

# Senso comune artificiale

## Egemonia culturale e intelligenza artificiale generativa\*

Maria Francesca Murru, Donatella Selva

### Abstract

This article explores the role of Generative AI (GenAI) as an “epistemic gatekeeping” technology capable of influencing collective knowledge and individual epistemic agency. Through an experimental and qualitative approach, the study analyzes interactions between two Large Language Models (ChatGPT 4.0 and Claude 3.5 Sonnet) and ideal-typical users with different gender and political orientation (right/left) profiles, prompted on the controversial topic of surrogacy. The analysis focuses on the “genericity” of the outputs as a socio-technical code that produces an artificial common sense: a plausible and familiar discursive repertoire that mediates between content standardization and personalization. The findings reveal that while standardization is achieved through recurring discursive frames (legal, ethical, socio-economic), personalization adapts the content to the user’s profile, reinforcing stereotypes and crystallizing ideological positions. The article concludes that the primary risk of GenAI is not so much disinformation, but rather the progressive sedimentation of standardized discursive repertoires that, while appearing pluralistic, normalize a “stereotyped variety” and reinforce cultural hegemony by marginalizing non-mainstream positions.

**Keywords:** Generative artificial intelligence; large language models; artificial common sense; cultural hegemony; surrogacy.

**ISSN:** 03928667 (print) 18277969 (digital)

**DOI:** 10.26350/001200\_000250

Creative Commons License CC-BY-NC-ND 4.0

### 1. Introduzione

Da diversi anni sappiamo che almeno la metà del traffico su internet è generato da bot, una stima certamente al ribasso viste le relative limitazioni

Contributo sottoposto a *double-blind peer review*.

\* Maria Francesca Murru, Università degli Studi di Bergamo – mariafrancesca.murru@unibg.it; Donatella Selva – Luiss Università di Roma, dselva@luiss.it. Il saggio è il risultato di un lavoro condiviso cui le autrici hanno preso parte con uguale responsabilità, e per questo motivo risultano qui elencate in semplice ordine alfabetico.

dei software di bot detection<sup>1</sup>. Con la diffusione dei sistemi di intelligenza artificiale generativa (d'ora in poi: GenAI), si comincia a delineare il quadro di una società in cui possono circolare indifferentemente discorsi prodotti da esseri umani e discorsi prodotti "artificialmente", con il supporto più o meno invasivo di tali sistemi. Sappiamo che la conoscenza prodotta dalle macchine è spesso depositaria di un surplus fiduciario, esito di una forma di autorità algoritmica<sup>2</sup> che si vuole oggettiva e infallibile, al contrario delle capacità di giudizio umane, per loro natura soggettive ed esposte a bias di tipo emotivo<sup>3</sup>. Ne consegue il cosiddetto *automation bias*, studiato soprattutto in relazione ai sistemi di automazione a supporto di processi decisionali<sup>4</sup> ma evidente anche rispetto ai *large language models* (LLMs). Non mancano tuttavia evidenze scientifiche che mostrano una realtà più sfumata e contraddittoria, in cui la credibilità della GenAI come fonte di informazione è mediata da altri fattori, come la fiducia nel sistema dei media nel suo complesso o il posizionamento politico degli attori sociali<sup>5</sup>. In generale, sembra più probabile che le persone siano portate a dichiarare di "fidarsi abbastanza" dei risultati piuttosto che "fidarsi pienamente"<sup>6</sup>. Oltre a segnalare un certo grado di scetticismo, questa fiducia sufficiente ma non piena sembra essere indizio di un compromesso incipiente tra le esigenze informative del pubblico e ciò che la GenAI concretamente offre: informazioni non perfette eppure *sufficientemente buone* per gli obiettivi più o meno impegnativi della vita quotidiana e professionale. Come queste prime evidenze empiriche sembrano suggerire, la GenAI si sta velocemente radicando anche presso il pubblico generalista come una "tecnologia epistemica"<sup>7</sup>, vale a dire come un sistema socio-tecnico in grado

<sup>1</sup> E. Ferrara *et al.*, "The Rise of Social Bots", *Communications of the ACM*, 59, 7 (2016): 96-104; E. Esposito, *Comunicazione artificiale*, Milano: Egea, 2022.

<sup>2</sup> M. Airoidi, *Machine habitus*, Roma: LUISS University Press, 2024.

<sup>3</sup> T. Araujo *et al.*, "In AI We Trust? Perceptions About Automated Decision Making by Artificial Intelligence", *AI & Society*, 35, 6 (2020): 611-623; S. Natale, *Macchine ingannevoli*, Torino: Einaudi, 2022.

<sup>4</sup> Databricks (2024) "Automation Bias" Databricks. <https://www.databricks.com/glossary/automation-bias>.

<sup>5</sup> T. Araujo *et al.*, "Humans vs. AI: The Role of Trust, Political Attitudes, and Individual Characteristics on Perceptions About Automated Decision Making Across Europe", *International Journal of Communication*, 17, 28 (2023); S. Toff, F. Simon, "«Or They Could Just Not Use It?»: The Dilemma of AI Disclosure for Audience Trust in News", *The International Journal of Press/Politics*, 2024. <https://doi.org/10.1177/19401612241308697>.

<sup>6</sup> R. Fletcher, R.K. Nielsen, "What Does the Public in Six Countries Think of Generative AI in News?", *Reuters Institute for the Study of Journalism*, 2024, DOI: 10.60625/risj-4zb8-cg87.

<sup>7</sup> R. Alvarado, "AI as an Epistemic Technology", *Science and Engineering Ethics*, 29, 32 (2023). DOI: <https://doi.org/10.1007/s11948-023-00451-3>.

di esercitare un processo multistratificato di “gatekeeping”<sup>8</sup> negli spazi e nei processi del discorso pubblico. Non si tratta d'altra parte di un'evoluzione imprevista, se si guarda al quadro socio-tecnico<sup>9</sup> che la contraddistingue e che appare, fin dalle sue origini, primariamente orientato alla valorizzazione e allo svolgimento di operazioni epistemiche. Non solo si tratta di una tecnologia progettata e ideata per condurre analisi e previsioni in ambienti sociali in cui si elabora e si produce sapere, ma le sue stesse capacità generative si basano su un processo computazionale di accumulazione, sintesi e replicazione di versioni analoghe ma non identiche di quei contenuti simbolici che, nel corso degli ultimi decenni, sono stati diffusi e condivisi in formato digitale. Lo stesso *prompting*, inteso come pratica che attinge al linguaggio naturale per orientare e istruire i foundation models<sup>10</sup> verso i task desiderati, rappresenta una forma di interazione umano-macchina che mette a disposizione del sistema potenzialmente infiniti repertori culturali e linguistici utili ad allineare, orientare e attivare la conoscenza latente contenuta nei modelli e incrementarne, di conseguenza, il potenziale commerciale<sup>11</sup>.

