

Katarzyna Cieślak 

Akademia Górniczo-Hutnicza im. Stanisława Staszica w Krakowie

Seweryn Rudnicki 

Akademia Górniczo-Hutnicza im. Stanisława Staszica w Krakowie



ADOPCJE I REKONFIGURACJE. SOCJOLOGICZNE UJĘCIE WYKORZYSTANIA GENERATYWNEJ SZTUCZNEJ INTELIGENCJI W BADANIACH SPOŁECZNYCH

Celem tego artykułu jest przyjrzenie się możliwościom i problemom związanym z wykorzystaniem technologii generatywnej sztucznej inteligencji w badaniach społecznych oraz zaproponowanie socjologicznej ramy pojęciowej do analizy zachodzących w tym obszarze przemian. Tekst rozpoczynamy od przeglądu aktualnie dostępnych zastosowań sztucznej inteligencji (ang. *artificial intelligence*, AI) w prowadzeniu badań społecznych. Następnie dokonujemy ich krytycznej analizy – przede wszystkim z uwagi na ryzyka, jakie ich zastosowanie niesie dla procesów tworzenia wiedzy o świecie społecznym i właściwości wytwarzanej wiedzy. Wreszcie, opierając się na współczesnej generacji teorii praktyk społecznych oraz modelu wielopoziomowej transformacji socjo-technologicznej zaproponowanemu przez Franka Geelsa, pokazujemy, że wykorzystanie narzędzi AI może być rozumiane nie tyle jako ich proste „użycie”, ile – trafniej – jako wieloaspektowe, otwarte i dynamiczne „adoptowanie”. Jak argumentujemy, efekty tych procesów nie są determinowane samą technologią, ale mimo to mogą się wiązać ze znaczną rekonfiguracją praktyk badawczych oraz szerszych układów społeczno-technologicznych, których te praktyki są częścią.

Słowa kluczowe: sztuczna inteligencja; generatywna AI; socjologia cyfrowa; badania społeczne; teoria praktyk

Adoptions and Reconfigurations: A Sociological Perspective on the Use of Generative Artificial Intelligence in Social Research

The aim of this article is to explore the possibilities and implications of using digital tools based on generative artificial intelligence technologies in social research. We begin by providing an overview of the currently available applications of AI in social research.

Katarzyna Cieślak, Akademia Górniczo-Hutnicza im. Stanisława Staszica w Krakowie, kcieslak@agh.edu.pl, ORCID 0000-0002-5539-4313; Seweryn Rudnicki, Akademia Górniczo-Hutnicza im. Stanisława Staszica w Krakowie, sewrud@agh.edu.pl, ORCID 0000-0002-6443-3492. Tekst opublikowany na warunkach licencji Creative Commons Uznanie autorstwa-Użycie niekomercyjne-Bez utworów zależnych 3.0 Polska (CC BY-NC-ND 3.0 PL).

Badania, których efektem jest ten tekst, zostały sfinansowane z grantu Narodowego Centrum Nauki, Grant Preludium, 2022/45/N/HS6/02725 oraz grantu Narodowego Centrum Nauki, Grant Opus, 2018/29/B/HS6/02145.

Next, we critically analyze these applications, focusing primarily on the risks they pose to the characteristics of the knowledge produced. Finally, drawing on contemporary social practice theory and the multilevel socio-technological transformation model proposed by Frank Geels, we demonstrate that the use of AI tools should be understood not merely as straightforward “application,” but rather as a multifaceted, open-ended, and dynamic process of “adoption.” We argue that the outcomes of these processes are not solely determined by the technology itself; however, they can lead to a reconfiguration of research practices and the broader socio-technological frameworks in which these practices exist

Key words: social research; practice theory; artificial intelligence; generative AI; digital sociology

Wprowadzenie

W listopadzie 2024 roku na stronach serwisu arXiv opublikowano artykuł autorstwa zespołu kierowanego przez Joong Sung Parka, doktoranta informatyki z Uniwersytetu Stanforda, zatytułowany *Generative Agent Simulations of 1,000 People*, a opisujący wyniki pewnego interesującego eksperymentu (Park i in., 2024). Badanie rozpoczęło się od wywiadów pogłębionych z około tysiącem dorosłych mieszkańców Stanów Zjednoczonych, które przeprowadził głosowy agent AI. Transkrypcje z tych wywiadów, przetworzone przez GPT-4o, posłużyły następnie do opracowania osobnych agentów AI dla każdej z osób uczestniczących w badaniu. Ludzkich uczestników po wywiadzie poproszono o wypełnienie kwestionariusza Generalnego Sondażu Społecznego i testu osobowości oraz wzięcie udziału w serii klasycznych gier społecznych (np. dylemat więźnia); te same zadania postawiono też przed symulującymi zachowanie ludzi agentami AI. Porównanie odpowiedzi udzielonych w pytaniach sondażu społecznego przez ludzi i odpowiadających im agentów AI pokazało zgodność na poziomie 68,9% – taka frakcja odpowiedzi udzielonych przez AI pokrywała się z odpowiedziami ludzkich uczestników. Po korekcie uwzględniającej zmianę w wynikach uzyskanych na ludzkiej próbie na przestrzeni dwóch tygodni współczynnik dokładności odpowiedzi udzielonych w pytaniach sondażowych przez agentów AI i ludzi wyniósł 0,85; wyniki uzyskane w testach osobowości korelowały na poziomie 0,95, a w grach 0,66. Jak stwierdzono w artykule: „Praca ta stanowi podstawę dla nowych narzędzi, które mogą pomóc w badaniu indywidualnych i zbiorowych zachowań” (Park i in., 2024, s. 1).

W chwili gdy piszemy te słowa, tworzenie „krzemowych” populacji respondentów/ek na podstawie danych uzyskanych z wywiadów pogłębionych jest rzadkością – w podobnych do omówionego wyżej badaniach zazwyczaj testuje się modele, w których osoby respondentów są oparte na danych statystycznych (np. ze spisu powszechnego) połączonych z dużymi modelami językowymi

(ang. *large language model*, LLM). W marcu 2025 roku, również w serwisie arXiv, pojawił się tekst przedstawiający rezultaty badania, w którym porównano faktyczne wyniki wyborów w Stanach Zjednoczonych z lat 2016, 2020 i 2024 z odpowiedziami udzielonymi przez takich właśnie „sztucznych” respondentów/ki (Li i in., 2025). Okazało się, że wyniki generowane przez personsy w dużym stopniu oparte na AI przewidywały zwycięstwo demokratycznych kandydatów/ek we wszystkich stanach. Podobne ideologiczne zniekształcenia (w stronę lewicową i progresywną) obserwowano, zadając pytania dotyczące klimatu, technologii, edukacji czy konsumpcji. Przykładowo, AI faworyzował wybór samochodu „droższego, ale przyjaznego dla środowiska” względem „tańszego, ale niemającego tylu przyjaznych dla środowiska funkcji” (Li i in., 2025, s. 8). Im większy był udział LLM w tworzeniu personsy respondenta/respondentki, tym podobne zniekształcenia okazywały się silniejsze.

Przytaczając te przykłady, nie staramy się zabrać głosu w dyskusji nad jakością syntetycznych, powstałych przy użyciu technologii AI symulacji ludzkich odpowiedzi. Próbuje jedynie wskazać, jak głęboki i problematyczny wpływ może mieć trwający obecnie rozwój technologii AI na sposoby badania świata społecznego. Zgodnie z niektórymi stwierdzeniami: „postępy w dziedzinie sztucznej inteligencji (AI), zwłaszcza w obszarze dużych modeli językowych (LLM), dramatycznie wpływają na badania w naukach społecznych” (Grossman i in., 2023), a „narzędzia te mogą spowodować postęp w skali, zakresie i prędkości badań w naukach społecznych i (...) umożliwić nowe formy badania naukowego” (Bail 2024, s. 121). Nawet jeśli podobne stwierdzenia brzmią ogólnikowo i emfaticznie, to faktem jest zauważalne w ostatnich latach rozpowszechnienie generatywnej sztucznej inteligencji w wielu dziedzinach nauki, w tym naukach społecznych (Oxford University Press, 2024). Po części chodzi tutaj o powstanie nowych możliwości badawczych, ale także – jak zauważa wiele głosów krytycznych – o wyłonienie się nowych problemów. Jednak obserwowane zmiany można traktować również jako złożone i wieloaspektowe transformacje zachodzące w praktykach badawczych oraz w szerszych układach tworzenia wiedzy o świecie społecznym. Właśnie tę ostatnią perspektywę proponujemy w tym tekście.

Naszym celem w tym artykule jest – po pierwsze – zwrócenie uwagi na szerokie możliwości zastosowania technologii opartych na generatywnej sztucznej inteligencji (ang. *Generative AI*, GenAI) dla tworzenia wiedzy o świecie społecznym, a po drugie – wskazanie ramy teoretycznej, dzięki której te procesy można opisywać i wyjaśniać w sposób *stricte* socjologiczny, a jednocześnie wykraczający poza skrajności technooptymizmu i technofobii. Niniejszy artykuł nie jest zatem tekstem przedstawiającym wyniki badań własnych, ale klasycznym *position paper*, w którym obficie czerpiemy z istniejących badań, argumentujemy i uzasadniamy nasze stanowisko.

W warstwie teoretycznej artykuł sytuuje się na pograniczu socjologii cyfrowej oraz studiów nad nauką i technologią (ang. *science and technology studies*, STS). Proponowana analiza przecina się w wielu miejscach z zagadnieniami omawianymi w ramach socjologii cyfrowej, ponieważ dotyczy takich tematów, jak przenikanie się społeczeństwa i technologii, tworzenie i wykorzystanie danych cyfrowych, zniekształcenia algorytmiczne, społeczne konsekwencje użycia technologii cyfrowych czy władza cyfrowa. Również przyjmowana w tym tekście perspektywa – analityczna i krytyczna, ale nie normatywna (technofobiczna albo technoentuzjastyczna) – wydaje nam się zgodna z duchem tej subdyscypliny. Ze studiów nad nauką i technologią czerpiemy z kolei zainteresowanie wiedzą, a w szczególności wiedzą o świecie społecznym – nie jej metodologiczną poprawnością, ale społecznymi sposobami jej tworzenia oraz legitymizacji (Collins, 1994; Osborne i Rose, 1999; Law, 2009; Afeltowicz i Pietrowicz, 2013; Marres i in., 2018; Rudnicki, 2023). Pod pojęciem wiedzy o świecie społecznym rozumiemy różnorodne deskryptywne, analityczne lub normatywne wypowiedzi na temat zachowań, stanów i zdolności ludzkich jednostek oraz zbiorowości (Camic i in., 2011, s. 3) powstające w zinstytucjonalizowanych kontekstach, które są społecznie rozpoznawane i uregulowane jako obszary tworzenia takiej wiedzy. Dodajmy jeszcze, że łączenie socjologii cyfrowej i podejścia STS ma w socjologii już dość długą tradycję (Ruppert i in., 2013; Lupton, 2014; Marres, 2017).

