



Finanziato
dall'Unione europea
NextGenerationEU



Ministero
dell'Università
e della Ricerca



Italiadomani
PIANO NAZIONALE
DI RIPRESA E RESILIENZA

Inclusione ed elaborazione del linguaggio naturale nell'era dell'intelligenza artificiale generativa

AA. VV.



Politecnico
di Torino

E  MIMIC



TOR VERGATA
UNIVERSITÀ DEGLI STUDI DI ROMA

La pubblicazione è finanziata su fondi del progetto PRIN 2022 PNRR
"Empowering Multilingual Inclusive Communication"
Finanziato dall'Unione Europea NextGenerationEU
PNRR - Missione 4 Istruzione e Ricerca
Componente 2 dalla Ricerca d'Impresa
Investimento 1.1., codice 2022WEFCFP - CUP J53D23007230006



Politecnico
di Torino

E  MIMIC



ALMA MATER STUDIORUM
UNIVERSITÀ DI BOLOGNA



TOR VERGATA
UNIVERSITÀ DEGLI STUDI DI ROMA



Quest'opera è distribuita con
Licenza Creative Commons - Attribuzione - Non opere derivate.
Condividi allo stesso modo 4.0 Internazionale.
Copyright © 2025



Ledizioni LediPublishing
Via Antonio Boselli, 10
20136 Milano - Italia
www.ledizioni.it
info@ledizioni.it

ISBN ebook: 9791256004133

Graphics and page layout
Silvio Ortolani, SISHO - fotografia & archivi

Inclusione ed elaborazione del linguaggio naturale nell'era dell'intelligenza artificiale generativa



Inclusione ed elaborazione del linguaggio naturale nell'era dell'intelligenza artificiale generativa

- 7 **Inclusione ed elaborazione del linguaggio naturale nell'era dell'intelligenza artificiale generativa**
Rachele Raus

I. Inclusione, elaborazione del linguaggio naturale e intelligenza artificiale

- 19 **The Hidden Voice of Artificial Intelligence: Navigating the Unseen Barriers to Inclusion in Scientific Discourse**
Genoveva Vargas-Solar
- 41 **Key-passages and artificial intelligence. How can we use hidden layers to explore new linguistic objects?**
Laurent Vanni, Damon Mayaffre, Hadi Mahmoudi
- 59 **L'intelligenza artificiale generativa al servizio della parità di genere: uno studio esplorativo sugli annunci di lavoro della Confederazione svizzera**
Anna-Maria De Cesare
- 85 **An Inclusive Language Proofreader: Enhancing Writing with Equity and Diversity Principles Through AI Algorithms**
Matteo Berta, Tania Cerquitelli

II. Il progetto E-MIMIC

- 109 **Esprimersi. Da un punto di vista linguistico, logico e simbolico**
Simone Arcari
- 129 **Per un linguaggio inclusivo e accessibile nella Pubblica Amministrazione**
Virginia Laconi
- 149 **Comunicación inclusiva en España y *Deep Learning*: adaptación metodológica del proyecto E-MIMIC al lenguaje de la administración española**
Natalia Peñín Fernández
- 165 **Communication inclusive et apprentissage profond : adaptation méthodologique du projet E-MIMIC à la langue de l'administration française**
Michela Tonti, Martina Ailén García

III. Riflessioni a *latere*

- 183 **Il linguaggio discriminatorio tra etica ed etimologia**
Danio Maldussi
- 193 **Chi ha dato i dati all'intelligenza artificiale?**
Una riflessione al tempo dei *Large Language Models*
Francesca Dragotto
- 201 **"The limits of LLMs are the limits of the world":**
considerations on the role of linguistics in enhancing
AI-generated language quality
Valeria Zotti
- 211 **Le langage inclusif : un savoir-faire traductologique.**
Pour une réflexion sur l'inclusion à l'aune
de l'intelligence artificielle
Ilaria Cennamo
- 217 **Intelligenze artificiali ma *bias* reali:**
***Large Language Models*, dati e possibili mitigazioni**
Chiara Russo