In quanto risorsa pubblica nel processo collettivo di elaborazione della conoscenza e di costruzione dei significati, la GenAI richiede quindi di essere studiata come una forma emergente e multistratificata di mediazione epistémica, in grado di condizionare ampiezza e varietà di quei repertori discorsivi che alimentano la conoscenza collettiva e che abitano l'ecosistema mediale e lo spazio pubblico. Nello specifico, a necessitare una chiara messa a fuoco è il modo in cui la mediazione dei discorsi e contenuti operata dalla GenAI incide sulla cosiddetta “agency epistémica”<sup>12</sup> degli attori sociali, vale a dire sulla capacità di esercitare una qualche forma di controllo sulle modalità di elaborazione delle proprie conoscenze e convinzioni. Questa possibilità di controllo riflessivo rappresenta uno dei pilastri dell'agency democratica. I principi ideali e partecipativi della democrazia richiedono infatti che i cit-

<sup>8</sup> M. Miragoli, “Conformism, Ignorance & Injustice: AI as a Tool of Epistemic Oppression”, *Episteme*, Published online, 2024: 1-19. DOI: doi:10.1017/epi.2024.11.

<sup>9</sup> P. Flichy, *L'innovazione tecnologica. Le teorie dell'Innovazione di fronte alla rivoluzione digitale*, Milano: Feltrinelli, 1996.

<sup>10</sup> R. Bommasani *et al.*, “On the Opportunities and Risks of Foundation Models”, Center for Research on Foundation Models (CRFM) Stanford Institute for Human-Centered Artificial Intelligence (HAI) Stanford University, 2022. DOI: <https://doi.org/10.48550/arXiv.2108.07258>.

<sup>11</sup> S. Burkhardt, B. Rieder, “Foundation Models Are Platform Models: Prompting and the Political Economy of AI”, *Big Data & Society*, 11, 2 (2024): 1-15.

<sup>12</sup> M. Coeckelbergh, “Democracy, Epistemic Agency, and AI: Political Epistemology in Times of Artificial Intelligence”, *AI Ethics*, 3 (2023): 1341-1350. DOI: <https://doi.org/10.1007/s43681-022-00239-4>.

tadini possano non solo esprimere pubblicamente le proprie convinzioni ma anche conoscere e, se necessario, modificare quelle condizioni di contesto che consentono loro di maturare opinioni più o meno accurate su questioni rilevanti per la vita individuale e collettiva.

Obiettivo di questo contributo è la presentazione e la discussione di un approccio sperimentale all'analisi di una particolare forma di "gatekeeping epistemico" esercitata dalla GenAI. Se finora buona parte della letteratura che ha indagato rischi e potenzialità democratiche della GenAI ha prediletto l'analisi dei bias (di presentazione e rappresentazione, politici e di genere), nell'analisi sperimentale qui proposta l'attenzione si sposta dai contenuti al setting interazionale tra un utente idealtipico e due LLMs (ChatGPT 4.0 e Claude 3.5 Sonnet). Riteniamo infatti che, solo aderendo alla natura interattiva e contingente<sup>13</sup> della GenAI, sia possibile raccogliere indizi utili a mappare le più ampie logiche culturali e sociali di "riduzione e ottimizzazione"<sup>14</sup> indotte dal paradigma della datificazione, di cui l'intelligenza artificiale rappresenta la più evoluta e recente incarnazione.

Nell'indagare una delle forme possibili di gatekeeping epistemico della GenAI, l'attenzione del presente contributo si concentra sulla "genericità" come codice socio-tecnico che informa l'elaborazione dei suoi contenuti e che è all'origine di quella sensazione di plausibilità e di familiarità ingannevole<sup>15</sup> propria degli LLM. Proponiamo di considerare tale genericità come produzione di una forma di *senso comune artificiale*, mostrandone la doppia articolazione tecnica e politica, ovvero in quanto parametro di misurazione della performance degli LLM e insieme base di conoscenza indiscussa e generale in cui si annida l'ideologia egemonica di un dato contesto sociale<sup>16</sup>. La sua concettualizzazione è inoltre accompagnata da un'analisi sperimentale volta in primo luogo a farlo emergere empiricamente, mediante un setting interazionale di interrogazioni idealtipiche, e in secondo luogo ad indagarne l'articolazione concreta sotto forma di impliciti discorsivi annidati nella plausibilità apparente degli output della GenAI. Mediante un'analisi dei contenuti generati da due LLMs su un tema polarizzante e divisivo come la gestazione per altri, lo studio apre a una riflessione sul rapporto tra la GenAI

<sup>13</sup> E. Esposito, *Comunicazione artificiale*. Milano: Egea, 2022; A.L. Guzman, "Talking about «Talking with Machines»", in *The SAGE Handbook of Human-Machine Communication*, edited by A.L. Guzman, R. McEven, S. Jones S., Thousand Oaks, CA: Sage, 2023, 243-251.

<sup>14</sup> A.L. Hoffmann, "Where Fairness Fails: Data, Algorithms, and the Limits of Anti-discrimination Discourse", *Information, Communication & Society*, 22, 7 (2019): 900-915.

<sup>15</sup> Natale, *Macchine ingannevoli*.

<sup>16</sup> A. Gramsci, *Quaderni del carcere*, Torino: Einaudi, 1977.

e i repertori discorsivi dello spazio pubblico, osservando in che modo la genericità trovi declinazioni differenziate (o meglio, personalizzate) a seconda dei diversi utenti e dei posizionamenti ideologici da essi sollecitati.

## 2. Senso comune e conoscenza generica

Sappiamo che l'IA attinge a corpora abbastanza ampi da risultare generalisti e perciò in linea con il senso comune diffuso; è la tesi che vede la GenAI come niente più che un “pappagallo stocastico”, che riproduce stereotipi e rappresentazioni in base alla loro maggiore o minore frequenza<sup>17</sup>. Dal momento che la GenAI non conosce e non comprende, le sue risposte si basano su modelli di calcolo della probabilità che dopo un determinato token appaia un altro<sup>18</sup>. Rispetto a questo, sono attuati dei correttivi per evitare che le risposte tendano a somigliarsi, utilizzando un range ristretto di token super-frequenti: la cosiddetta “temperatura” regola l'equilibrio tra calcolo delle probabilità e uso di token insoliti, in modo che l'output risulti in qualche modo originale e diverso da tutti gli altri<sup>19</sup>. Nel caso in cui la risposta non sia contenuta nel database utilizzato per l'addestramento del modello, la GenAI tenderà a selezionarne una tra quelle più probabili, andando incontro al rischio di “allucinazioni”, assegnando priorità alla scorrevolezza sintattica più che all'accuratezza fattuale<sup>20</sup>. Dal momento che tali testi sono già il risultato di un calcolo delle probabilità, la selezione basata su di essi privilegerà in modo sempre più marcato le forme comunicative più frequenti, attestandosi così su un “*fairly banal standard*”<sup>21</sup> a scapito di espressioni più creative

<sup>17</sup> E. Bender *et al.*, “On the Dangers of Stochastic Parrots: Can Language Models Be Too Big? 🦜”, *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*, 2021; T. Teubner *et al.*, “Welcome to the Era of ChatGPT”, *Business & Information Systems Engineering*, 65 (2023): 95-101. DOI: <https://doi.org/10.1007/s12599-023-00795-x>.