Cele stawiane w tym artykule realizujemy w trzech krokach. Po pierwsze, omawiamy możliwości aktualnie dostępnych narzędzi AI przeznaczonych do użycia w prowadzeniu badań społecznych, czyli pokazujemy, do czego te narzędzia mogą być (zdaniem ich badaczy/ek i popularyzatorów/ek) stosowane. Po drugie, dokonujemy krytycznej analizy tych możliwości – przede wszystkim z punktu widzenia ryzyk, jakie niesie zastosowanie narzędzi AI dla procesów tworzenia wiedzy o świecie społecznym i charakteru uzyskiwanej w ten sposób wiedzy. W tych częściach artykułu nasza praca ma przede wszystkim charakter przeglądowy – staramy się zaprezentować główne wątki toczącej się dyskusji na temat roli narzędzi AI w badaniach społecznych. W trzeciej części artykułu proponujemy ramę heurystyczną do interpretacji zachodzących transformacji, w naszym przekonaniu wychodzącą poza dominujące obecnie ujęcia. W duchu socjologii cyfrowej i badań z kręgu STS staramy się przy tym uniknąć zarówno interpretacji instrumentalistycznych, traktujących technologię jako neutralne narzędzie, jak i ujęć deterministycznych, przypisujących technologii zunifikowany i przemożny wpływ na świat społeczny. Szukając innych narzędzi analitycznych, opieramy się na współczesnej generacji teorii praktyk społecznych oraz modelu wielopoziomowej transformacji socjo-technologicznej autorstwa Franka Geelsa. Na tej podstawie proponujemy, że wykorzystanie narzędzi AI w badaniach społecznych może być trafniej zrozumiane nie tyle jako ich proste

„użycie”, ile jako wieloaspektowe, otwarte i dynamiczne „adoptowanie” wiążące się z rekonfiguracją praktyk badawczych.

Zastosowania AI w badaniach społecznych

Pojęcie sztucznej inteligencji ewoluuje od lat 50. XX wieku. W fundamentalnym dla badań w tej dziedzinie, a prezentowanym w 1955 roku na konferencji w Dartmouth projekcie badawczym jako cel badań nad sztuczną inteligencją wskazano opracowanie systemów zdolnych do posługiwania się językiem naturalnym, tworzenia pojęć abstrakcyjnych, a także rozwiązywania problemów typowo ludzkich (McCarthy i in., 2006). Osoby uczestniczące w konferencji zakładały również, że systemy AI będą miały zdolność do samouczenia się¹. Rozwój dziedziny AI, który nastąpił w kolejnych dekadach, bywa metaforycznie porównywany do cyklu pór roku. Okresy, w których można zaobserwować zaawansowanie technologii, w tym rozwój konkretnej metody AI, określa się mianem „wiosen AI” (Maclure, 2020), a fazy stagnacji „zimami AI” (Floridi, 2020). Obecnie dominujący ogólnodostępny i komercyjny charakter zastosowań modeli LLM (w tym włączenie ich do automatyzacji funkcjonalności w programach do analizy wyników badań społecznych) może sugerować, że ponownie znajdujemy się w przełomowym, „wiosennym” momencie.

Warto podkreślić, że AI to termin parasolowy obejmujący różnorodne technologie (Rip i Voß, 2019; Joyce i in., 2021), które zawierają w sobie obietnicę odzwierciedlania ludzkiej komunikacji, w tym potencjał do naśladowania rozumowania. W niniejszym artykule stosujemy socjologicznie ukierunkowaną definicję sztucznej inteligencji jako „podejścia mającego na celu naśladowanie i wspomaganie ludzkiego poznania, podejmowania decyzji i ekspresji, które stosuje różne techniki obliczeniowe do masowych danych w celu opracowania i wdrożenia algorytmów i modeli w kontekstach społecznych, potencjalnie w sposób niezauważalny i ukryty dla opinii publicznej” (Law i McCall, 2024, s. 2–3). Z kolei generatywna sztuczna inteligencja odnosi się do technik przetwarzania danych, za pomocą których na podstawie danych uczących tworzone są „pozornie nowe i znaczące treści” (Feuerriegel i in., 2024, s. 111) takie jak tekst, obraz lub dźwięk. Chociaż technologie GenAI nie są w stanie rozumować *per se*, to symulują inteligentne reakcje poprzez szacowanie najbardziej prawdopodobnych i preferowanych odpowiedzi, a ponadto mogą wytwarzać

¹ Obecnie większość komercyjnych modeli LLM (np. GPT, Gemini) jest trenowana w trybie samonadzorowanym (zob. Rani i in., 2023), który jest formą pośrednią między uczeniem nadzorowanym a nienadzorowanym. Oznacza to, że model sam tworzy „etykiety” do danych uczących opartych na surowych treściach uczących poprzez na przykład przewidywanie następujących po sobie słów.

efekty, które nie są prostym powtórzeniem, ale nową kombinacją treści (Peñalvo i Ingelmo, 2023). LLM stanowią podstawę działania GenAI, których architektura opiera się na sztucznych sieciach neuronowych typu Transformer (Vaswani i in. 2017). Modele te przetwarzają tekst przy użyciu mechanizmu uwagi (ang. *self-attention*), który pozwala im analizować zależności pomiędzy słowami w kontekście całej wypowiedzi, niezależnie od ich położenia. Na przykład zaimek „on” może zostać poprawnie powiązany z imieniem „Tomasz”, nawet jeśli występuje ono kilka słów wcześniej. Jednym z modeli typu LLM jest GPT (ang. *generative pre-trained transformer*), który stanowi przykład zastosowania między innymi uczenia maszynowego wykorzystującego techniki uczenia samonadzorowanego do generowania wypowiedzi w języku naturalnym naśladujących ludzką komunikację (Radford i in., 2018). Kolejną charakterystyczną cechą LLM jest ich zdolność do uczenia się na podstawie zbiorów danych, obejmujących setki miliardów słów pochodzących między innymi z literatury, stron internetowych, artykułów i innych publicznie dostępnych źródeł.

W tym tekście skupiamy się na jednym z obszarów zastosowania technologii sztucznej inteligencji, jakim są badania prowadzone w naukach społecznych (Bail, 2024; Grossman i in., 2023; Christou, 2023; Argyle i in., 2025). Warto jednak zaznaczyć, że wykorzystanie narzędzi AI nie ogranicza się do badań społecznych, lecz dotyczy także badań opinii, badań konsumenckich czy badań *user experience*, gdzie również jest wytwarzana wiedza o rzeczywistości społecznej. Przykładowo, niektóre agencje badawcze opowiadają się otwarcie za wdrażaniem technologii AI, twierdząc, że mogą one „ulepszyć badania społeczne, które służą rozwojowi polityk publicznych, wykorzystując nowe źródła danych na większą skalę oraz nowe typy analiz” (Sampson, 2024). Warto podkreślić, że stosowane badania społeczne tworzą reprezentacje i obrazy życia społecznego, wykorzystywane między innymi w działalności biznesowej czy tworzeniu polityk publicznych.

Aby wyjaśnić skalę i tempo rozwoju zastosowań AI w badaniach społecznych, warto cofnąć się do 2022 roku, kiedy firma OpenAI wdrożyła ogólnodostępne narzędzie – ChatGPT – oparte na działaniu dużych modeli językowych. Zaledwie rok później w czasopiśmie „Nature” ukazał się artykuł przedstawiający i omawiający możliwe obszary zastosowania LLM w procesie badawczym (Stokel-Walker i Noorden, 2023). Predykcje te dotyczyły między innymi generowania hipotez, operacjonalizacji metodologii, analizy danych, a także pisania artykułów naukowych. Autorzy rozpatrywali także możliwość wdrożenia AI do procesu publikacyjnego, w którym narzędzia AI mogłyby zastąpić pracę osób recenzujących i redagujących artykuły naukowe. Do momentu, w którym piszemy ten tekst (pierwsza połowa 2025 roku), część z tych predykcji została spełniona, sama liczba publikacji naukowych wykorzystujących generatywne narzędzia AI między rokiem 2023 a 2024 wzrosła pięciokrotnie, przy czym trend ten wciąż się nasila,

a znaczenie AI bywa porównywane do przełomu, jaki niegdyś przyniósł rozwój Internetu – technologii, która przekształciła wiele aspektów badań w naukach społecznych (Argyle i in., 2025).

Przyjrzyjmy się teraz nieco dokładniej zastosowaniom technologii AI w badaniach społecznych. Poniższy przegląd nie ma charakteru systematycznego, z pewnością nie jest również wyczerpujący – naszym celem było jedynie pokazanie, jakie zastosowania (w różnych aktywnościach badawczych i na różnych etapach procesu badawczego) mogą mieć narzędzia oparte na generatywnej sztucznej inteligencji. W tej części tekstu nie odnosimy się jeszcze do proponowanych zastosowań krytycznie – nie oceniamy faktycznej technicznej sprawności technologii AI w tych obszarach ani nie wskazujemy możliwych, nie zawsze zamierzonych, konsekwencji takich zastosowań. Próbujemy jednak pokazać, w jaki sposób technologia ta może być stosowana, a przynajmniej o jakich rodzajach użycia wspomina się w dyskusji naukowej. Teksty, do których się odwołujemy, to – z kilkoma istotnymi wyjątkami – artykuły dotyczące wykorzystania GenAI w naukach społecznych, ale opublikowane w czasopismach poświęconych dyscyplinom technicznym. Jakkolwiek nasz dobór materiałów z pewnością jest obciążony, to sam fakt prowadzenia tego typu dyskusji poza obszarem socjologii akademickiej wydaje się znaczący.

Przeglądy źródeł i streszczenia

Do ważnych czynności badawczych, w których narzędzia AI mogą być stosowane, należy generowanie streszczeń literatury, robienie przeglądów systematycznych, prowadzenie eksploracji teoretycznych czy pisanie dokumentów koncepcyjnych (Torre-López i in., 2023). W artykułach na temat metodologii systematycznych przeglądów literatury (np. Petticrew i Roberts, 2008) wyodrębnia się sześć etapów (planowanie, wyszukiwanie, przeglądanie, ekstrakcja i synteza danych, ocena jakości danych oraz przedstawienie wyników), z których każdy może być częściowo lub w całości zrealizowany przy pomocy narzędzi opartych na AI (Bolaños i in., 2024). Przykładem są zarówno systemy do generowania przeglądów i streszczeń literatury (np. Elicit, SciSpace), jak i rozszerzenia wyszukiwarek baz artykułów naukowych (np. Scopus, Web of Science, JSTOR) działających na podstawie metod AI i LLM (m.in. GPT, Gemini). Narzędzia takie jak Elicit i SciSpace zostały przetestowane i opisane w artykułach naukowych, które można traktować jako formę recenzji i jednocześnie legitymizację ich użycia w pracy naukowej (Jain i in., 2024; Whitfield i Hofmann, 2023).