Introduzione

Inclusione ed elaborazione del linguaggio naturale
nell'era dell'intelligenza artificiale generativa

Inclusione ed elaborazione del linguaggio naturale nell'era dell'intelligenza artificiale generativa

Rachele Raus, rachele.raus@unibo.it
Università di Bologna

Il presente volume è frutto di una riflessione collettiva di autori e autrici di diverse provenienze universitarie che hanno voluto confrontarsi sul tema dell'inclusione, dell'elaborazione del linguaggio naturale e dell'intelligenza artificiale generativa nell'ambito dei lavori di un progetto di rilevante interesse nazionale (PRIN2022)¹, finanziato nel quadro del Piano italiano di Ripresa e Resilienza (PNRR) e coordinato dal Politecnico di Torino in partenariato con l'Università di Bologna e l'Università di Roma Tor Vergata.

Nell'era dell'intelligenza artificiale (IA), infatti, temi come l'inclusione e il linguaggio diventano ancora più significativi a causa dell'impatto linguistico, sociale, politico e non solo che i nuovi dispositivi supportati da tale tipo di tecnologia possono implicare rispetto a una società equa e rispettosa del multilinguismo e di ogni differenza.

Il testo si articola in tre sezioni: nella prima sono stati raccolti dei contributi che aprono delle piste di ricerca sulla base di alcune delle sfide filosofiche, linguistiche, informatiche e/o sociali poste dall'IA generativa di ultima generazione. Questa sezione si chiude con la presentazione del suddetto progetto PRIN, il progetto *Empowering Multilingual Inclusive Communication* (E-MIMIC) che intende riflettere su come l'IA possa salvaguardare il multilinguismo, proponendo modelli alternativi (cfr. anche Crosthwaite, Baisa, 2023) e fornendo l'esempio della creazione di un dispositivo basato su reti neurali, *Inclusively*, che riformula i testi amministrativi in chiave inclusiva. Tale saggio introduce la seconda parte del volume, dove sono presentati i lavori dell'*équipe* di E-MIMIC, e di persone che dall'esterno collaborano con essa, per le lingue prese in considerazione dal progetto, ovvero l'italiano d'Italia (cfr. Simone Arcari, Virginia Laconi), il francese di Francia (cfr. Michela Tonti e Marti-

¹ Si tratta del progetto *Empowering Multilingual Inclusive Communication* (E-MIMIC), finanziato dall'Unione europea nell'ambito del programma *NextGenerationEU* — PNRR (Missione 4 Istruzione e ricerca — Componente 2 Dalla ricerca all'impresa — Investimento 1.1, codice 2022WEFCFP - CUP J53D23007230006).

na Ailén García) e lo spagnolo di Spagna (Natalia Peñín Fernández). Nella terza parte, infine, sono stati inclusi dei saggi che intendono inaugurare dibattiti e riflessioni di tipo etico, formativo e linguistico sull'IA e sulle forme d'inclusione non solo di genere ma anche linguistico e di comunicazione inclusiva più generalmente intesa².

1. L'IA tra nuove forme di conoscenza, linguaggi e immaginari

Il saggio di Genoveva Vargas-Solar, che apre la prima sezione del volume, illustra come la generazione di testi prodotta dall'IA vada considerata in maniera molto più estesa, ovvero non solo in relazione alla questione, già nota nella letteratura di settore, della diffusione dei *bias* presenti nei dati di addestramento dei dispositivi (Bartoletti, 2020), ma anche e soprattutto per la diffusione di una conoscenza e la creazione d'immaginari "colonizzati" (nel senso dell'occidentalizzazione del sapere).

Questa riflessione sollecita questioni, da un lato relativamente al modo in cui i dispositivi d'IA generativa contribuiranno alla conoscenza comune e al forgiare appunto l'immaginario, dall'altro rispetto al fatto che questo tipo di conoscenza richiede strumenti interpretativi diversi che sottolineino e diano visibilità a nuovi oggetti d'analisi distinguendoli dal sapere "naturale".