<sup>18</sup> Il token è una parola, ma può essere anche un insieme di parole, o, nel caso di immagini e audiovisivi, pixel e note.

<sup>19</sup> M. Gillings, T. Kohn, G. Mautner, “The Rise of Large Language Models: Challenges for Critical Discourse Studies”, *Critical Discourse Studies*, 2024: 1-17. DOI: <https://doi.org/10.1080/17405904.2024.2373733>.

<sup>20</sup> *Ibid.*

<sup>21</sup> Padmakumar e He hanno dimostrato che lo stesso modello (ChatGPT 3), se istruito sulla base di dati di feedback, conduceva a una significativa riduzione della diversità degli output e a una progressiva indistinguibilità degli autori. Dunque, il contrario di generico è specifico, laddove la specificità è relativa alla “cifra autoriale”, cioè a uno stile di scrittura specificamente riconducibile a un determinato autore (ad es. vocabolario, sintassi, temi ecc.). V. Padmakumar, H. He, “Does Writing with Language Models Reduce Content Diversity?”, *Computer Science arXiv preprint*, 2023. DOI: <https://doi.org/10.48550/arXiv.2309.05196>.

e originali<sup>22</sup>. Intervengono contestualmente ripetuti (e opachi) processi di fine-tuning volti ad allineare gli output a criteri di desiderabilità sociale mediante un rinforzo dei pattern positivi e un'attenuazione di quelli negativi, soprattutto in relazione a temi e stili discorsivi politicamente controversi. A tal proposito, Noels *et al.*<sup>23</sup> hanno osservato forme variegiate di auto-censura da parte della GenAI, distinguendo tra “hard censorship” (ad es. messaggi di errore) e “soft censorship”, più insidiosa perché basata sullo stile di presentazione degli argomenti (ad es. alcuni temi, attori o posizioni possono essere minimizzati o omessi). Tutti questi interventi mirano ad allineare i LLMs a un “commonsense” generale trasformato dalle scienze informatiche in parametro di misurazione della performance e inteso come “comprensione ampia e fondativa del mondo”<sup>24</sup>, condivisa dalla maggior parte delle persone. Nel momento in cui questi modelli sono usati per produrre contenuti politici, è la stessa ambiguità semantica del concetto di “common sense” a offrire nuovi spunti per una teoria critica della GenAI<sup>25</sup>. Come aveva già avuto modo di osservare Antonio Gramsci<sup>26</sup>, nella lingua italiana si distingue tra *senso comune*, per identificare un bagaglio di conoscenza popolare condivisa e in fondo un po' banale, e *buon senso*, che prende da questo bagaglio alcune indicazioni pratiche per agire nel mondo nel miglior modo possibile, cioè nel modo più vicino al senso comune perché quest'ultimo è un “senso per il giusto e per il bene comune”<sup>27</sup>. L'assenza di tale distinzione nella lingua inglese fa sì che l'espressione *common sense* racchiuda entrambi gli aspetti, tanto di normatività implicita quanto di apparente naturalezza dovuta alla sua diffusione generalizzata. Teologi, filosofi e antropologi nel corso del tempo hanno definito il senso comune in molti modi, con una stagione di particolare interesse che si registra intorno agli anni Settanta e Ottanta<sup>28</sup>. Da una prospettiva

<sup>22</sup> Gillings, Kohn, Mautner, “The Rise of Large Language Models: Challenges for Critical Discourse Studies”.

<sup>23</sup> Noels *et al.*, “What Large Language Models Do Not Talk About: An Empirical Study of Moderation and Censorship Practices”, *Computer Science arXiv preprint* (2025). DOI: <https://doi.org/10.48550/arXiv.2504.03803>.

<sup>24</sup> S. Shen *et al.*, “Understanding the Capabilities and Limitations of Large Language Models for Cultural Commonsense”, *Proceedings of the 2024 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies* (Volume 1: Long Papers), Association for Computational Linguistics, Mexico City, Mexico, 2024, 5668-5680.

<sup>25</sup> S. Lindgren, *Handbook of Critical Studies in Artificial Intelligence*, Cheltenham: Edward Elgar, 2023.

<sup>26</sup> Gramsci, *Quaderni del carcere*.

<sup>27</sup> P. Jedlowski, *Il sapere dell'esperienza*, Roma: Carocci, 2008.

<sup>28</sup> S. Rosenfeld, *Common Sense: A Political History*, Cambridge: Harvard University Press, 2011.

sociologica, l'aspetto fondamentale del senso comune risiede nei suoi effetti relazionali, ovvero nel fatto di rappresentare una risorsa conoscitiva comune funzionale al riconoscimento reciproco e a un'intesa primaria tra gli attori sociali. Secondo Harold Garfinkel<sup>29</sup>, il senso comune si riferisce a una forma di conoscenza delle strutture sociali basata sull'esperienza; più precisamente, coincide con quella base di azioni e inferenze socialmente legittimata che le persone utilizzano nella loro vita quotidiana, assumendo che gli altri facciano lo stesso. Nel senso comune si trovano i proverbi, le narrative e gli archetipi riprodotti nella cultura popolare in forma di fiabe, letteratura, film e serie televisive, le tradizioni, i valori per distinguere il bene dal male, ma anche una conoscenza pratica, come le istruzioni per reagire agli eventi della vita, per lo più fondate su un sapere che nasce dall'esperienza (anche quella mediata) e che si è tramandato attraverso le generazioni. Nell'ambito della sociologia dei media, l'attenzione si è spostata sui processi di visibilizzazione e invisibilizzazione, identificando un doppio binario: da una parte, la tendenza dei media a marginalizzare la diversità (di genere, culturale, religiosa, politica ecc.) attraverso processi di silenziamento o annichilimento simbolico; dall'altra, la tendenza a rappresentarla sotto forma di stereotipi talvolta stigmatizzanti<sup>30</sup>. Nell'attuale ecosistema comunicativo, lo spazio per la rappresentazione di opinioni, posizioni e gruppi diversificati è decisamente più ampio rispetto al passato, anche se nell'ambito di una negoziazione continua con la tendenza all'omogeneizzazione culturale<sup>31</sup>. Come in altre industrie culturali, il senso comune prodotto dalla GenAI è il risultato emergente di una forma di genericità che si radica tanto nel modello economico sottostante quanto nella sua dimensione prettamente tecnica. A livello culturale, la genericità è raggiunta attraverso l'adesione a generi letterari, artistici o discorsivi riconoscibili (ad es. il comizio elettorale, l'articolo di cronaca o l'editoriale), codificati e ben radicati nell'immaginario comune; economicamente, la genericità consente l'adattamento a molti e diversi segmenti di pubblici; infine, tecnicamente, la genericità è una strategia per simulare l'intelligenza umana attraverso risposte che potrebbero appartenere a un qualsiasi essere umano, dove "qualsiasi" sta per "ogni essere umano" ma anche per "essere umano comune", cioè un idealtipo che si presume non caratterizzato e con una media capacità di accesso alle risorse culturali.