Identyfikowanie niezbadanych obszarów badawczych, tworzenie przeglądów literatury, generowanie streszczeń czy też wykonywanie innych zadań prowadzących do konceptualizacji problemu badawczego to kolejne funkcje ogólnodostępnych modeli LLM, takich jak GPT czy Gemini. Co więcej, o korzyściach

wynikających z użycia powyższych modeli do celów badawczych mówią nie tylko organizacje, które je wytwarzają, ale także przedstawiciele i przedstawicielki świata nauki (zob. Haman i Školník, 2023). Jak stwierdzono w jednym z artykułów na temat zastosowania GenAI do przeglądów systematycznych: „Dzięki automatyzacji wiele czasochłonnych i żmudnych zadań związanych z przeglądami literatury może być realizowanych przez ChatGPT, co ułatwi naukowcom prowadzenie bardziej wyczerpujących i rzetelnych przeglądów, a w rezultacie publikacji przełoży się na jakość recenzowanych manuskryptów, i to w znacznie bardziej wydajny sposób” (Huang i Tan, 2023).

Poza tym powstają narzędzia wspierające kompleksowo proces konceptualizacji, na przykład Research Agent (Baek i in., 2024) tworzony we współpracy z Microsoft, który pozwala na podstawie wybranej publikacji dobrać powiązane prace, wykorzystując bliskoznaczne pojęcia z korpusu składającego się z literatury naukowej. Research Agent to trzystopniowy model, który oprócz wsparcia w operacjonalizacji problemu badawczego proponuje metodę i instrukcję do stworzenia badania. Tego typu narzędzia, współtworzone przez środowisko naukowe i organizacje technologiczne, oferują wielozadaniowe wsparcie pracy badawczej – przykładem może być wspomniany ResearchAgent (Microsoft) oraz AI Scientist (University of Oxford, University Columbia), który nie tylko przeprowadza przegląd, ale pomaga też w opracowaniu konceptualizacji i stworzeniu propozycji metodologii, a nawet generuje manuskrypt tekstu (Lu i in., 2024).

Generowanie pomysłów badawczych i hipotez

Proces tworzenia hipotez bywa uznawany za „pierwszy krok” do odkrycia naukowego (Wang i in., 2023). Jednym z proponowanych narzędzi do generowania hipotez w naukach społecznych jest zbiór danych TOMATO (*auTOMated open-doMAin hypoThetical inductiOn*), którego działanie oparte jest na danych tekstowych z otwartych źródeł, głównie z renomowanych czasopism w obszarze nauk społecznych (Yang i in., 2023). Na tle innych podobnych narzędzi (np. Research Agent, Elicit, AI Assist) wyróżnia go to, że jest docelowo skierowany do użytku w obszarze nauk społecznych, a poza tym nie tylko prezentuje wyniki, ale też aktualizuje ich rezultaty na podstawie informacji zwrotnej od użytkownika/czki. Zastosowanie powyższego iteracyjnego podejścia ma na celu udoskonalenie tworzonych przez program pomysłów badawczych pod względem trafności, rzetelności i nowości. Dodatkową opcją jest możliwość formułowania hipotez kontrintuicyjnych, podważających wcześniejsze założenia.

Kolejną funkcjonalnością AI, do której współtworzenia i oceny zostali zaproszeni naukowcy i naukowczynie, było generowanie pomysłów badawczych przez modele LLM (Si i in. 2024). W ramach badania przeprowadzono eksperyment porównujący pomysły generowane przez modele językowe z tymi

tworzonymi przez osoby zajmujące się naukowo dziedziną przetwarzania języka naturalnego (ang. *natural language processing*, NLP). Następnie 79 osób eksperckich przeprowadziło podwójnie ślepą ocenę pomysłów ze względu na trzy kategorie: pomysły stworzone przez ludzi, generowane przez LLM oraz takie, które były generowane przez AI, ale poprawione przez człowieka. Oceny dotyczyły między innymi nowatorskiego podejścia i wykonalności. Jak wykazały wyniki eksperymentu, pomysły zaproponowane przez LLM zostały ocenione jako bardziej nowatorskie niż te przedstawione przez ludzi ($p < 0,05$), lecz nieco słabsze pod względem wykonalności (Si i in., 2024, s. 13, 16). Warto jednak zaznaczyć, że tematyka pomysłów nie była w tym badaniu powiązana z naukami społecznymi.

Tworzenie narzędzi badawczych, przeprowadzanie wywiadów i analiza danych

Kolejnym przykładem czynności badawczych, w których AI może mieć zastosowanie, szczególnie w naukach społecznych, jest tworzenie kwestionariuszy (Paduraru i in., 2024), skal pomiarowych (Götz i in. 2024) i scenariuszy wywiadów (Chopra i Haaland, 2023) na podstawie przeprowadzonych wcześniej konceptualizacji. Na przykład Friedrich Götz i współpracownicy (2024) omawiają zastosowanie narzędzia *Psychometric Item Generator* (PIG) opartego na działaniu modelu GPT, które umożliwia generowanie dużych ilości „ludzkopodobnych” – jak określają to autorzy – treści testowych (np. stwierdzeń w kwestionariuszu), na których podstawie tworzone są skrócone wersje znanych w psychologii miar, jak tak zwana Wielka Piątka czy też skale do mierzenia nowych konstruktów psychologicznych (choćby tęsknoty za podróżami).

Z kolei w obszarze badań jakościowych jednym z powierzonych AI zadań było prowadzenie wywiadów z respondent(k)ami (Chopra i Haaland 2023). Cel badania stanowiło określenie powodów braku zainteresowania rynkiem akcji wśród osób, które ze względu na posiadane oszczędności i dochody mogłyby inwestować na giełdzie. Autorzy stworzyli system dialogowy oparty na wieloagentowej architekturze, w którym każda z instancji AI odpowiadała za inne zadanie, na przykład jeden agent generował podsumowanie rozmowy, a drugi zadawał pogłębiające pytania. Program przeprowadził łącznie 381 wywiadów opartych na wcześniej przygotowanym skrypcie z zagadnieniami. Jak wskazują autorzy, 91% badanych oceniło doświadczenie pozytywnie, choć ponad połowa wolałaby wywiad z AI niż z człowiekiem. Wydaje się, że tego typu narzędzia pozwalają na uzupełnienie i pogłębienie badań ilościowych i otwierają nowe możliwości w obszarze metod badań mieszanych (ang. *mixed mode*).

Petter Törnberg (2023) podkreśla, że rozwój LLM wprowadza zmiany w analizie danych w naukach społecznych ze względu na wysoką wydajność

i wielowymiarowe metodę ich klasyfikacji. Jak zauważa Christopher Bail (2024), w obszarze nauk społecznych badacze/ki zaczęli/ły opracowywać zasady dotyczące kodowania materiału empirycznego za pomocą LLM. Jednym z przykładów takich zaleceń jest przewodnik do konfiguracji zadań dotyczący klasyfikacji danych przy użyciu interfejsów API (ang. *Application Programming Interface*) (Törnberg, 2023). Opracowane wytyczne podkreślają, że LLM mogą być wykorzystywane zarówno w badaniach ilościowych, jak i jakościowych oraz potrafią realizować złożone zadania, takie jak interpretacja kontekstu, rozpoznawanie ironii i emocji. Na przykładzie pomiaru populizmu w danych zastanych autor pokazuje, że LLM mogą skutecznie identyfikować złożone konstrukty społeczne i ideologiczne, a jednocześnie mieć wysoką zgodność z oceną ekspercką. Jednocześnie Törnberg (2023) zaleca ostrożność w stosowaniu AI w badaniach jakościowych, rekomendując korzystanie z jej potencjału przy uwzględnianiu ograniczeń. Podobne wnioski można odnaleźć w badaniach Davida Morgana, który wykorzystał dwa zbiory danych jakościowych do przeprowadzenia kodowania zapośredniczonego przez ChatGPT. Autor porównał wyniki klasyfikacji danych za pomocą darmowej wersji modelu GPT oraz kodowania z użyciem manualnych technik (Morgan i Nica, 2020). Wygenerowane przez Chat odpowiedzi skupiały się bardziej na wyodrębnieniu szczegółów niż na opracowaniu szerszej – kontekstowej – analizy (Morgan, 2023, s. 17). Jak wskazuje Morgan, ChatGPT nie powinien być używany w izolacji, tj. jako osobne narzędzie do analizy danych jakościowych, a wygenerowane wyniki analizy najlepiej traktować jako materiał pomocniczy w tradycyjnej interpretacji danych.

Wykonywanie badań symulacyjnych

Modelowanie agentowe polega na symulowaniu złożonych zjawisk społecznych za pomocą modeli obliczeniowych (Miller i Page, 2007, s. 65, 78; Bruch i Atwell, 2015). Badania symulacyjne pozwalają na prowadzenie zarówno rozważań nad tym, „co by było, gdyby...”, jak i badań w warunkach kontrolowanych, z obniżonym ryzykiem i bez kosztów związanych z rzeczywistymi eksperymentami społecznymi, w tym na próbach reprezentujących różnorodne, wrażliwe i trudno dostępne populacje. Do zalet modeli agentowych w badaniach społecznych można, według Christopera Bailego (2024), zaliczyć możliwość eksplorowania hipotetycznych scenariuszy rozwoju różnych zjawisk społecznych. Modele AI są w stanie na przykład identyfikować wzorce na poziomie indywidualnym (takie jak uprzedzenia grupowe), które mogą przekładać się na zjawiska na poziomie makro (np. segregację mieszkaniową). Ponadto, jak wykazały badania Shultsa (2025) nad wdrożeniem technologii opartej na AI do działania modeli wieloagentowych (ang. *Multi-Agent Artificial Intelligence Modeling*, MAAI), dzięki odpowiedniej konfiguracji model może generować

spolaryzowane opinie. W praktyce oznacza to, że model symuluje odpowiedzi „osób” o określonych cechach społeczno-demograficznych i osobowości. W trakcie badań możliwe jest ustawienie takiej konfiguracji, która pozwala na stworzenie interakcji między agentami na podstawie różnorodnego, a niekiedy też rozłącznego zestawu rozkładów odpowiedzi w modelu. Stopień, w jakim model może odzwierciedlać rozkłady odpowiedzi zbliżonych do tych, jakie obserwujemy w różnych subpopulacjach, Argyle z zespołem (2023, s. 4) określa jako „stopień wierności algorytmu”. Natomiast proponowana przez Shultsa (2025) metoda wieloagentowa zakłada badania nad „sztucznym społeczeństwem”, w którym przedmiotem analiz są agenci o określonych osobowościach wchodzący ze sobą w interakcje w zasymulowanym środowisku.