In merito alla prima questione, ricordiamo che la ricercatrice brasiliana Eni Puccinelli Orlandi (2010) ha parlato di "memoria metallica" in relazione a quelle "memorie circolari" inaugurate dai media e dalle nuove tecnologie informatiche emergenti che, decontestualizzando le frasi, le riducono a meri dati³ da accumulare in modo "orizzontale", eliminando così la "storia" della loro produzione e il loro contesto discorsivo. Quella che viene prodotta quindi è una

² Cfr. al riguardo la guida sulla *Comunicazione inclusiva* del Segretariato generale del Consiglio dell'Unione europea. <https://www.consilium.europa.eu/it/documents-publications/publications/inclusive-comm-gsc/>

³ In merito alla distinzione tra i testi raccolti in *corpora*, che richiedono l'interpretazione umana, rispetto ai dati trattati a livello informatico, si veda il saggio di Laurent Vanni Damon Mayaffre e Hadi Mahmoudi in questo volume, nonché François Rastier (2021). « Data vs corpora », In D. Mayaffre, L. Vanni (a cura di), *L'intelligence artificielle des textes. Des algorithmes à l'interprétation*, Paris : Honoré Champion, 203-246.

conoscenza “a specchio”, che diffonde estratti dei discorsi prodotti da un’umanità assente da questo processo. Oltre a diffondere i *bias*, tale reiterazione crea un sapere superficiale, tramite narrazioni tautologiche che producono nuovi immaginari collettivi, intendendo questi ultimi come rappresentazioni mentali evocate dall’immaginazione sulla base di discorsi o immagini (Raus, 2017: VII). In tal senso, occorre comprendere quanto queste nuove forme di conoscenza, che contribuiscono a forgiare determinati immaginari, siano una questione di fondamentale importanza, se consideriamo che, nell’ottica della fisica quantistica, «ogni pensiero è (...) informazione che condiziona gli eventi oltre lo spazio e il tempo» o che, in altre parole, «ciò che pensiamo è più importante di ciò che facciamo» (Conoscenti 2016: 441)⁴, dato che è proprio il pensiero a creare la realtà.

Vogliamo tornare, quindi, alla seconda questione, alla quale abbiamo precedentemente accennato, ovvero al fatto che, per capire appieno questa nuova forma di conoscenza, occorrono nuovi oggetti di studio da osservare e categorie “artificiali”. D’altronde, diversamente dalla metafora biologica che accomuna l’IA all’umanità per il fatto stesso di condividere il tratto dell’essere “intelligenti”, la differenza tra le due necessita di esser posta da più punti di vista complementari, tra i quali vanno citati almeno i due seguenti:

- un diverso tipo di conoscenza che, come abbiamo detto, presenta caratteristiche e tratti propri specifici;
- un diverso modo di “apprendere”, sul quale riflette, nell’ambito linguistico, Francesca Dragotto nella terza parte del volume.

Rispetto al primo punto, il contributo di Vanni, Mayaffre e Mahmoudi da un lato e quello di Anna-Maria De Cesare nella prima sezione dell’opera, introducono nuovi oggetti di studio e di conoscenza. Nel primo, si tratta di un “osservabile” linguistico-informatico definito con il termine francese “*passage*” che potremmo tradurre in italiano con “estratto”. Si tratta, secondo gli autori, di sequenze di dati estratti dalla macchina da *corpora* omogenei, quali quelli letterari analizzati nel loro contributo, che presentano caratteristiche simili, permettendo ai dispositivi d’IA, in questo caso di

⁴ La traduzione dall’inglese è la nostra.

tipo non generativo, di ottenere *performance* interessanti rispetto all'analisi dei testi e all'individuazione di caratteristiche discorsive degli scrittori o dei politici francesi esaminati. Nel caso di Anna-Maria De Cesare, è la nozione di "morfema artificiale" che viene introdotta come nuovo oggetto di studio per caratterizzare dei fenomeni di generazione di testo da parte dell'IA quando, nella traduzione automatica, crea forme inventate dal punto di vista morfologico, quali quelle che marcherebbero il femminile negli esempi riportati nel suo saggio.