L'ipotesi che guida questo studio, dunque, è che una delle dimensio-

<sup>29</sup> H. Garfinkel, *Studies in Ethnomethodology*, Englewood Cliffs: Prentice-Hall, 1967.

<sup>30</sup> S. Hall, *Representation*, London: Sage, 1977.

<sup>31</sup> M. Sorice, *Partecipazione disconnessa*, Roma: Carocci, 2021.

ni decisive nel processo di gatekeeping epistemico<sup>32</sup> esercitato dalla GenAI consista nella veicolazione di un senso comune artificiale, pensato come un repertorio discorsivo generico di volta in volta personalizzato a seconda della cornice comunicativa instaurata nell'interazione umano-macchina e della personalizzazione dell'output sulla base della profilazione degli utenti. La personalizzazione offerta dagli LLM non contraddice la genericità, piuttosto può essere intesa come una sua particolare declinazione: il sistema seleziona e presenta all'utente le varianti di senso comune più coerenti con i posizionamenti ideologici e le caratteristiche demografiche che l'utente sceglie di rivelare. Per verificare questa ipotesi, l'approccio scelto è quello di simulare interazioni differenziate per tipo di utente, caratterizzandolo di volta in volta per posizionamento ideologico e genere in modo da sollecitare le "versioni" personalizzate di senso comune artificiale offerte dai modelli interpellati. Mediante un'analisi dei contenuti così ottenuti, ci chiediamo dunque quale sia la forma discorsiva assunta dal senso comune artificiale prodotto a seguito di interrogazioni riguardanti questioni politiche controverse.

### 3. Il design della ricerca

L'approccio sperimentale qui presentato scaturisce dall'obiettivo di indagare empiricamente il "senso comune artificiale". L'obiettivo è confrontare le risposte della GenAI fornite a diversi profili di interrogante, in modo da evidenziare quelle somiglianze di contenuti che, nel nostro quadro concettuale, possiamo far risalire al senso comune artificiale. Abbiamo privilegiato un metodo di indagine prettamente qualitativo (Denzin, 2016) ma che attinge alle tradizioni dell'*algorithmic auditing* e del *reverse engineering*<sup>33</sup>, in particolare nelle forme utilizzate da Gillespie<sup>34</sup> e da Volk *et al.*<sup>35</sup>. Pur avendo ra-

<sup>32</sup> Si noti che questo ruolo è concepibile come un'agency potenziale, almeno fino a quando gli utenti non decidono di utilizzare i testi prodotti dalla GenAI per produrre/riprodurre conoscenza.

<sup>33</sup> L'*algorithmic auditing* è un insieme di approcci usati per osservare il comportamento degli algoritmi in base a criteri definiti di volta in volta dagli obiettivi della ricerca. All'interno di questo insieme, è possibile rintracciare il *reverse engineering*, un procedimento che si basa sull'input di comandi (coerenti con la sintassi della piattaforma o degli algoritmi da studiare) in modo da sollecitare output che costituiranno un dataset, da analizzare secondo tecniche più o meno consolidate per dedurre quali sono i meccanismi di funzionamento degli algoritmi. Ad esempio, nel caso della GenAI conversazionale, gli input sono i prompt di istruzioni e gli output sono le risposte.

<sup>34</sup> T. Gillespie, "Generative AI and the Politics of Visibility", *Big Data & Society*, 11, 2 (2024): 1-14. DOI: <https://doi.org/10.1177/20539517241252131>.

<sup>35</sup> S.C. Volk *et al.*, "How Generative Artificial Intelligence Portrays Science: Interview-



dici comuni, i due lavori sono esemplificativi della diversità di approcci allo studio della GenAI. In particolare, Gillespie segue un uso che si sta attestando in letteratura, ovvero quello di reiterare le interrogazioni un certo numero di volte (in questo caso 50 per ciascun prompt), in modo tale da minimizzare il rischio che l'output sia un semplice "errore" dell'IA (le cosiddette "allucinazioni") e dare maggiore validità ai risultati; Volk *et al.*, invece, non adotta questo modello e usano il prompt solo una volta perché interessati a condurre un'intervista in profondità con l'AI. Nel nostro caso, abbiamo scelto di non reiterare le interrogazioni perché adottiamo un paradigma metodologico di tipo interpretativo, coerente con la prospettiva di sociologia della cultura entro cui ci muoviamo, con la tradizione di studi sul senso comune e con le finalità specifiche di questo progetto di ricerca. Il nostro interesse è dunque legato agli output potenziali che si possono ottenere in contesti reali: proprio perché la GenAI funziona in modo stocastico, qualsiasi risultato ottenuto rientra nei repertori discorsivi che potenzialmente possono essere fruiti e usati dagli utenti. Segnaliamo inoltre che l'opacità strutturale del modello implica la difficoltà strutturale di stabilire una soglia minima standard di iterazioni necessarie a garantire validità statistica ai risultati.

A partire da questa letteratura, lo studio ha deciso di concentrarsi su due chatbot, ChatGPT e Claude. Se il primo è stato selezionato perché tra più diffusi e riconoscibili, la scelta di Claude risponde invece a una logica di differenziazione radicale rispetto al benchmark rappresentato da ChatGPT: se quest'ultimo è il servizio di GenAI più popolare, in termini di utenti Claude occupa una posizione medio-bassa sul mercato<sup>36</sup>. Inoltre, al momento della selezione dei servizi di GenAI, Claude veniva indicato da più parti come il principale e più promettente rivale commerciale di ChatGPT, differenziandosi anche sulla base di una filosofia detta di "Constitutional AI", più attenta agli aspetti etici<sup>37</sup>. Per entrambi i servizi abbiamo utilizzato versioni a pa-

ing ChatGPT from the Perspective of Different Audience Segments", *Public Understanding of Science*, 34, 2 (2024): 132-153.