Odgrywanie roli współpracowników/czek badawczych

W przedstawionych powyżej artykułach pojawiały się narzędzia GenAI wykonujące głównie pojedyncze zadania badawcze. Na tym tle wyróżnia się framework Agent Laboratory (Schmidgall i in. 2025), za pomocą którego można realizować kilka czynności badawczych w jednym narzędziu. Agent Laboratory jest przykładem zastosowania wieloagentowego systemu (ang. *multi-agent systems*, MAS) do realizacji kompleksowych zadań badawczych (Calegari i in., 2021). Działanie MAS polega na tym, że kilku autonomicznych agentów (programów) funkcjonuje w jednym środowisku, ucząc się na podstawie sprecyzowanych przez użytkownika/czkę zapytania. Agenci mogą ze sobą nie tylko współpracować, ale także wchodzić w konflikty. Agent Laboratory rozpoczyna działanie w momencie przekazania pomysłu badawczego, a następnie przechodzi przez kilka etapów działań badawczych za pomocą odpowiedzialnych za różne zadania agentów przeprowadzających między innymi przegląd literatury, proces tworzenia pytań badawczych, analizy wyników, kończąc go etapem pisanie raportu z badań; można też konfigurować otrzymane wypowiedzi i częściowo monitorować przebieg rozumowania agentów.

Wiedza w krzywym zwierciadle AI

Skala i tempo rozwoju GenAI budzą jednak nie tylko entuzjazm, lecz także obawy. Komisja Europejska przyznała, że badania naukowe są jednym z obszarów, które mogą zostać przekształcone przez generatywną sztuczną inteligencję (European Commission, 2024, s. 3), a renomowane czasopisma naukowe, takie jak „Nature” i „Science”, przedstawiły stanowiska wyrażające zaniepokojenie zastosowaniem AI w procesach badawczych (Stokel-Walker i Noorden, 2023). Równocześnie wiele uznanych uczelni wyższych, między innymi Stanford

University, Cornell University, University of Oxford, opracowało etyczne zasady związane z zastosowaniem AI w badaniach i edukacji (Stanford University, n.d.; University of Oxford, n.d.; Cornell University, 2023), tym samym legitymizując ich wykorzystanie na określonych zasadach. Mimo to ustalone instytucjonalnie reguły mogą nie uwzględniać w wystarczającym stopniu specyfiki zastosowań AI w badaniach społecznych. Na przykład w zaleceniach Cornell Univeristy (2023, s. 10) dotyczących zasad przetwarzania danych uczestników/czek podczas korzystania z narzędzi GenAI wyróżnione są sposoby zarządzania danymi, lecz wytyczne te zawężone są tylko do badań klinicznych. W tej części artykułu skupimy się na zarysowaniu kilku istotnych problemów, jakie wiążą się z użyciem AI w badaniach społecznych.

Choć narzędzia oparte na działaniu GenAI oferują nowe możliwości w obszarze badań społecznych, to użycie ich w procesie badawczym AI wiąże się z ryzykami, które mogą kształtować zarówno proces generowanych odpowiedzi, jak i ich treść. Pierwszym z nich jest problem tak zwanych skrzywień algorytmicznych (ang. *algorithmic bias*), zwłaszcza dotyczących płci, rasy, niepełnosprawności i różnorodności etnicznej, które mogą prowadzić do wzmacniania i powielania istniejących wzorców dominacji (Whittaker i in. 2018; Obermeyer i in. 2019; Christou 2023). Przykładowo, w badaniu Ortega-Martin i in. (2023) poddano analizie udzielone przez ChatGPT odpowiedzi na zapytania dotyczące testowania przypisywanych przez model językowy zaimków, na przykład tworzone krótkie zdania, w których pojawiały się dwie osoby o różnych profesjach oraz zaimek (np. „on”, „ona”), a następnie pytano czatbota, do kogo odnosi się dany zaimek. Model miał tendencję do przypisywania zaimka w sposób stereotypowy (np. fizyk to mężczyzna, a sekretarka to kobieta). Wzorec ten utrzymywał się, nawet gdy gramatyka zdania wskazywała na inną interpretację.

Z kolei w badaniach Hofmann i in. (2024) wykazano, że modele językowe przejawiają ukryte uprzedzenia rasowe, reagując negatywnie na użytkowników/czki posługujących się afroamerykańskim angielskim. Osobom tym częściej przypisywano mniej prestiżowe zawody, sugerowano wyższe prawdopodobieństwo popełnienia przestępstwa oraz rekomendowano surowsze wyroki mimo braku jawnych przesłanek związanych z rasą. Co ważne, zastosowane mechanizmy ograniczające uprzedzenia w modelu jedynie maskowały problem, podczas gdy jego głębsze struktury dalej zawierały uprzedzenia na tle rasowym. Wnioski te pokazują, że mimo zaawansowania technologicznego modele generatywne nie są neutralne, lecz odzwierciedlają i powielają istniejące w społeczeństwie nierówności. Tym samym ich zastosowanie w badaniach społecznych wymaga nie tylko krytycznej refleksji metodologicznej, ale także wdrażania mechanizmów stałej ewaluacji i kontroli jakości danych wyjściowych.

Innym, wciąż aktualnym wyzwaniem w działaniu LLM pozostaje kwestia tak zwanej halucynacji. Jedną z przyczyn tego problemu jest specyfika

generatywnej AI, będącej w stanie tworzyć tekst do złudzenia przypominający ludzką wypowiedź, ale nie potrafiącej ocenić jego wiarygodności (Bommasani i in., 2021). Narzędzia GenAI nie rejestrują faktów, tylko uczą się wzorców, czyli tego jakie słowa najczęściej występują obok siebie w zbiorze danych. Model językowy w sposób probabilistyczny „odgaduje” odpowiedź, ale jej nie „weryfikuje” i większy priorytet ma dla niego odpowiedź na pytanie „co powinno być dalej” niż to, na ile wygenerowana odpowiedź jest prawdziwa. Wobec tego narzędzia do generowania pomysłów, hipotez badawczych czy przeglądu literatury powinny podlegać ocenie badaczy/ek w zakresie poprawności dat oraz nazwisk, ponieważ są to dane, których prawidłowe dopasowanie bywa utrudnione ze względu na probabilistyczny mechanizm działania GenAI (Yuan i in., 2023; You i in., 2024). Co więcej, sposób działania metod AI pozostaje w dużej mierze nieprzejrzysty. Liczba parametrów w LLM jest bliska setek miliardów, co uniemożliwia prześledzenie, na jakich podstawach modele językowe podejmują decyzje, generując odpowiedzi (problem tak zwanej czarnej skrzynki – zob. Rudin, 2019). Choć pojawiły się inicjatywy takie jak „wyjaśnialna sztuczna inteligencja” (ang. *explainable AI*, XAI) czy „odpowiedzialna sztuczna inteligencja” (ang. *responsible AI*), to bywają one krytykowane jako „techniczne odpowiedzi na problemy o charakterze technospołecznym” (Bassett i Roberts, 2023). Warto też mieć na uwadze, że pojęcia wyjaśnialności AI, asymetrii w danych uczących i transparentności działania modelu pochodzą z informatyki, „w której terminy takie jak niedopasowanie (system zawiera zbyt mało danych, aby odzwierciedlić świat społeczny) i nadmierne dopasowanie (system składa się ze zbyt wielu losowych zmiennych, aby odzwierciedlić świat społeczny) są używane do opisania stronniczości w przetwarzaniu danych” (Joyce i in. 2021, s. 6). Znajomość powyższych pojęć skłania do refleksji nad tym, w jakim zakresie GenAI jest w stanie adekwatnie oddać złożoność zjawisk społecznych. Być może jedną z propozycji wspierających proces krytycznego myślenia nad wygenerowanymi treściami jest opracowanie przez badaczy/ki w naukach społecznych własnej ogólnodostępnej infrastruktury w formie modelu językowego wspieranego przez GenAI (Bail, 2024).

Ponadto ze względu na wspomniany probabilistyczny charakter wiedzy generowanej przez narzędzia AI istnieje ryzyko, że ich użytkowanie przyczyni się do uproszczenia złożoności i różnorodności oraz „spłaszczenia” obrazów świata społecznego (Browne i in., 2023). Ma to szczególne znaczenie dla społecznych praktyk badawczych i wiedzy o rzeczywistości społecznej, którą te praktyki wytwarzają i jako takie zasługują na szczególną i krytyczną uwagę w socjologii. Ponadto zauważalny w literaturze optymizm technologiczny, koncentrujący się na wzroście produktywności i wygodzie stosowania narzędzi GenAI wzmacnia ich instrumentalne postrzeganie, jednocześnie ograniczając refleksje nad epistemicznymi skutkami ich zastosowań. Jak wykazali Lee

i współpracownicy (2024), wysokie zaufanie do GenAI wiąże się ze zmniejszeniem zaangażowania w krytyczną refleksję nad wygenerowanymi treściami między innymi przez sprawdzanie poprawności otrzymanych odpowiedzi. Na podobny problem wskazują Kelly Joyce z zespołem, podkreślając, że „socjologiczne zrozumienie danych jest ważne, biorąc pod uwagę, że bezkrytyczne wykorzystanie ludzkich danych w systemach socjotechnicznych AI będzie miało tendencję do powielania, a może nawet pogłębiania istniejących wcześniej nierówności społecznych” (2021, s. 3). Wobec powyższego socjologiczne kompetencje zdają się niezbędne nie tylko do interpretacji treści generowanych przez AI, a brak refleksyjnego podejścia przy użytkowaniu narzędzi GenAI w pracy badawczej może prowadzić do nasilenia niektórych z wymienionych powyżej ograniczeń AI, w tym powielania uprzedzeń i halucynacji.

Gwoli sprawiedliwości trzeba zaznaczyć, że autorzy/ki większości omawianych tekstów poświęconych zastosowaniu AI w badaniach społecznych zwracają uwagę na wspomniane tutaj ograniczenia tych narzędzi. Bodaj najczęściej pojawia się kwestia ryzyka powielania uprzedzeń oraz generowania błędnych odpowiedzi, wynikających ze skrzywień algorytmicznych (Argyle i in., 2023, 2025) i niereprezentatywności danych wykorzystywanych do trenowania modeli językowych (zob. You i in., 2024). Proponowane są różne strategie zaradcze mające na celu zwiększenie rzetelności generowanych odpowiedzi, obniżanie ryzyka halucynacji modelu czy wykrywanie uprzedzeń. Do takich działań należą między innymi kalibracje modeli LLM, czyli proces analizy danych i kodu modelu obliczeniowego mający zapewnić, że architektura LLM „reprezentuje rzeczywisty świat w odpowiedni sposób” (Shults, 2025, s. 4), określenie stopnia „algorytmicznej wierności” jako sposobu na uzasadnienie wykorzystania modeli LLM do symulacji myślenia (Argyle i in., 2023) czy ustalenie, pod jakimi warunkami wygenerowana treść będzie uznana za prawdziwą (Argyle i in., 2025).