Questi spunti di ricerca impongono due evidenze che ribadiamo con forza: la necessità di porre l'IA come un qualcosa di fondamentalmente diverso dall'umanità⁵, superando la sovrapposizione concettuale dovuta alla metafora alla base stessa del sintagma, e l'ovvia conseguenza di dover disporre di categorie ermeneutiche originali per poter comprendere le nuove forme di conoscenza inaugurate dai dispositivi generativi. In tal senso, sarebbe peraltro opportuno evitare di essenzializzare l'IA, preferendo piuttosto parlare di "intelligenze artificiali" diverse, basate su architetture e modelli distinti. Il saggio di Berta e Cerquitelli, che chiude la prima sezione e inaugura la seconda dedicata al progetto E-MIMIC, mostra, ad esempio, che tra le IA possibili, i modelli di apprendimento supervisionato dall'umano potrebbero evitare la diffusione di *bias* come avviene invece in quelli non supervisionati.

2. IA al servizio del linguaggio inclusivo: sì, ma di quale inclusione parliamo?

L'elaborazione del linguaggio naturale pone al centro la questione dell'automazione prodotta dai dispositivi d'IA nell'ambito del linguaggio (generazione, traduzione, analisi del testo, ecc.) anche relativamente al tema, molto discusso recentemente, dell'inclusione. D'altronde, il progetto E-MIMIC, presentato alla fine della prima sezione del volume e al quale è dedicata tutta la seconda, intende produrre un dispositivo capace di segnalare all'utenza delle parti del testo amministrativo che sarebbero non inclusive e di proporre delle riformulazioni inclusive. Nella sezione dedicata *ad hoc*

⁵ Su questo cfr. anche i recenti lavori di Federico Faggin e in particolare il volume del 2022 (*L'irriducibile*. Milano: Mondadori).

al progetto emerge quanto occorra anzitutto chiarire i termini della questione già solo nel porre la dicotomia “segmento inclusivo vs non inclusivo”, come fa Simone Arcari nel suo saggio. Una volta rifondata la questione, l’inclusione è stata intesa anzitutto dal punto di vista linguistico in relazione alle donne. In tal senso, il saggio di Anna-Maria De Cesare, nella prima parte del volume, introduce la parità linguistica nella scrittura delle offerte di lavoro in Svizzera, come d’altronde pronato dalle politiche linguistiche svizzere che vengono citate nel suo testo. Una riflessione di tipo etico rispetto proprio all’inclusione delle donne, prima ancora di quella più ampia di genere, è quella proposta da Danio Maldussi nella terza sezione del volume, che ci invita a una riflessione e apre un vero e proprio dibattito partendo dall’etimologia delle parole che usiamo in relazione a queste tematiche. Di nuovo, quindi, come con Arcari, si cerca di rifondare i termini della questione. Nella terza parte del volume, Valeria Zotti e Chiara Russo amplieranno la discussione, parlando del modo in cui i dispositivi d’IA possono contribuire o meno alla diffusione degli stereotipi di genere.

Gli altri saggi della seconda sezione, invece, pongono l’inclusione nell’ottica non solo del genere ma anche degli stereotipi concernenti categorie diverse come le persone disabili, ipovedenti ecc., più in linea quindi con le raccomandazioni dell’Unione europea sulla comunicazione inclusiva che non si limita alle donne o al genere. Nel saggio di Virginia Laconi, oltre all’inclusione viene presa in considerazione anche la nozione di accessibilità semantica, che accosterebbe maggiormente la cittadinanza ai testi della pubblica amministrazione. Nel saggio di Michela Tonti e Martina Ailén García e in quello di Natalia Peñín Fernández vengono presentati i criteri per la riformulazione inclusiva rispettivamente per il francese di Francia e lo spagnolo di Spagna, altre lingue prese in considerazione dal progetto E-MIMIC, nonché le politiche linguistiche diverse delle pubbliche amministrazioni di questi paesi che impattano sulle riformulazioni proposte e sui modi d’intendere l’inclusione, tema che resta spesso controverso e dibattuto anche in sede legislativa. La presenza di altre lingue del progetto ci porta a un’altra evidenza quando si parla d’inclusione: l’imperativo di avere “rispetto”, per riprendere il termine usato da Maldussi, per le varie lingue esistenti.