<sup>36</sup> ChatGPT si distingue come leader indiscusso del mercato con oltre l'82,5% del traffico totale su servizi di GenAI, ovvero oltre due miliardi di visite globali e 500 milioni di utenti al mese (World Bank, 2024). Al momento dell'avvio della ricerca (novembre 2024), ChatGPT deteneva una posizione dominante sul mercato mondiale, con 4,8 miliardi di visite al mese per la versione web. Claude è stato rilasciato nel marzo 2023 con la finalità più o meno dichiarata di scalare questo primato, ma al momento della ricerca accoglieva 18,9 milioni di visite al mese (fonte Semrush, disponibile a: <https://backlinko.com/claude-users>).

<sup>37</sup> I. Belcic, C. Stryker, "What Is Claude AI?", IBM, 24/9/2024, available at: <https://www.ibm.com/think/topics/claude-ai>; W. Davis, "OpenAI Rival Anthropic Makes Its Claude Chatbot Even More Useful", *The Verge*, 21/11/2023, available at: <https://www.theverge.com/2023/11/21/23971070/anthropic-claude-2-1-openai-ai-chatbot-update-beta-tools>, ac-

gamento in modo da massimizzare la qualità degli output testuali; tuttavia, questo ha implicato la creazione di un account e il consenso a registrare i dati di utilizzo da parte degli stessi servizi, con il conseguente rischio che la personalizzazione<sup>38</sup> distorcesse i risultati. Per ovviare alle distorsioni che ne derivano, in linea con gli studi già citati, ogni interazione è stata svolta su una nuova schermata con cancellazione della precedente.

Rispetto ai contenuti, si è deciso di sollecitare i due chatbot su un tema che, in questo preciso momento storico, appare divisivo e polarizzante: la gestazione per altri (o maternità surrogata). La questione della gestazione per altri è stata oggetto di una guerra culturale<sup>39</sup> che ha portato di recente all'identificazione in Italia della maternità surrogata come reato universale, perseguibile anche se compiuta al di fuori dei confini italiani. L'ipotesi è che una questione così divisiva sia in grado non solo di massimizzare la varianza tra le molte rappresentazioni "generiche" possibili a cui la GenAI può attingere e che può riprodurre, ma anche di testare il limite tra il mantenimento della genericità e la forzatura introdotta dalla sollecitazione ad assumere un profilo personalizzato in termini di polarità politiche (destra/sinistra).

### 3.1. Una prospettiva performativa

Nell'ambito di questo approccio sperimentale, abbiamo predisposto la ricerca in modo tale da poter valutare gli output della GenAI dal punto di vista della loro performance ordinaria, legata a situazioni di scambio comunicativo verosimili con interroganti variamente profilati, così che potesse emergere la

cessed; W. Henshall, "What to Know About Claude 2, Anthropic's Rival to ChatGPT", *Time*, 18/7/2023, available at: <https://time.com/6295523/claude-2-anthropic-chatgpt/> accessed.

<sup>38</sup> L'interrogazione ai due chatbot è stata svolta nel corso del mese di novembre 2024. È noto che OpenAI investe nell'ottimizzazione della performance di ChatGPT attraverso la possibilità, per gli utenti registrati, di fornire istruzioni di personalizzazione e di registrarle in modo da essere automaticamente implementate ad ogni interazione (OpenAI "Custom Instructions for ChatGPT", 20/7/2023. Available at: <https://openai.com/index/custom-instructions-for-chatgpt/>), accessed [ottobre 2025]. Sappiamo di avere a che fare con una tecnologia che evolve molto rapidamente e in modo poco trasparente; molto probabilmente le interrogazioni ripetute nel mese corrente (luglio 2025), con versioni più evolute dello stesso chatbot e tramite API avrebbero restituito risultati differenti; tuttavia, riteniamo che il contributo dell'articolo risieda nella concettualizzazione e nella conseguente sperimentazione di un modello di analisi di quella background knowledge che abbiamo definito "senso comune artificiale" e che è per definizione suscettibile di tutte quelle variazioni derivanti dal setting interazionale, dalla contingenza strutturale propria del modello e dalla sua rapida e non trasparente evoluzione.

<sup>39</sup> E. De Blasio, D. Selva, *Gender and Culture Wars: A Genealogy of Media Representations*, Cham: Palgrave, 2024.

“lettura preferita” come esito del processo di personalizzazione attivato dal modello. La ricerca ha contemplato tre tipi di interroganti – i cittadini, i giornalisti, i politici – a partire dal presupposto che tutti e tre rappresentino figure idealtipiche della sfera pubblica, ciascuno a suo modo coinvolto nel processo di gatekeeping epistemico e nell’elaborazione e riproduzione dell’egemonia culturale. Per ragioni di spazio, ci concentriamo in questa sede esclusivamente sull’analisi degli output sollecitati dal primo dei tre interroganti. I prompt utilizzati sono riportati nella Tabella 1 e mostrano l’articolazione del setting secondo profili costanti dal punto di vista della nazionalità ma variabili in base al genere e al posizionamento politico-ideologico.

**Tabella 1** - Elenco dei prompt utilizzati per l’interrogante cittadino/a

|                                                                                            |                                                                                             |
|--------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------|
| Sono un <i>cittadino italiano</i> e sono interessato alla maternità surrogata.             | Sono una <i>cittadina italiana</i> e sono interessata alla maternità surrogata.             |
| Sono un <i>cittadino italiano di destra</i> e sono interessato alla maternità surrogata.   | Sono una <i>cittadina italiana di destra</i> e sono interessato alla maternità surrogata.   |
| Sono un <i>cittadino italiano di sinistra</i> e sono interessato alla maternità surrogata. | Sono una <i>cittadina italiana di sinistra</i> e sono interessato alla maternità surrogata. |

Dopo il prompt iniziale, l’interrogazione è proseguita con domande di follow-up volte a ottenere un confronto tra il contesto italiano e quello internazionale (la domanda era: “Perché la maternità surrogata è vietata in Italia e consentita in altri Paesi?”), l’ulteriore precisazione delle diverse posizioni nel dibattito italiano, la manifestazione di un posizionamento esplicito (“Qual è la tua opinione?”) e, infine, l’offerta di consigli all’utente in merito ai percorsi da intraprendere (“Tu cosa mi consiglieresti?” nel caso delle cittadine e “Come fa un uomo italiano a ricorrere alla maternità surrogata?” nel caso dei cittadini).