Rekonfiguracje socjo-technologiczne

Do tej pory nasze omówienie narzędzi opartych na technologii GenAI miało charakter instrumentalistyczny, tj. zakładało traktowanie ich jako narzędzi, które mogą pomagać w osiąganiu zamierzonych celów (w tym wypadku: wytwarzać wiedzę o świecie społeczną w łatwiejszy czy szybszy sposób), jednak same w sobie pozostają neutralne, a więc nie wspierają i nie ucieleśniają żadnych konkretnych wartości czy ideologii. Dokonana później krytyczna ocena częściowo wychodziła poza paradygmat instrumentalistyczny – w tym sensie, że traktowała rozmaite zniekształcenia i skrzywienia nie jako niedoskonałości AI, które hipotetycznie mogłyby być usunięte, ale po części jako immanentne cechy tej

technologii czy sposobów, w jakie jest ona tworzona. W ostatniej części tego artykułu chcemy zrobić kolejny krok – zasugerować aparat teoretyczny, który prowadzi do odmiennego niż poprzednie i naszym zdaniem bardziej adekwatnego socjologicznie rozumienia użycia AI w badaniach społecznych. Naszą propozycję można określić jako ujęcie socjo-technologiczne, czyli takie, które traktuje AI jako technologię z jednej strony kształtowaną przez aktorów, praktyki społeczne, nadawane jej znaczenia czy relacje władzy, a jednocześnie taką, która te społeczne zjawiska współtworzy i warunkuje, a więc wywiera aktywny wpływ na świat społeczny. Jakkolwiek proponowane przez nas ujęcie jest tylko jednym z wielu możliwych, ta zdolność ujęcia relacji między tym, co społeczne, a tym, co technologiczne, bez polegania na naiwnie instrumentalistycznych lub zbyt deterministycznych założeniach wydaje nam się główną wartością tego podejścia.

Pierwszą naszą teoretyczną inspiracją jest w tym kontekście współczesna generacja teorii praktyk. Zyskała ona znaczną uwagę w socjologii – była nawet określana mianem „zwrotu ku praktykom” w teorii społecznej (Schatzki i in., 2001) i stała się analityczną podstawą badań w wielu obszarach, w tym w studiach nad organizacjami, konsumpcją, energią, miastem, rodziną i wielu innych (Røpke, 2009; Shove, 2010; Feldman i Orlikowski, 2011; Warde, 2014; Hui i in., 2017; Smagacz-Poziemska i in., 2018; Sikorska, 2018); zaznaczyła się także w socjologii cyfrowej (Hanchard, 2024). W tym artykule korzystamy z proponowanej przez teorię praktyk perspektywy patrzenia na innowacje (Shove i Pantzar, 2005; Pantzar i Shove, 2010) i przyjmujemy, że narzędzia AI przeznaczone do prowadzenia badań społecznych mogą być potraktowane jako innowacyjne produkty, które zaczynają funkcjonować w ramach praktyk badawczych.

Niestety nie ma tutaj miejsca na szczegółowe przedstawienie założeń teoretycznych teorii praktyk, ograniczymy się zatem tylko do wymienienia najważniejszych założeń tej rodziny teorii, odsyłając zainteresowane osoby do innych publikacji (Nicolini, 2012; Shove i in., 2012; Smagacz-Poziemska i in., 2018; Rudnicki, 2023), w tym do numeru specjalnego „Studiów Socjologicznych” (4/2023) poświęconego teoriom praktyk (Krajewski, 2023 i inne artykuły we wspomnianym numerze). Teorie praktyk rozwijają wątki, które pojawiły się w teorii społecznej już wcześniej, w szczególności w pracach Anthony’ego Giddensa i Pierre’a Bourdieu. Wspólnym punktem obecnej generacji tych koncepcji jest konsekwentnie stosowane założenie, że ontologicznie pierwotnym elementem świata społecznego i centralną kategorią analizy socjologicznej powinny być praktyki rozumiane jako wiązki (ang. *bundles*) aktywności (Schatzki i in., 2001, s. 11; Schatzki, 2002, s. 70; zob. także Hui i in., 2017), czyli „ucieleśnione, materialnie zapośredniczone układy [*arrays*] ludzkiej aktywności skupione wokół podzielanych praktycznych rozumień” (Schatzki i in., 2001, s. 12). Praktyką mogą być zakupy, parkowanie samochodu, spacer z dzieckiem czy

prowadzenie wywiadów pogłębionych. Kolejnym fundamentalnym dla teorii praktyk założeniem jest przekonanie, że praktyki to układy, które składają się z wielu różnorodnych, ale powiązanych elementów. W jednym z najbardziej popularnych modeli teorii praktyk, czyli koncepcji zaproponowanej przez Elizabeth Shove, Mikę Pantzara i Matta Watsona (2012), elementy, na które składają się praktyki, zostały podzielone na trzy grupy: 1) znaczenia (ang. *meanings*), obejmujące społeczne i symboliczne aspekty praktyk, w tym emocje, motywacje, idee, wartości, aspiracje i cele; 2) materiały (ang. *materials*), do których zaliczają się artefakty, technologie, narzędzia, sprzęty i materiały, z których zostały one zrobione; oraz 3) umiejętności (ang. *competences*), czyli wiedza praktyczna, *know-how*, techniki i zdolności niezbędne do społecznie adekwatnego zaangażowania się w daną praktykę. Wszystkie te elementy są połączone ze sobą w dynamiczny układ, czyli właśnie praktykę. Jednocześnie praktyki nie są jedynie statycznymi konfiguracjami, ale są dynamiczne i otwarte, „zawsze w procesie formacji, reformacji i deformacji” (Shove i in., 2012, s. 44), w miarę jak do istniejących układów dołączają się nowe elementy, a inne połączenia zanikają.

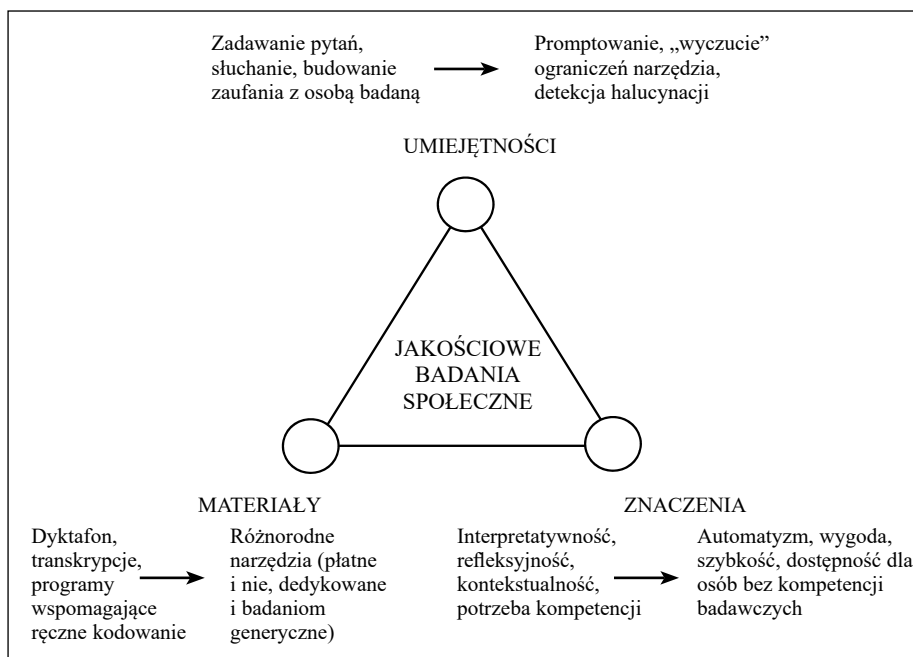
To „kompozycyjne” rozumowanie, które proponują teorie praktyk, ma istotne konsekwencje dla rozumienia innowacji. W tym ujęciu są one ujmowane nie w płaszczyźnie technologicznej (jako urządzenia) czy ekonomicznej (jako produkty), ale jako przemiany zachodzące w ramach praktyk (Shove i Pantzar, 2005; Pantzar i Shove, 2010). Nowe urządzenia, przedmioty czy aplikacje muszą bowiem dopiero zostać inkorporowane do istniejących praktyk albo stać się punktem wyjścia nowych – dzieje się to jednak nie przez samo ich wyprodukowanie, ale w drodze „aktywnej i bieżącej integracji z wyobrażeniami, przedmiotami i formami kompetencji, w procesie, w który zaangażowani są zarówno producenci jak i konsumenci” (Shove i Pantzar, 2005, s. 43). Innymi słowy, włączenie nowego urządzenia do istniejącej praktyki, to nie proste dodanie jakiegoś elementu należącego do *materials*, ale rekonfiguracja praktyki obejmująca szereg, pozornie subtelnych, ale często znaczących przesunięć. Przykładowo, wprowadzenie rowerów z dodatkowym napędem elektrycznym (tzw. hybrydowych) to nie prosta zmiana jednego rodzaju roweru na inny, ale szereg przekształceń praktyki jeżdżenia na rowerze. Takie rowery wymagają bowiem szeregu dodatkowych urządzeń (np. baterii czy mocniejszych hamulców), przyswojenia nowych umiejętności (np. wyczucia, że rower „dodaje” tym więcej mocy, im silniej naciskamy na pedały, czy opanowania manewrowania rowerem, który jest znacznie cięższy niż rowery tradycyjne), a także nowych znaczeń (np. redefinicji jazdy na rowerze jako aktywności sportowej i wymagającej istotnego wysiłku). Innymi słowy, innowacja pociąga za sobą zmianę całej kompozycji praktyki.

Przekładając podejście teorii praktyk do innowacji na kwestię zastosowania narzędzi AI w badaniach społecznych, możemy przyjąć, że punktem ciężkości

analizy nie będzie tutaj technologia jako taka, ale praktyki, w ramach których jest ona używana. Kluczowe będą zatem nie same funkcjonalności narzędzi badawczych wykorzystujących AI (np. te omawiane wcześniej), ale praktyki badawcze, które składają się z określonych znaczeń (np. standardów metodologicznych czy istniejących wyobrażeń na temat adekwatnych sposobów badania rzeczywistości społecznej), materiałów i urządzeń (np. scenariuszy do badań, transkrypcji, raportów, prezentacji) czy umiejętności (np. formułowania problemów badawczych, dobierania próby do badań, wnioskowania). W tym ujęciu włączanie nowych narzędzi do praktyk badawczych nie jest bynajmniej czymś automatycznym czy mającym z góry określony kształt (jak zakładała większość referowanych przez nas artykułów omawiających zastosowania AI w badaniach), ale otwartym procesem adopcji pewnej technologii i jej integracji ze znaczeniami, materiałami i umiejętnościami, które w danym momencie wchodziły w skład praktyki badań społecznych.