3. Includere anzitutto il multilinguismo: una vera sfida per le IA

Il progetto E-MIMIC intende prendere in considerazione la diatopia linguistica di tre lingue, — l'italiano, il francese e lo spagnolo, — due delle quali presentano molte varietà geolocalizzate (il francese e lo spagnolo). L'addestramento dei dispositivi su dati indifferenziati (La Quatra, 2023: 75), infatti, non consente di capire da quale varietà linguistica essi stiano imparando (di un paese specifico, europea, internazionale, veicolare, ecc.) e quale variante diatopica della lingua stiano contribuendo, a loro volta, a diffondere generando del testo. La maggior parte dei traduttori automatici, con rarissime e limitate eccezioni, permette di selezionare un generico "francese" a dispetto dei francesi esistenti, e lo stesso avviene anche per lo spagnolo o per molte altre lingue diffuse. Come annuncia Valeria Zotti nel saggio nella terza parte del volume, è spesso l'inglese la lingua dominante nella traduzione automatica; Anna-Maria De Cesare mostra come l'utilizzo di *Large Language Models* nei chatbot che supportano output in inglese o in cinese siano meno performanti quando utilizzati per le altre lingue. Nel n°237 della rivista *Langages* (Raus, Tonti, 2025) emerge come sia proprio l'assenza di *corpora* realmente multilingui e/o rispettosi della variazione diatopica delle lingue a mettere in discussione i dispositivi d'IA, che contribuiscono a diffondere conoscenza tramite concetti spesso calcati sull'inglese internazionale, contribuendo così a quella che in letteratura è stata definita l'"omogeneizzazione" (Bommasani *et al.*, 2021) o la "standardizzazione" (Larsonneur, 2021) linguistica.

La necessità quindi di rispettare le varianti linguistiche e di utilizzare risorse nazionali o multilingui è un imperativo per realizzare anzitutto l'inclusione linguistica. In tal senso, Chiara Russo, nella terza parte del volume, analizza il caso del modello di LLM "Italia", creato in Italia e addestrato su fonti nazionali italiane. Lo stesso dispositivo *Inclusively* del progetto E-MIMIC è stato specializzato su *corpora* nazionali, per il francese, e dispone di un modello addestrato su *corpora* italiani per la versione italiana (La Quatra, Cagliero, 2023).

Tuttavia, queste sono ancora iniziative circoscritte, sebbene restino molto significative. Nell'ottica di un'inclusione linguistica e non solo, sarebbe utile concertare globalmente una "politica dei dati" che sia realmente etica e che possa affiancare le iniziative legislative,

anch'esse rare e spesso tardive a causa delle lungaggini dei tempi di redazione e approvazione delle leggi, sia a livello nazionale sia a livello sovranazionale. Come ben ribadito da Francesca Dragotto nella terza parte del volume, tale politica resta fondamentale.

Anche la formazione, però, è essenziale. In tal senso, occorre, come annuncia Ilaria Cennamo sempre nella terza parte della presente opera, una “competenza traduttiva” all'inclusione, considerando quindi la riformulazione inclusiva come una traduzione di tipo intralinguistico (cfr. Raus, 2023).

Diventa, infine, necessario creare una consapevolezza critica nei confronti dei dispositivi d'IA, come già sollevato da più parti in ambito traduttivo (cfr. Raus, Mattioda, 2024 : 230). In questa prospettiva, ci sembra ancor più significativo ribadire l'esigenza di introdurre nuovi oggetti di studio per cogliere appieno la portata delle nuove forme di sapere, di linguaggi e d'immaginari prodotti dalle IA, che caratterizzeranno questi e gli anni a venire.

Bibliografia

Bartoletti Ivana (2020). *An artificial revolution: on power, politics and AI*. Black Spot Books.