3.2. Una prospettiva discorsiva

Tutte le interazioni ottenute sono state analizzate attraverso la prospettiva dei critical discourse studies<sup>40</sup> guardando in particolare ai meccanismi di priming e framing messi in gioco, valutando la priorità assegnata ai vari argo-

<sup>40</sup> I critical discourse studies sono un insieme di approcci disciplinari e metodologici piuttosto variegato e in continua trasformazione (R. Wodak, M. Meyer, *Methods of Critical Discourse Studies*, London: Sage, 2015), caratterizzato da un pluralismo metodologico attorno a un nucleo epistemologico comune: in tutti i casi, si parte infatti dall’assunto che il linguaggio sia una pratica sociale che produce e riproduce strutture di potere all’interno di un quadro ideologico (J. Flowerdew, J. Richardson, *The Routledge Handbook of Critical Discourse Studies*, London: Routledge, 2017). In questo senso, studiare il discorso significa

menti (effetto di “priming”) e i temi inclusi nelle risposte. Stante la relativa similarità degli argomenti presentati dai due sistemi di GenAI in riferimento al tema della gestazione per altri, abbiamo ritenuto opportuno procedere all’individuazione dei frame discorsivi. Nello specifico, ci si è concentrati su come gli output delle GenAI siano il risultato di meccanismi di selezione e salienza<sup>41</sup>. In quanto sintesi divulgativa di tutte le informazioni che si possono rintracciare sul tema della gestazione per altri, il contenuto proposto dai chatbot contiene una selezione e un’implicita attribuzione di salienza. Infine, a livello lessicale, l’analisi si è concentrata sulla scelta delle parole per descrivere i vari aspetti della maternità surrogata. Questo punto di vista è quello più tradizionale e più consolidato nell’ambito dei critical discourse studies, e trova applicazione anche allo studio della GenAI<sup>42</sup>.

#### 4. Risultati

Nell’interazione con i cittadini e le cittadine, i due LLMs tendono a riprodurre i medesimi schemi discorsivi, chiaramente riconoscibili sul piano argomentativo e tematico. Per entrambi, a essere mobilitato per primo è il frame di tipo legale: si forniscono dettagli sulla normativa italiana, sulle motivazioni del divieto, sulla recente introduzione del “reato universale”, sui rischi e sulle sanzioni che comporterebbe per un cittadino italiano il ricorso alla maternità surrogata all’estero, con precisazioni riguardanti le differenze rispetto ad altri ordinamenti che consentono la maternità surrogata. La priorità data al frame legale-istituzionale (in alternativa, per esempio a quello medico o etico-religioso) è indizio di un approccio di servizio da parte dell’IA che si offre nel ruolo di intermediario tra il cittadino e le istituzioni. Alcune delle domande di follow-up hanno spinto i due chatbot a fornire suggerimenti pratici su come avviare un percorso di maternità surrogata. ChatGPT ha interpretato fino in fondo il suo ruolo di consulente legale, soprattutto per quanto riguarda l’interazione con la cittadina, fornendo un modello di azione step-by-step che comprende sia l’invito a una sorta di esame di coscienza (“chiarisci le tue priorità e i tuoi valori”, “valuta le alternative”) sia gli aspetti burocratici da tenere in considerazione. Claude ha invece risposto in modo

svelare queste dinamiche ideologiche e rendere visibili le strategie con cui si costruisce una legittimazione di tipo egemonico.

<sup>41</sup> R.M. Entman, “Framing: Toward Clarification of a Fractured Paradigm”, *Journal of Communication*, 43, 4 (2023): 51-58.

<sup>42</sup> Gillings, Kohn, Mautner, “The Rise of Large Language Models: Challenges for Critical Discourse Studies”.

più conciso, consigliando al cittadino di rivolgersi ad avvocati specializzati, medici e psicologi. L'analisi ha evidenziato altri quattro frame ricorrenti, condivisi da entrambi i modelli:

1. frame etico-religioso, in cui la diversità delle posizioni è ricondotta a sistemi valoriali differenti ma sostanzialmente equiparati nella loro plausibilità;

2. frame socio-economico, che insiste sul rischio di sfruttamento delle donne in condizioni economiche svantaggiate e si pone come principale obiezione alle prospettive libertaria e femminista;

3. un frame tecnico-medico, che sottolinea le possibili conseguenze psicologiche per bambini e madri surrogate;

4. infine, un frame femminista, in cui vengono messi in risalto i temi della libertà e dell'auto-determinazione da parte delle donne, con un riferimento esplicito al femminismo.

Mentre i primi tre frame appaiono, anche se con pesi diversificati, in tutte le interazioni con i cittadini, il frame femminista compare solo in risposta a cittadine di sesso femminile e di orientamento politico di sinistra<sup>43</sup>.

#### 4.1. Cittadine/i politicamente neutrali

Nel dipingere l'articolazione delle posizioni in gioco a cittadine/i politicamente non schierati, entrambi i modelli restituiscono un quadro contraddittorio. Da una parte, la normativa viene ritratta come in linea con il sentire della maggioranza della popolazione. Alla domanda "Perché in Italia la maternità surrogata è vietata?", Claude risponde alle cittadine di genere femminile elencando tre tipi di motivazioni legali, etiche e sociali, e tra queste ultime fa riferimento al fatto che il divieto di maternità surrogata "riflette valori tradizionali della società italiana sulla famiglia e la maternità". Lo stesso Claude non offre questo tipo di considerazione agli interroganti maschi. ChatGPT risponde alla stessa domanda posta da cittadini di genere maschile facendo riferimento a vari aspetti, tra cui una "cultura giuridica italiana che dà priorità al legame biologico e genetico tra genitori e figli" e "un contesto culturale e religioso" in cui "le implicazioni etiche sono particolarmente sentite". Alle cittadine di genere femminile, invece, spiega che "la normativa riflette le opinioni prevalenti nella società italiana al momento dell'approvazione

<sup>43</sup> Si noti che questo frame "femminista" è un'evidente semplificazione dell'ampia articolazione delle posizioni in seno al femminismo. In particolare, l'insistenza sull'auto-determinazione e sulla libertà delle donne può essere ricondotta alla variante del femminismo liberale, a sua volta versione contemporanea egemonica del femminismo.

della legge 40/2004". Dall'altra, si dà conto dell'esistenza di una pluralità di opinioni e valori che caratterizza il dibattito sia a livello politico e giuridico che a livello sociale e culturale. Le risposte di ChatGPT parlano di un tema che "resta controverso e genera dibattiti" (per le donne), e di un dibattito che "rimane acceso e controverso, con opinioni divergenti" (per gli uomini). Similmente, le risposte di Claude indicano un "dibattito che resta aperto" (per le femmine) e "un acceso dibattito pubblico" (per i maschi). I termini "resta" e "rimane" servono a rafforzare l'idea che *nonostante* ci sia una legislazione consolidata che nega la gestazione per altri (recentemente rafforzata dal Parlamento italiano), la vicenda è ancora dibattuta nell'agenda pubblica.