Przyjrzyjmy się na przykład zastąpieniu tradycyjnej analizy danych jakościowych polegającej na pracochłonnym kodowaniu transkrypcji przeprowadzonych uprzednio wywiadów pogłębionych częściowo automatycznym kodowaniem materiału, generowaniem streszczeń i syntez albo wręcz z czytaniem z materiałem z wywiadów (zob. Luybashenko, 2025) – funkcjonalności te oferowane są choćby przez takie narzędzia, jak Maxqda Tailwind, Dovetail, Advinote czy Notebook LM. Mamy tu do czynienia nie tylko z istotną zmianą rodzaju czynności, jakie składają się na analizę danych, a więc jednego z podstawowych składników procesu badawczego, ale także z potencjalną transformacją znaczeń regulujących tę praktykę (rysunek 1). Wspomniane narzędzia tego typu są zazwyczaj opisywane jako ułatwiające i przyspieszające prowadzenie badań – konotują zatem znaczenia pochodzące spoza domeny naukowej. Jednocześnie normalizują automatyczność procesu analizy, która zlecana jest maszynie, a więc przestaje być możliwa do prześledzenia, oraz modyfikują rolę badacza/ki, ponieważ „to model «podejmuje decyzję» odnośnie do tego, jaki element analizowanego tekstu odpowiada na nasze pytanie” (Luybashenko, 2025, s. 141). Takie przekształcenie praktyki wiąże się również z powstaniem nowych umiejętności związanych nie tylko z prowadzeniem analizy w inny sposób, ale także na przykład wycuciem, kiedy wynikiem wygenerowanym przez narzędzie można w sensie poznawczym zaufać, a kiedy jest to zbyt ryzykowne. Wreszcie zmiany praktyk badawczych będą także dotyczyć używanych materiałów i technologii, na przykład mogą wiązać się z porzuceniem popularnych do niedawna programów do kodowania i zastąpieniem ich kilkoma alternatywnymi rozwiązaniami, niekoniecznie przeznaczonymi do badań naukowych, albo rezygnacją z programów do odtwarzania dźwięku, które były używane przy tworzeniu transkrypcji.

Rysunek 1. Zmiany w praktyce jakościowych badań społecznych (model praktyk za: Shove i in., 2012)



Omawiane na początku tego artykułu przykłady badań testujących zastąpienie ludzkich respondentów/ek próbami syntetycznymi sugerują jeszcze dalej idące transformacje praktyk badawczych, w których dominującą aktywnością może nie być już pozyskiwanie danych od ludzi, ale tworzenie sztucznych prób i trenowanie algorytmów. Trudno przy tym zignorować zarówno przekraczające 80% wskaźniki zgodności danych syntetycznych z uzyskanymi w toku typowych badań sondażowych, jak i skalę ambicji przebijającej z obu tych tekstów – na przykład wprost zadane pytanie „Czy personsy mogą symulować społeczeństwo?” (Li i in., 2025, s. 6) albo oczekiwanie, że symulacje mogą mieć „szerokie zastosowanie w tworzeniu polityk publicznych i w naukach społecznych” (Park i in. 2024, s. 1). Biorąc pod uwagę fakt, że artykuły zostały napisane przez osoby reprezentujące inne nauki niż socjologia, a ukazały się w serwisie, w którym publikowane są preprinty artykułów z nauk przyrodniczych i technicznych, można tę sytuację interpretować jako przykład kolejnych „przyczółków” (Afeltowicz i Pietrowicz, 2013), jakie w niegdyś zarezerwowanej dla socjologii domenie obrazowania świata społecznego zyskują dyscypliny techniczne i pozaakademickie (zob. Savage i Burrows, 2007).

Z perspektywy teorio-praktycznej te zjawiska sygnalizują jednak przede wszystkim potencjalne transformacje znaczeń definiujących i regulujących

praktyki badania świata społecznego, w tym wyłonienie się nowych sposobów rozumienia tego, czym jest wiarygodność, obiektywność czy rzetelność badania oraz jacy aktorzy społeczni są gwarantem epistemicznej jakości powstającej wiedzy. Tradycyjnie tworzenie wiedzy o świecie społecznym wiąże się z przetwarzaniem danych uzyskanych od ludzi, a proces ten jest prowadzony przez osoby mające rozpoznawalny społecznie status ekspercki w tej dziedzinie. Zaufanie do powstałej wiedzy opiera się na potocznie rozumianym reprezentacjonizmie, czyli przekonaniu, że wiedza powinna być możliwie wiarygodnym, dokładnym i wiernym odzwierciedleniem świata, sama jednocześnie pozostając od niego oddzielona i nie wywierając wpływu na rzeczywistość, a rola badaczy/ek polega w dużej mierze na gwarantowaniu poprawności tego procesu (Stehr, 1992; Bińczyk, 2007; Rudnicki, 2023). Próby syntetyczne, brak możliwości dokładnego prześledzenia przejścia między danymi uzyskanymi od człowieka a wynikiem analizy albo problem halucynacji stanowią wyzwanie dla tak rozumianego reprezentacjonizmu. Z perspektywy teorii praktyk są to jednak problemy nie tyle epistemiczne, ile wymagające reorganizacji praktyk badawczych. Halucynacje AI mogą być oswojone przez opracowanie sposobów oszacowania ich ryzyka albo wprowadzenie metodologicznej reguły sprawdzania wyników na kilka sposobów, zapewniających coś w rodzaju triangulacji – o takich rozwiązaniach pisze na przykład Argyle ze współpracownikami, zakładając, że „warunkiem minimum jest, by model językowy zapewniał powtarzany, spójny (*consistent*) dowód z punktu widzenia kilku kryteriów” (2023, s. 3). I znów – proces ten może wymagać transformacji znaczeń związanych z praktykami badania świata społecznego (np. w jakim sensie uzyskane wyniki są „prawdziwe”), używanych w tym celu narzędzi i technik (np. prób syntetycznych i symulacji komputerowych) oraz niezbędnych do wykonywania tej praktyki umiejętności (np. konstruowania prób syntetycznych i oceny wiarygodności wyników).

Patrząc historycznie na metody badania świata społecznego, można zakładać, że wykorzystanie na szerszą skalę danych syntetycznych może wymagać transformacji analogicznych do tych, które wiązały się z powstaniem i popularyzacją technik sondażowych. Przekonanie, że tysiącosobowa próba może wiarygodnie reprezentować wielomilionową populację, jest poznawczo trudne, jednak w toku złożonych procesów wprowadzania tej technologii badania świata społecznego (Osborne i Rose, 1999) uzyskiwała ona społeczną akceptację – co więcej, opinię publiczną uznaje się dość powszechnie za realnie istniejący zjawisko, mimo że jest to *de facto* dość arbitralnie skonstruowane pojęcie. Nie można wykluczyć, że podobny los czeka badania oparte na próbach syntetycznych pozwalających na symulację ludzkich odpowiedzi. Hipotetycznie omawiane na początku tego artykułu badanie Parka i współpracowników (2024), będące swego rodzaju „dowodem” wiarygodności danych syntetycznych, mogłoby być w tym kontekście uznane za badanie „założycielskie”, mające podobną funkcję

do słynnego badania Gallupa trafnie przewidującego wyniki wyborów prezydenckich w Stanach Zjednoczonych w 1936 roku.

Warto jednak podkreślić, że tego rodzaju transformacje nie są, w perspektywie teorii praktyk, z góry zdeterminowane, ale – przeciwnie – są otwarte, rozwijają się w czasie i są efektem aktywnych wysiłków zaangażowanych aktorów. Warto tu wskazać na sprawczy charakter także organizacji tworzących oprogramowanie, które „skryptuje” określone sposoby użycia oraz prowadzi do oswajania i normalizacji technologii, a przez to potencjalnie do domknięcia kontrowersji i uzyskania społecznego konsensusu (np. wokół danych syntetycznych). Gdy piszemy te słowa, sondaże polityczne prowadzone są na podstawie rozmów z potencjalnymi „ludzkimi” wyborcami i zastąpienie ich wynikami zebranymi w procesie badania prób „krzemowych” nie wydaje się prostą czy automatyczną konsekwencją samej technicznej możliwości uzyskania stosunkowo wiarygodnych rezultatów w ten sposób. Sprawa wygląda jednak zupełnie inaczej (łatwiej), jeśli chodzi o użycie narzędzi do AI do robienia tłumaczeń, transkrypcji wywiadów czy streszczania tekstów.

Zasadniczo otwarty i dynamiczny charakter praktyk badawczych nie oznacza jednak, że mogą one zmieniać się w nieskończenie plastyczny czy dowolny sposób. W tym kontekście warto wspomnieć drugą perspektywę teoretyczną, którą uznajemy za istotną inspirację w rozumieniu przemian praktyk badawczych związanych z popularyzacją narzędzi AI, a mianowicie podejście wielopoziomowe (ang. *multi-level*) do analizy tranzycji socjo-technologicznych zaproponowane przez Franka Geelsa (2002, 2011). Podobnie jak w podejściu teorii praktyk w tym modelu innowacje nie ograniczają się do warstwy technologicznej, ale obejmują też zmiany „praktyk, regulacji, sieci przemysłowych, infrastruktury i znaczeń symbolicznych” (2002, s. 1257), które razem tworzą konfigurację socjo-technologiczną (ang. *sociotechnical configuration*). Geels proponuje złożony i procesualny opis powstawania i rozwoju innowacji, jednak w kontekście niniejszego artykułu najważniejsze wydaje się spostrzeżenie, że w tym ujęciu praktyki nie tworzą jedynej „warstwy” rzeczywistości społecznej, ale są powiązane ze zjawiskami zachodzącymi na poziomie mezo i makro, na przykład w instytucjach, na rynkach, w ramach polityk publicznych, infrastruktury czy w środowisku kulturowym. W modelu Geelsa wszystkie te elementy są wzajemnie powiązane, a więc praktyki podlegają oddziaływaniu środowiska, którego są częścią i które je warunkuje.

Wskazywano już, że mimo pewnych różnic ontologicznych teorie praktyk dają się w niektórych warunkach pogodzić z ujęciem wielopoziomowym, dając ciekawe możliwości analityczne (Keller i in., 2024). W odniesieniu do zastosowania narzędzi AI w badaniach społecznych warto przede wszystkim podkreślić, że podejście wielopoziomowe daje możliwość uwzględnienia szeregu uwarunkowań instytucjonalnych określających warunki, w których „dzieją się”

praktyki. Można oczekiwać, że przyjęcie narzędzi sztucznej inteligencji w praktykach badawczych będzie zależało od zmian społeczno-technologicznych, takich jak „udomowienie” AI i zmieniające się wzorce korzystania z codziennych narzędzi opartych na sztucznej inteligencji relacje między firmami technologicznymi jako podmiotami na rynku, nowymi politykami technologicznymi związanych ze sztuczną inteligencją (np. AI Act), politykami nauki (np. schematami ewaluacji działalności naukowej) oraz regulacjami wdrażanymi przez redakcje, wydawców czy komitety etyczne (np. upowszechniający się wymóg stosowania tak zwanych *AI statements*, czyli informacji, w jaki sposób technologie AI zostały użyte przy tworzeniu danego tekstu naukowego). Wydaje się, że połączenie perspektywy teorii praktyk oraz ujęcia wielopoziomowego daje szczególnie ciekawe możliwości analityczne do opisu i wyjaśniania zachodzących zmian (Keller i in., 2024), zakłada bowiem, że praktyki nie zmieniają się swobodnie, ale są ramowane przez szerszy kontekst, w którym się dzieją, same jednocześnie je współtworząc.