Bommasani Rishi *et alii* (2021). “On the Opportunities and Risks of Foundation Models”, Report page with citation guidelines, Center for Research on Foundation Models (CRFM) at the Stanford Institute for Human-Centered Artificial Intelligence (HAI). <https://crfm.stanford.edu/assets/report.pdf>

Conoscenti Michelangelo (2016). “In Transit between Two Wor(l)ds: NATO Military Discourse at a Turning Point”, in S. Campagna, E. Ochse, V. Pulcini, M. Solly (a cura di), *Languaging in and across Communities: New Voices, New Identities*. Berna: Peter Lang, 425-449.

Crosthwaite Peter, Baisa Vit (2023), “Generative AI and the end of corpus-assisted data-driven learning? Not so fast!”, *Applied Corpus Linguistics* 3 (3), <https://www.sciencedirect.com/science/article/pii/S2666799123000266>

La Quatra Moreno (2023). “Technological independence and cultural diversity in European artificial intelligence”, in R. Raus (a cura di), *How Artificial Intelligence Can Further European Multilingualism*, Artificial Intelligence for European Integration – AI4EI | Report – 2023. Milano: Ledizioni, 73-78.

La Quatra Moreno, Cagliero Luca (2023). “BART-IT: An efficient Sequence-to-Sequence model for Italian text summarization”, *Future Internet*, 15 (1), 15.

Larsonneur Claire (2021). “Intelligence artificielle ET/OU diversité linguistique : les paradoxes du traitement automatique des langues”, *Hybrid*, 7. <https://journals.openedition.org/hybrid/650>

Puccinelli Orlandi Eni (2010). “A contrapelo: Incursão teórica na tecnologia: Discurso eletrônico, escola, cidade”, *RUA*. <https://periodicos.sbu.unicamp.br/ojs/index.php/rua/article/view/8638816/6422>

Raus Rachele (2017). “Préface”, in R. Raus, G. Cappelli, C. Flinz (a cura di). *Le guide touristique : lieu de rencontre entre lexique et images du patrimoine culturel*, Vol. 2, VII-XVI. https://media.fupress.com/files/pdf/24/3360/3360_11640

Raus Rachele (2023). “Deep learning e traduzione intralinguistica: riformulare i testi della pubblica amministrazione in modo inclusivo”, *Studi Italiani di Linguistica Applicata*, 3/2023, 552-569.

Raus Rachele, Maria Margherita Mattioda (2024). “L’intelligence artificielle en salle de classe : la perception des étudiantes et des étudiants”, in R. Raus,

F. Bisiani, M. M. Mattioda, M. Tonti (a cura di). *Multilinguismo europeo e IA tra diritto, traduzione e didattica delle lingue*. Milano: Ledizioni, 223-233. <https://www.collane.unito.it/oa/items/show/195#?c=0&m=0&s=0&cv=0>

Raus Rachele, Michela Tonti (a cura di) (2025). Intelligence artificielle, corpus et diversité linguistique : enjeux et perspective. *Langages*, 237.



Inclusione, elaborazione del linguaggio naturale e intelligenza artificiale

The Hidden Voice of Artificial Intelligence: Navigating the Unseen Barriers to Inclusion in Scientific Discourse

Genoveva Vargas-Solar, genoveva.vargas-solar@cnrs.fr
CNRS, INSA Lyon, Université Claude Bernard Lyon 1, LIRIS, UMR5205

*There is no time for despair, no place for self-pity,
no need for silence, no room for fear.
We speak, we write, we do language.
That is how civilizations heal. [...]
Like failure, chaos contains information
that can lead to knowledge even wisdom. Like art.*
Beloved, Toni Morrison

Introduction

Discourse is the hidden voice behind any text, subtly whispering the perspective, cultural, social, and historical context of its author. This implicit voice reveals that no discourse is ever complete; it is shaped by the intricate processes through which humans gather images and populate their imaginary via interpretation. Language serves as the medium through which the imaginary materializes into new texts, perpetuating a cycle of meaning-making. As Gilbert Durand suggests,

the imaginary is the foundation of mental life and an essential dimension of human existence, the power of dreams, the potency of symbols, the generative capacity of images, forming a kind of transcendental fantasy integral to human identity. It is a vital component of how individuals engage with and make sense of the world (Durand, 2021).