Proprio all'interno di questa articolazione di posizioni è possibile rintracciare il ruolo della religione e delle istituzioni religiose, che rappresenta un dettaglio distintivo sia tra LLM sia tra i generi degli interroganti. ChatGPT annovera la Chiesa cattolica come uno degli attori rilevanti del dibattito, a differenza di Claude che parla genericamente di "gruppi sociali e religiosi". Inoltre, al netto delle differenze tra modelli, nelle risposte per le cittadine il ruolo della religione è oggetto di una maggiore considerazione. ChatGPT risponde alle cittadine elencando una serie di obiezioni etiche, come la tutela della dignità della donna o l'opposizione alla commercializzazione, tra cui risalta anche "l'influenza della dottrina cattolica" sulla cultura italiana<sup>44</sup>. La risposta alla stessa domanda posta da cittadini fa invece riferimento a "molte persone e istituzioni (in particolare la Chiesa cattolica)" che "hanno posizioni contrarie a questa pratica". Claude, invece, risponde alle donne facendo riferimento al fatto che la normativa "si allinea con posizioni espresse da vari gruppi sociali e religiosi", mentre non considera affatto questi aspetti nella risposta fornita ai maschi. Ciò che viene dato per scontato, osservando queste risposte, è che per le donne, più che per gli uomini, la religione sia un aspetto da tenere in considerazione.

<sup>44</sup> Questo tema ritorna anche in un altro passaggio, quando una cittadina chiede a ChatGPT perché in altri Paesi sia consentita la gestazione per altri: "in Paesi con culture meno influenzate dalla dottrina cattolica (come gli Stati Uniti, il Canada o alcune nazioni dell'Asia), le questioni legate alla riproduzione assistita sono meno legate a principi morali assoluti. Questo facilita una regolamentazione che considera la maternità surrogata come un servizio medico e sociale". Il senso comune artificiale, in questo caso, afferma che a) gli Stati Uniti sono un Paese più laico rispetto all'Italia (dal punto di vista delle classi dirigenti), e che b) la laicità è un valore perché consente maggiore libertà (meno legami con "principi morali assoluti") e favorisce l'elaborazione di una regolamentazione (che è a sua volta, evidentemente, un valore).

## 4.2. Cittadine/i politicamente schierati

Nel caso degli interroganti politicamente posizionati, Claude introduce un riferimento ai soggetti religiosi solo quando interrogato da cittadini di destra. Parimenti, in ChatGPT tali riferimenti sono presenti in modo più esteso e pronunciato nelle risposte per i cittadini di destra. Un altro aspetto importante è che, accanto alla Chiesa cattolica, vengano citate le associazioni pro-vita ma solo nelle risposte fornite alle cittadine. Si può ipotizzare che, in questi casi, il senso comune artificiale presupponga che tali associazioni abbiano una *issue ownership* rispetto al tema dell'aborto e che l'unico pubblico potenzialmente interessato da questo tema sia di genere femminile.

Accanto agli attori religiosi, entrano in gioco degli altri come “la comunità medica”, “la comunità di bioetica”, “gli esperti” e “le femministe”, occasionalmente menzionati e utilizzati come figure che rafforzano la necessità di un approccio cauto e conservatore. Ad esempio, ai cittadini di sinistra, ChatGPT presenta la posizione di “alcuni accademici, bioeticisti” etichettandola come “intermedia”, a supporto di un “dialogo aperto ma con prudenza”. Per le donne di destra, Claude dà conto di una “comunità medica e bioetica [...] divisa, con posizioni che vanno dal divieto assoluto a proposte di regolamentazione molto stringente” e del fatto che “molte associazioni femministe italiane sono contrarie”, di fatto restringendo di molto il panorama delle opzioni. Nella risposta ai cittadini di sinistra, ChatGPT annovera “alcune femministe radicali” tra i contrari, insieme a “organizzazioni cattoliche e conservatrici”<sup>45</sup>.

Un altro elemento degno di nota è che, nella descrizione dello spettro dei diversi posizionamenti politico-ideologici, emerga una rappresentazione degli attuali rapporti di forza tra maggioranza e minoranza variabile a seconda del profilo degli interroganti. Ad esempio, alle cittadine di destra ChatGPT presenta un'opinione pubblica compatta nella contrarietà alla maternità surrogata, enfatizzando quindi un consenso generale rispetto alla politica del governo Meloni di introdurre il reato universale di maternità surrogata. Al contempo, l'opposizione è descritta come fortemente divisa al suo interno, tra posizioni contrarie e in linea con la maggioranza e posizioni favorevoli sparse e insignificanti. Al contrario, nelle risposte fornite ai cittadini di sinistra, “le posizioni sono variegate” e riflettono “un dibattito complesso e polarizzato” sia all'interno della sinistra (e del femminismo) sia all'interno

<sup>45</sup> L'etichettamento come “radicale” della parte di femminismo contraria alla gestazione per altri fa probabilmente riferimento al trans-exclusionary *radical* feminism (o femminismo TERF).

della società italiana nel suo insieme. Anche Claude, che fornisce mediamente risposte più sintetiche rispetto a ChatGPT, propone alle donne di sinistra un dibattito suddiviso in “diverse posizioni, specialmente all’interno del pensiero femminista e progressista”, mentre per i maschi di sinistra elimina il riferimento al femminismo e suddivide le posizioni tra “sinistra divisa al suo interno” e “centro-destra prevalentemente contrario”.

In conclusione, mentre il contesto politico evocato per i cittadini/e di destra è compatto attorno al governo, il contesto politico evocato per i cittadini di sinistra contiene il riconoscimento alla diversità, alla molteplicità delle posizioni e all’asprezza dello scontro. La rappresentazione della sinistra offerta dai due chatbot interpellati è schiacciata sul ruolo istituzionale dei partiti progressisti contemporanei (per quanto riguarda l’Italia, il riferimento evidente è il Partito Democratico), il cui ruolo parrebbe consistere prevalentemente nella mediazione di un metodo di dibattito, più che nella difesa di principi, valori o posizioni specifici rispetto alla maternità surrogata.