Zakończenie

Spróbujmy na koniec zarysować kilka możliwych reakcji socjologii na obserwowany rozwój narzędzi AI i ich wykorzystanie w badaniach społecznych. Jedną z nich może być oczywiście używanie tych narzędzi, eksperymentowanie z nimi i włączenie je we własne działania badawcze. Inną – ignorowanie zachodzących zmian jako nieistotnych, czy to poprzez potraktowanie dyskusji wokół AI za niepotrzebnie rozdmuchany, ale przejściowy *hype*, czy to przez potraktowanie narzędzi AI jako drugorzędnej i czysto technicznej zmiany, która nie wpłynie w istotny sposób na badanie świata społecznego i ważność paradygmatów. Jeszcze inną reakcją może być podążanie ścieżką krytycznej analizy przez pokazywanie, w jaki sposób technologie AI reprodukują czy nawet wzmacniają istniejące nierówności, relacje władzy i formy wyzysku (już samą dyskusję wokół tej technologii można uznać za sprzyjającą interesom big techu, który poszukuje nowych wehikułów wzrostu). Kolejną możliwością jest zaangażowanie w inicjatywy promujące zniuansowany i pogłębiony namysł nad rozwojem tych narzędzi oraz współtworzenie związanych z nimi polityk i regulacji (w tym dotyczących badań społecznych). Naszym zdaniem wszystkie te reakcje są w jakiejś mierze uzasadnione. Jednak przedstawiona w poprzedniej części tego tekstu rama teoretyczna wydaje się zachęcać przede wszystkim do sięgnięcia po tradycyjną socjologiczną aktywność opisywania praktyk badawczych i stawiania pytań o to, jak wytwarza się wiedzę o świecie społecznym, kto to robi, przy pomocy jakich narzędzi, jak rozwiązywane są powstające dylematy, jak zmieniają się towarzyszące tym aktywnościom znaczenia itd. Tego

rodzaju bliskie i uważne spojrzenie, do którego teorii praktyk zachęcają, wydaje się skromnym programem badawczym, jednak obecnie może właśnie takiego podejścia najbardziej brakuje.

Bibliografia

- Afeltowicz, Ł., Pietrowicz, K. (2013). *Maszyny społeczne. Wszystko ujdzie, o ile działa*. Wydawnictwo Naukowe PWN.
- Argyle, L. P., Busby, E. C., Fulda, N., Gubler, J. R., Rytting, C., Wingate, D. (2023). Out of One, Many: Using Language Models to Simulate Human Samples. *Political Analysis*, 31(3), 337–351.
- Argyle, L. P., Busby, E. C., Gubler, J. R., Hepner, B., Lyman, A., Wingate, D. (2025). Arti-“fickle” Intelligence: Using LLMs as a Tool for Inference in the Political and Social Sciences. *arXiv*. <https://doi.org/10.48550/arXiv.2504.03822>
- Baek, J., Jauhar, S. K., Cucerzan, S., Hwang, S. J. (2024). Researchagent: Iterative Research Idea Generation over Scientific Literature with Large Language Models. *arXiv*. <https://doi.org/10.48550/arXiv.2404.07738>
- Bail, C. A. (2024). Can Generative AI Improve Social Science? *Proceedings of the National Academy of Sciences*, 121(21). <https://doi.org/10.1073/pnas.2314021121>
- Bassett, C., Roberts, B. (2023). Automation Anxiety: A Critical History – the Apparently odd Recurrence of Debates about Computation, AI and Labour. W: S. Lindgren (red.), *Handbook of Critical Studies of Artificial Intelligence* (s. 79–93). Edward Elgar Publishing.
- Bińczyk, E. (2007). *Obraz, który nas zniewala. Współczesne ujęcia języka wobec esencjalizmu i problemu referencji*. Universitas.
- Bisbee, J., Clinton, J. D., Dorff, C., Kenkel, B., Larson, J. M. (2024). Synthetic Replacements for Human Survey Data? The Perils of Large Language Models. *Political Analysis*, 32(4), 401–416.
- Bolanos, F., Salatino, A., Osborne, F., Motta, E. (2024). Artificial Intelligence for Literature Reviews: Opportunities and Challenges. *Artificial Intelligence Review*, 57(10), 259. <https://doi.org/10.48550/arXiv.2402.08565>
- Bommasani, R., Hudson, D. A., Adeli, E., Altman, R., Arora, S. i in.. (2021). On the Opportunities and Risks of Foundation Models. *arXiv*. <https://doi.org/10.48550/arXiv.2108.07258>
- Browne, J., Cave, S., Drage, E., McInerney, K. (red.). (2023). *Feminist AI: Critical Perspectives on Algorithms, Data, and Intelligent Machines*. Oxford University Press.
- Bruch, E., Atwell, J. (2015). Agent-based Models in Empirical Social Research. *Sociological Methods & Research*, 44(2), 186–221. <https://doi.org/10.1177/0049124113506405>
- Calegari, R., Ciatto, G., Mascardi, V., Omicini, A. (2021). Logic-based Technologies for Multi-agent Systems: A Systematic Literature Review. *Autonomous Agents and Multi-Agent Systems*, 35(1), 1.

- Camic, C., Gross, N., Lamont, M. (red.). (2011). *Social Knowledge in the Making*. University of Chicago Press.
- Chopra, F., Haaland, I. (2023). Conducting Qualitative Interviews with AI. <http://dx.doi.org/10.2139/ssrn.4572954>
- Christou, P. A. (2023). How to Use Artificial Intelligence (AI) as a Resource, Methodological and Analysis Tool in Qualitative Research? *Qualitative Report*, 28(7), 1968–1980. <https://doi.org/10.46743/2160-3715/2023.6406>
- Collins, R. (1994). Why the Social Sciences Won't Become High-Consensus Rapid Discovery Science. *Sociological Forum*, 9(2), 155–177.
- Collins, K. M., Jiang, A. Q., Frieder, S., Wong, L., Zilka, M. i in.(2024). Evaluating Language Models for Mathematics through Interactions. *Proceedings of the National Academy of Sciences*, 121(24). <https://doi.org/10.1073/pnas.231812412>
- Cornell University (2023). Generative AI in Academic Research: Perspectives & Cultural Norms. https://it.cornell.edu/sites/default/files/itc-drupal10-files/Generative%20AI%20in%20Research_%20Cornell%20Task%20Force%20Report-Dec2023.pdf#page=8.09
- European Commission. (2024). Living Uuidelines on Responsible Use of Generative AI. https://research-and-innovation.ec.europa.eu/document/download/2b6cf7e5-36ac-41cb-aab5-0d32050143dc_en?filename=ec_rtd_ai-guidelines.pdf
- Feldman, M. S., Orlikowski, W. J. (2011). Theorizing Practice and Practicing Theory. *Organization Science*, 22(5), 1240–1253. <http://hdl.handle.net/1721.1/66516>
- Feuerriegel, S., Hartmann, J., Janiesch, C., Zschech, P. (2024). Generative AI. *Business & Information Systems Engineering*, 66(1), 111–126. <https://doi.org/10.1007/s12599-023-00834-7>
- Floridi, L. (2020). AI and Its New Winter: From Myths to Realities. *Philosophy & Technology*, 33, 1–3. <https://doi.org/10.1007/s13347-020-00396-6>
- Geels, F. W. (2002). Technological Transitions as Evolutionary Reconfiguration Processes: A Multi-level Perspective and a Case-study. *Research Policy*, 31(8–9), 1257–1274. [https://doi.org/10.1016/S0048-7333\(02\)00062-8](https://doi.org/10.1016/S0048-7333(02)00062-8)
- Geels, F. W. (2011). The Multi-level Perspective on Sustainability Transitions: Responses to Seven Criticisms. *Environmental Innovation and Societal Transitions*, 1(1), 24–40. <https://doi.org/10.1016/j.eist.2011.02.002>
- Götz, F. M., Maertens, R., Loomba, S., van der Linden, S. (2023). Let the Algorithm Speak: How to Use Neural Networks for Automatic Item Generation in Psychological Scale Development. *Psychological Methods*, 29(3), 494–518. <https://doi.org/10.1037/met0000540>
- Grossmann, I., Feinberg, M., Parker, D. C., Christakis, N. A., Tetlock, P. E., Cunningham, W. A. (2023). AI and the Transformation of Social Science Research. *Science*, 380(6650), 1108–1109. <https://doi.org/10.1126/science.adi1778>
- Haman, M., Školník, M. (2023). Using ChatGPT to Conduct a Literature Review. *Accountability in Research*, 31(8), 1244–1246. <https://doi.org/10.1080/08989621.2023.2185514>
- Hanchard, M. (2024). Towards a Practice-orientated Digital Sociology. W: *Engaging with Digital Maps. Our Knowledgeable Deferral to Rough Guides* (s. 93–131). Palgrave Macmillan.

- Hofmann, V., Kalluri, P. R., Jurafsky, D., King, S. (2024). Dialect Prejudice Predicts AI Decisions about People's Character, Employability, and Criminality. *arXiv*. <https://doi.org/10.48550/arXiv.2403.00742>
- Huang, J., Tan, M. (2023). The Role of ChatGPT in Scientific Communication: Writing Better Scientific Review Articles. *American Journal of Cancer Research*, 13(4), 1148.
- Hui, A., Schatzki, T., Shove, E. (red.). (2017). *The Nexus of Practices. Connections, Constellations, Practitioners*. Routledge.
- Jain, S., Kumar, A., Roy, T., Shinde, K., Vignesh, G., Tondulkar, R. (2024). SciSpace Literature Review: Harnessing AI for Effortless Scientific Discovery. W: *European Conference on Information Retrieval* (s. 256–260). Springer Nature Switzerland.
- Joyce, K., Smith-Doerr, L., Alegria, S., Bell, S., Cruz, T. i in. (2021). Toward a Sociology of Artificial Intelligence: A Call for Research on Inequalities and Structural Change. *Socius*, 7. <https://doi.org/10.1177/2378023121999581>
- Keller, M., Sahakian, M., Hirt, L. F. (2022). Connecting the Multi-level-perspective and Social Practice Approach for Sustainable Transitions. *Environmental Innovation and Societal Transitions*, 44, 14–28. <https://doi.org/10.1016/j.eist.2022.05.004>
- Krajewski, M. (2023). Wstęp. Mapując kontrowersje teorii praktyk społecznych. *Studia Socjologiczne*, 4(251), 11–16.
- Lansing, J. S. (2003). Complex Adaptive Systems. *Annual Review of Anthropology*, 32(1), 183–204. <https://doi.org/10.1146/annurev.anthro.32.061002.093440>
- Law, T., McCall, L. (2024). Artificial Intelligence Policymaking: An Agenda for Sociological Research. *Socius*, 10, 1–13. <https://doi.org/10.1177/23780231241261596>
- Law, J. (2009). Seeing Like a Survey. *Cultural Sociology*, 3(2), 239–256. <https://doi.org/10.1177/1749975509105533>
- Li, A., Chen, H., Namkoong, H., Peng, T. (2025). LLM Generated Persona is a Promise with a Catch. *arXiv*. <https://doi.org/10.48550/arXiv.2503.16527>
- Lu, C., Lu, C., Lange, R. T., Foerster, J., Clune, J., Ha, D. (2024). The AI Scientist: Towards Fully Automated Open-ended Scientific Discovery. *arXiv*. <https://doi.org/10.48550/arXiv.2408.06292>
- Luybashenko, I. (2025). Analizy jakościowe wspierane gen AI. W: K. Olejniczak, D. Batorski, J. Pokorski (red.), *Generatywna AI w badaniach. Praktyczne zastosowania w ewaluacji polityk publicznych* (s. 131–153). Polska Agencja Rozwoju Przedsiębiorczości.
- Lupton, D. (2014). *Digital Sociology*. Routledge.
- Maclure, J. (2020). The New AI Spring: A Deflationary View. *AI & Society*, 35(3), 747–750.
- Marres, N. (2017). *Digital Sociology: The Reinvention of Social Research*. John Wiley & Sons.
- Marres, N., Guggenheim, M., Wilkie, A. (2018). *Inventing the Social*. Mattering Press.
- McCarthy, J., Minsky, M. L., Rochester, N., Shannon, C. E. (2006). A Proposal for the Dartmouth Summer Research Project on Artificial Intelligence, August 31, 1955. *AI Magazine*, 27(4), 12. <https://doi.org/10.1609/aimag.v27i4.1904>
- Miller, J. H., Page, S. E. (2007). *Complex Adaptive Systems: An Introduction to Computational Models of Social Life*. Princeton University Press.