This imaginative process forms the bedrock of human creativity and understanding, as individuals navigate and interpret the complexities of their environments. It underscores the inherent subjectivity of all discourse and its dependence on shared symbols and cultural references to communicate meaning effectively (Ricaurte *et al.*).

The implications of a populated imaginary on the construction of scientific discourse are profound, as it influences not only the way knowledge is produced and communicated but also the values and power dynamics embedded within it. Scientific discourse often aspires to be objective, but it is inevitably shaped by the cultural, his-

torical, and social imaginary of its creators. The images, metaphors, and symbols that populate the imaginary of scientists subtly inform the framing of research questions, the interpretation of results, and the way findings are presented.

The imaginary that drives scientific discourse often reflects the dominant socio-political and cultural narratives, leading to biases in research priorities, methodologies, and applications (*idem*). For example, the overrepresentation of Western, male, and affluent perspectives in scientific institutions and publications influences what is deemed important to study, often overlooking issues pertinent to marginalized communities. Indeed, the discourse employed in scientific communication, pivotal in disseminating observations and knowledge, is inherently influenced by the thought systems of its creators, contrary to the presumed objectivity in the hard sciences. This discourse is often interwoven with biases and potentially exclusionary perspectives shaped by the communicator's perceptions of gender, socio-economic background, race, culture, and physical attributes (Figueiroa, 2023).

Besides, scientific discourse does not exist in isolation; it influences and is influenced by public understanding of science. A populated imaginary affects how scientific ideas are communicated to and received by the public. For instance, the portrayal of artificial intelligence as either a "threat" or a "savior" reflects and shapes societal attitudes toward the technology, impacting policy and funding decisions (*idem*; Laws *et al.*, 2024).

Biases are prevalent in human-generated content and ingrained in Artificial Intelligence (AI) models, reflecting collective prejudices. With the advent of generative text systems, AI is increasingly used in crafting scientific content. However, lacking an inherent understanding of diversity and inclusiveness, these systems often propagate biased narratives under the guise of neutral discourse. This perpetuates a scientific dialogue that is resistant to diversity and inclusion, creating a cyclical reinforcement of exclusionary practices. Thus, it is imperative to reimagine how scientific discourse is structured, integrating safeguards that actively promote diversity, equity, and inclusion (DEI).

To address these implications, it is crucial to diversify the imaginary that informs scientific discourse. This requires recognizing and integrating multiple epistemologies, including non-Western and in-

digenous knowledge systems, which have long been marginalized in favour of dominant Western paradigms. By embracing diverse ways of knowing, science can move beyond its current limitations and gain new insights into global challenges. Developing inclusive language practices is equally essential, as language shapes not only how knowledge is communicated but also whose voices are heard and whose perspectives are validated. This involves reflecting diverse perspectives, cultural contexts, and avoiding exclusionary jargon that alienates underrepresented communities. Encouraging interdisciplinary collaboration is another key strategy, bringing together experts from different fields to challenge and expand the dominant scientific imaginary, ensuring that it reflects the complexity and interconnectedness of real-world phenomena. Finally, embedding ethics, equity, and sustainability as foundational principles in scientific research and communication ensures that scientific endeavours align with broader societal goals, prioritizing collective well-being over profit or narrow utility. By reimagining the imaginary itself, scientific discourse can transcend its historical biases and become a transformative tool for addressing the complexities of the world, fostering innovation that is equitable, inclusive, and reflective of humanity's diverse cultural and intellectual heritage.