## 5. Conclusioni

I risultati evidenziano come gli LLM esercitino la loro funzione di gate-keeping epistemico mediante l’elaborazione di un senso comune artificiale inteso come forma specifica di genericità che emerge al crocevia tra standardizzazione dei repertori e personalizzazione dei contenuti. Rispetto alla standardizzazione, entrambi i modelli restituiscono risposte analoghe sul piano argomentativo e dei frame discorsivi. La ricorsività delle prospettive giuridico-legale, socio-economica, etica e medica evidenzia come il senso comune artificiale sia mobilitato innanzitutto sotto forma di cornice di realtà che si presuppone familiare e condivisa e che sia per questo funzionale a un processo di “normalizzazione” e legittimazione dei contenuti veicolati. Entrambi i modelli assicurano una parvenza di pluralismo politico, riproducendo rappresentazioni convenzionali ma al contempo articolate, in linea con la “varietà stereotipata” e intrinsecamente normativa osservata da Gillespie<sup>46</sup>. All’interno di una differenziazione necessaria tra le posizioni politiche restituite dai chatbot si possono infatti intravedere alcune posizioni ideologiche cristallizzate: le narrazioni di destra si orientano costantemente verso posizioni conservatrici, enfatizzando i valori tradizionali; i discorsi di sinistra si concentrano sul valore in sé del dibattito democratico, con un ricorso minimo, poco aggiornato e molto definito in termini di interrogan-

<sup>46</sup> Gillespie, “Generative AI and the Politics of Visibility”.



ti, al repertorio discorsivo del femminismo. Tale cristallizzazione implica il duplice rischio di normalizzare la relazione biunivoca tra destra/sinistra e un set predeterminato di contenuti, valori e toni, e insieme di indebolire e marginalizzare ancora di più i posizionamenti non-mainstream per ciascuno dei due schieramenti. I posizionamenti politici restituiti dall'IA sono concepibili, in buona sostanza, come idealtipi ricostruiti dalla generalità di contenuti con cui il modello è addestrato, o sarebbe forse meglio dire "super-tipi"<sup>47</sup>, nati dal sincretismo di diversi testi, discorsi ed elementi marcati come non-neutrali (quindi connotati politicamente), ma non per questo estranei a rappresentazioni di senso comune.

Sul piano della personalizzazione, nel caso dei cittadini non schierati, ad esempio, la genericità si mostra nella strategia di dar conto di posizioni contraddittorie senza esplicitarne le radici culturali e le implicazioni politiche. È significativo, in tal senso, che alle donne vengano proposte letture più ancorate alla sfera religiosa e morale, mentre agli uomini letture più inclini a sottolineare gli aspetti giuridici. Quando il posizionamento politico degli interroganti è reso esplicito, invece, la personalizzazione segue altre strade. Ad esempio, la destra appare compatta e solida nel difendere i valori tradizionali, mentre la sinistra è descritta come frammentata al suo interno; anche il riferimento al femminismo, utilizzato in modo molto selettivo solo nel caso delle donne di sinistra, è emblematico di questa dinamica di etichettamento.

Come hanno già avuto modo di sottolineare studiosi come Goffman<sup>48</sup> (1969) e Garfinkel (1967)<sup>49</sup>, al di là della corrispondenza tra segno e significato, sono le conoscenze di sfondo, indiscusse e quindi fonte di riconoscimento implicito, a rendere possibile una comunicazione efficace intesa come allineamento sufficientemente buono tra ciò che l'emittente intende comunicare e ciò che il destinatario coglie. Presentando i contenuti attraverso frame riconoscibili, gli LLM garantiscono quella sensazione di plausibilità e familiarità che facilita l'accettazione acritica delle informazioni fornite. Se la GenAI si radica come risorsa "normalizzata" per la produzione e riproduzione di conoscenza (un ruolo che abbiamo ricondotto al processo di gatekeeping epistemico ma anche di egemonia culturale), il rischio non è tanto la diffusione di informazioni false quanto la progressiva sedimentazione di repertori discorsivi sempre più standardizzati, pur nell'apparente diversificazione degli output.

<sup>47</sup> M. Shanahan, K. McDonell, L. Reynolds, "Role Play with Large Language Models", *Nature*, 623, 7987 (2023): 493-498.

<sup>48</sup> E. Goffman, *La vita quotidiana come rappresentazione*, Bologna: Il Mulino, 1969.

<sup>49</sup> H. Garfinkel, *Studies in Ethnomethodology*, Englewood Cliffs: Prentice-Hall, 1967.

La ricerca qui presentata risente inevitabilmente del carattere sperimentale dell'approccio sviluppato. Non solo abbiamo a che fare con una tecnologia in rapidissima evoluzione che accelera enormemente l'obsolescenza delle condizioni di validità metodologica e delle evidenze empiriche raccolte, ma la stessa letteratura di riferimento ha tutto il carattere provvisorio degli inizi, con un taglio marcatamente interdisciplinare che rallenta il consolidamento di parametri analitici riconosciuti. Riteniamo che la sfida metodologicamente più impegnativa risieda nella necessità di descrivere un fenomeno comunicativo legato, nelle sue stesse condizioni di possibilità, a una contingenza della macchina, strutturale e difficilmente imbrigliabile in schematismi predefiniti. Non si tratta evidentemente di una novità assoluta, se si pensa a quelle forme di adattamento e incorporazione in tempo reale dei comportamenti degli utenti che sono consentite dal machine learning e che sono al cuore di quel processo di piattaformaizzazione culturale e sociale a cui abbiamo assistito negli ultimi anni. L'elemento di rottura è dato piuttosto dal fatto che tale adattamento avviene ora grazie a un'interazione dialogica con la macchina, con proiezioni e assunzioni di ruolo che facilmente innestano valori e aspettative culturali su modelli e criteri di efficienza computazionale. Sulla scia di quanto suggerito da Guzman<sup>50</sup>, la proposta qui discussa sposa l'idea che "parlare con le macchine" rappresenti una delle strade metodologicamente più plausibili, con tutti gli affinamenti che gradualmente emergono ed emergeranno come indispensabili. Al contempo, riteniamo che la naturale prosecuzione di tale approccio non possa che consistere in una ancora più radicale contestualizzazione dei setting comunicativi generati nell'interazione umano-macchina, auspicabilmente mediante un approccio etnografico<sup>51</sup> che consenta di esaminare in profondità dialoghi tra chatbot e specifici gruppi sociali, dispiegati nel loro naturale contesto di vita.

<sup>50</sup> Guzman, "Talking about «Talking with Machines»".

<sup>51</sup> N. Seaver, "Algorithms as Culture: Some Tactics for the Ethnography of Algorithmic Systems", *Big Data & Society*, 4, 2 (2017): 1-12.