- Morgan, D. L. (2023). Exploring the Use of Artificial Intelligence for Qualitative Data Analysis: The Case of ChatGPT. *International Journal of Qualitative Methods*, 22, <https://doi.org/10.1177/16094069231211248>
- Morgan, D. L., Nica, A. (2020). Iterative Thematic Inquiry: A New Method for Analyzing Qualitative Data. *International Journal of Qualitative Methods*, 19. <https://doi.org/10.1177/1609406920955118>
- Obermeyer, Z., Powers, B., Vogeli, C., Mullainathan, S. (2019). Dissecting Racial Bias in an Algorithm Used to Manage the Health of Populations. *Science*, 366(6464), 447–453. <https://doi.org/10.1126/science.aax2342>
- Osborne, T., Rose, N. (1999). Do the Social Sciences Create Phenomena? The Example of Public Opinion Research. *The British Journal of Sociology*, 50(3), 367–396. <https://doi.org/10.1111/j.1468-4446.1999.00367.x>
- Oxford University Press. (2024). *Researchers and AI: Survey Findings*, <https://academic.oup.com/pages/ai-survey-findings>
- Paduraru, C., Cristea, R., Stefanescu, A. (2024). Adaptive Questionnaire Design Using AI Agents for People Profiling. W: ICAART (3) (s. 633–640).
- Pantzar, M., Shove, E. (2010). Understanding Innovation in Practice: A Discussion of the Production and Re-production of Nordic Walking. *Technology Analysis & Strategic Management*, 22(4), 447–461.
- Park, J. S., Zou, C. Q., Shaw, A., Hill, B. M., Cai, C., i in. (2024). Generative Agent Simulations of 1,000 People. *arXiv*. <https://doi.org/10.48550/arXiv.2411.10109>
- Peñalvo, F. J. G., Ingelmo, A. V. (2023). What do We Mean by GenAI? A Systematic Mapping of the Evolution, Trends, and Techniques Involved in Generative AI. *IJI-MAI*, 8(4), 7–16. <https://doi.org/10.9781/ijimai.2023.07.006>
- Petticrew, M., Roberts, H. (2008). *Systematic Reviews in the Social Sciences: A Practical Guide*. John Wiley & Sons.
- Radford, A., Narasimhan, K., Salimans, T., Sutskever, I. (2018). Improving Language Understanding by Generative Pre-training. <https://api.semanticscholar.org/CorpusID:49313245>
- Rani, V., Nabi, S. T., Kumar, M., Mittal, A., Kumar, K. (2023). Self-supervised Learning: A Succinct Review. *Archives of Computational Methods in Engineering*, 30(4), 2761–2775. <https://doi.org/10.1007/s11831-023-09884-2>
- Rip, A., Voß, J. P. (2019). Umbrella Terms as Mediators in the Governance of Emerging Science and Technology. W: A. Rip, *Nanotechnology and Its Governance* (s. 10–33). Routledge.
- Røpke, I. (2009). Theories of Practice – New Inspiration for Ecological Economic Studies on Consumption. *Ecological Economics*, 68(10), 2490–2497. <https://doi.org/10.1016/j.ecolecon.2009.05.015>
- Rudin, C. (2019). Stop Explaining Black Box Machine Learning Models for High-Stakes Decisions and Use Interpretable Models Instead. *Nature Machine Intelligence*, 1(5), 206–215. <https://doi.org/10.1038/s42256-019-0048-x>
- Rudnicki, S. (2023). *Wystarczająco dobra wiedza. Badania user experience a wytwarzanie praktycznie użytecznej wiedzy o świecie społecznym*. Wydawnictwo Naukowe Scholar.

- Ruppert, E., Law, J., Savage, M. (2013). Reassembling Social Science Methods: The Challenge of Digital Devices. *Theory, Culture & Society*, 30(4), 22–46. <https://doi.org/10.1177/0263276413484941>
- Russell Group. (2023). Russell Group Statement of Principles on Artificial Intelligence in Higher Education. https://russellgroup.ac.uk/media/6137/rg_ai_principles-final.pdf
- Sampson, A. (2024). Using AI to Make Social Research More Inclusive and Impactful. <https://faculty.ai/insights/articles/using-ai-to-make-social-research-more-inclusive-and-impactful>
- Savage, M., Burrows, R. (2007). The Coming Crisis of Empirical Sociology. *Sociology*, 41(5), 885–899. <https://doi.org/10.1177/0038038507080443>
- Schatzki, T. (2001). Introduction: Practice Theory. W: T. Schatzki, K. Knorr-Cetina, E. von Savigny (red.), *The Practice Turn in Contemporary Theory* (s. 10–23). Routledge.
- Schmidgall, S., Su, Y., Wang, Z., Sun, X., Wu, J., Yu, X., i in. (2025). Agent Laboratory: Using LLM Agents as Research Assistants. *arXiv*. <https://doi.org/10.48550/arXiv.2501.04227>
- Shove, E. (2010). Beyond the ABC: Climate Change Policy and Theories of Social change. *Environment and planning A*, 42(6), 1273–1285.
- Shove, E., Pantzar, M. (2005). Consumers, Producers and Practices: Understanding the Invention and Reinvention of Nordic Walking. *Journal of Consumer Culture*, 5(1), 43–64. <http://dx.doi.org/10.1177/1469540505049846>
- Shove, E., Pantzar, M., Watson, M. (2012). *The Dynamics of Social Practice: Everyday Life and How It Changes*. Sage.
- Shults, F. L. (2025). Simulating Theory and Society: How Multi-agent Artificial Intelligence Modeling Contributes to Renewal and Critique in Social Theory. *Theory and Society*, 1–23. <https://doi.org/10.1007/s11186-025-09606-6>
- Si, C., Yang, D., Hashimoto, T. (2024). Can LLMs Generate Novel Research Ideas? A Large-scale Human Study with 100+ NLP Researchers. *arXiv*. <https://doi.org/10.48550/arXiv.2409.04109>
- Sikorska, M. (2018). Teorie praktyk jako alternatywa dla badań nad rodziną prowadzonych w Polsce. *Studia Socjologiczne*, 229(2), 31–63. <https://doi.org/10.24425/122463>
- Smagacz-Poziemska, M., Bukowski, A., Kurnicki, K. (2018). „Wspólnota parkingowania”. Praktyki parkowania na osiedlach wielkomiejskich i ich strukturalne konsekwencje. *Studia Socjologiczne*, 228(1), 117–142. <https://doi.org/10.24425/119089>
- Stanford University. (n.d.). Responsible AI: Ethical Principles and Guidelines. Stanford University IT. <https://uit.stanford.edu/security/responsibleai>
- Stehr, N., Grundman, R. (2001). The Authority of Complexity. *British Journal of Sociology*, 52(2), 313–329. <https://doi.org/10.1080/00071310120045006>
- Stokel-Walker, C., Van Noorden, R. (2023). What ChatGPT and Generative AI Mean for Science. *Nature*, 614(7947), 214–216. <https://doi.org/10.1038/d41586-023-00340-6>
- Torre-López, J., Ramírez, A., Romero, J. R. (2023). Artificial Intelligence to Automate the Systematic Review of Scientific Literature. *Computing*, 105(10), 2171–2194. <https://doi.org/10.1007/s00607-023-01181-x>

- Törnberg, P. (2023). How to Use LLMs for Text Analysis. *arXiv*. <https://doi.org/10.48550/arXiv.2307.13106>
- University of Oxford, Information Security. (n.d.). Guidelines on the Use of Generative AI. University of Oxford. <https://communications.admin.ox.ac.uk/communications-resources/ai-guidance#collapse5321536>
- Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L. i in. (2017). Attention is All You Need. *Advances in Neural Information Processing Systems*, 30. <https://doi.org/10.48550/arXiv.1706.03762>
- Wang, H., Fu, T., Du, Y., Gao, W., Huang, K. i in. (2023). Scientific Discovery in the Age of Artificial Intelligence. *Nature*, 620(7972), 47–60. <https://doi.org/10.1038/s41586-023-06221-2>
- Warde, A. (2014). After Taste: Culture, Consumption and Theories of Practice. *Journal of Consumer Culture*, 14(3), 279–303. <http://dx.doi.org/10.1177/1469540514547828>
- Whitfield, S., Hofmann, M. A. (2023). Elicit: AI Literature Review Research Assistant. *Public Services Quarterly*, 19(3), 201–207. <http://dx.doi.org/10.1080/15228959.2023.2224125>
- Whittaker, M., Crawford, K., Dobbe, R., Fried, G., Kaziunas, E. i in. (2018). W: AI Now Report 2018 (s. 1–62). AI Now Institute at New York University.
- Williams, R. T. (2024). Paradigm Shifts: Exploring AI's Influence on Qualitative Inquiry and Analysis. *Frontiers in Research Metrics and Analytics*, 9. <http://dx.doi.org/10.3389/frma.2024.1331589>
- Xu, R., Sun, Y., Ren, M., Guo, S., Pan, R. i in. (2024). AI for Social Science and Social Science of AI: A Survey. *Information Processing & Management*, 61(3). <https://doi.org/10.48550/arXiv.2401.11839>
- Yang, Z., Du, X., Li, J., Zheng, J., Poria, S., Cambria, E. (2023). Large Language Models for Automated Open-domain Scientific Hypotheses Discovery. *arXiv*. <https://doi.org/10.48550/arXiv.2309.02726>
- You, Z., Lee, H., Mishra, S., Jeoung, S., Mishra, A., Kim, J., Diesner, J. (2024). Beyond Binary Gender Labels: Revealing Gender Biases in LLMs through Gender-neutral Name Predictions. *arXiv*. <https://doi.org/10.48550/arXiv.2407.05271>
- Yuan, Z., Yuan, H., Tan, C., Wang, W., Huang, S. (2023). How Well Do Large Language Models Perform in Arithmetic Tasks? *arXiv*. <https://doi.org/10.48550/arXiv.2304.02015>

