Zob. Albert-László Barabási, *Why the World Needs 'Dataism,' the New Art Movement That Helps Us Understand How Our World Is Shaped by Big Data*, <https://news.artnet.com/art-world/archives/introducing-dataism-2181005>, dostęp: 25.06.2025.
 - 20 Zob. Piotr Zawojski, *Technokultura i jej manifestacje artystyczne. Medialny świat hybryd i hybrydyzacji*, Wydawnictwo Uniwersytetu Śląskiego, Katowice 2016, s. 216–253.
 - 21 Pisałem o tym szerzej w innym miejscu. Zob. Piotr Zawojski, *Revolucja algorytmiczna. We władzy cyfr i liczb*, „Opcje” 2018, nr 4, s. 14–21.
 - 22 Byung-Chul Han, *Infocracy. Digitalization and the Crisis of Democracy*, Polity Press, Cambridge–Medford 2022, s. 9.
 - 23 Anthony O'Hara, *Art and Technology: An Old Tension*, „Royal Institute of Philosophy Supplement” 1995, nr 38, s. 143–158.
 - 24 Jon McCormack, Camilo Cruz Gambardella, Nina Rajcic, Stephen James Krol, Maria Teresa Llano, Meng Yang, *Is Writing Prompts Really Making Art?*, [w:] *Artificial Intelligence in Music, Sound, Art and Design. 12th International Conference, EvoMUSART 2023 Held as Part of EvoStar 2023 Brno, Czech Republic, April 12–14, 2023 Proceedings*, red. Colin Johnson, Nereida Rodríguez-Fernández, Sérgio M. Rebelo, Springer, Cham 2023, s. 202.
 - 25 Zob. Michel Serres, *The Parasite*, The Johns Hopkins University Press, Baltimore–London 1982.
 - 26 Andreas Broeckmann, *Sztuka sieci...*, dz. cyt., s. 31.
 - 27 Max Tegmark, *Życie 3.0...*, dz. cyt., s. 58.
 - 28 Zob. Kate Crawford, *Atlas sztucznej inteligencji. Władza, pieniądze i środowisko naturalne*, przeł. Tadeusz Chawziuk, Wydawnictwo Uniwersytetu Jagiellońskiego, Kraków 2024. Zob. też komentarz autorki do tej publikacji: Piotr Cieśliński, *AI ma ciało*, „Magazyn Gazety Wyborczej”, 25–26.05.2024, s. 12–13.
 - 29 Kate Crawford, *Atlas...*, dz. cyt., s. 17.
 - 30 Jussi Parikka, *A Geology of Media*, University of Minnesota Press, Minneapolis–London 2015.
 - 31 Byung-Chul Han, *Nie(do)rzeczy*, przeł. Marcin J. Leszczyński, Wydawnictwo Uniwersytetu Łódzkiego, Łódź 2023, s. 19.
 - 32 Luciano Floridi, *The Fourth Revolution. How the Infosphere Is Reshaping Human Reality*, Oxford University Press, Oxford 2014, s. 96.
 - 33 Zob. [https://pl.wikipedia.org/wiki/Agent_\(programowanie\)](https://pl.wikipedia.org/wiki/Agent_(programowanie)), dostęp: 25.06.2025.
 - 34 Byung-Chul Han, *Nie(do)rzeczy*, dz. cyt., s. 71.
 - 35 Han odwołuje się tutaj do stwierdzenia Gilles’a Deleuze’a, że filozofia zaczyna się od „robienia z siebie idioty” (*faire l'idiot*); idiotyzm jest cechą charakterystyczną myślenia, nawet więcej: każdy filozof, który proponuje nowy język i nową myśl, jest idiotą. Ten wątek pojawił się już we wcześniejszej publikacji Byung-Chul Hana *Psychopolitics. Neoliberalism and New Technologies of Power*, Verso, London–New York 2017, s. 81–87.
 - 36 Byung-Chul Han, *Nie(do)rzeczy*, dz. cyt., s. 77.