This paper addresses the challenges posed by AI in maintaining an unbiased scientific narrative. It explores strategies to cultivate a more inclusive discourse that respects and represents diverse perspectives, ensuring that AI serves as a tool for enhancing, rather than undermining, scientific integrity and societal values. The paper critically examines the construction of the “official” scientific discourse by interrogating the way it shapes a scientific imaginary and exploring its potential for fostering a more inclusive narrative (*cf.* Section 1). In the era of AI, the methods employed — data science pipelines implementing *in silico* data-driven experiments — have effectively transformed into an automated factory of knowledge (*cf.* Section 2). However, this automation raises a fundamental tension: while AI ostensibly aims to democratize knowledge, it simultaneously centralizes it within a single, globalized framework, perpetuating existing inequities. This prompts the pivotal question: can AI truly serve as a vehicle for inclusive and responsible knowledge production? Central to inquiry is the role of language as the primary medium for knowledge dissemination (*cf.*

Section 3). If AI-generated (or influenced) textual content aspires to embody Diversity, Equity, and Inclusion (DEI) values, then the use of language must be approached with utmost care and awareness. Yet, the design and development of AI systems are deeply rooted in Western epistemologies, often reinforcing entrenched power dynamics and exclusionary practices (*cf.* Section 4). Finally, the paper explores fairness, responsibility, and accountability as guiding principles to counteract the extractivist tendencies of AI systems in constructing “official” scientific discourse. By embedding these principles into AI design and practice, the paper advocates for a shift toward a more equitable and sustainable knowledge production paradigm — one that not only recognizes but actively integrates diverse epistemologies and cultural contexts into the global scientific narrative.

1. The scientific “official” speech

“Words are not blown by the wind”, combined creatively, words can create art and realities: some believe that the world emerged by the words expressed by divinity, a price has been put on people's heads due to the words spoken and the attitudes that 'stain' the purity of the texts. Words “colonise imaginary” (P. Ricaurte *et al.*, 2025) and determine the individual and collective agreement about science, beliefs, social contract in different cultures. The question is: *how does this complex mechanics work celebrating, excluding & harming?*

1.1 The Scientific Imaginary

Speech inherently relies on an epistemic foundation, shaped by a series of critical decisions. These include determining the perspective through which an object of study is observed and formulating hypotheses to explore it. Additionally, speech builds upon pre-existing knowledge, such as the provenance of documents and language, which serve as the basis for generating new insights. The choice of words to communicate this knowledge further reflects these foundational decisions, influencing how ideas are constructed, interpreted, and shared within a given context.

AI models are increasingly populated and constructed using speech and narratives embedded in various forms of imagery, encompassing natural, abstract, and artificial visual elements. These images serve not only as data for machine learning algorithms but also as mediums shaping the collective imaginary (Miceli *et al.*, 2024). This imaginary, influenced by such images, risks becoming colonized by predefined notions of what is "natural" or abstract, potentially limiting diversity and reinforcing dominant perspectives. This process underscores the need for critical evaluation of how imaginary and narratives inform AI systems and their broader cultural impact.

The prevailing approach in AI often reduces reality into quantitative representations, using numerical, statistical, and probabilistic models to capture the complexity of phenomena. Within this framework, non-inclusive discourses are treated as biases in the input data, which can be measured and rectified statistically or through AI-driven adjustments. Algorithms can be fine-tuned to protect and promote diversity, equity, and inclusion (DEI) values. However, this methodology raises critical questions: Are DEI values interpreted universally or through a Western-centric lens? Do these adjustments genuinely embrace diverse cultural, socio-economic, and de/post-colonial perspectives, or do they risk imposing reductive and uniform notions of DEI? Moreover, how are multiple, potentially conflicting discourses incorporated to address inequalities in a holistic and equitable manner? (Vargas-Solar *et al.*, 2022). These unresolved tensions highlight the limitations of an exclusively quantitative approach to fostering inclusivity.

The scientific imaginary is deeply rooted in Western values, emphasizing the conquest and control of nature, the environment, and phenomena. Official scientific discourse, described by Pierre Bourdieu's concept of capital, often questions or marginalizes models that deviate from Western norms or originate from non-Western perspectives (Bourdieu, 1992). Pierre Bourdieu discusses the concept of symbolic capital and its role in scientific discourse in *The Rules of Art: Genesis and Structure of the Literary Field* (*Les Règles de l'Art: Genèse et Structure du Champ Littéraire*). He argues that symbolic capital, which encompasses recognition, prestige, and authority, plays a pivotal role in shaping scientific discourse and the hierarchy of the intellectual field. Bourdieu states:

Symbolic capital is nothing other than capital — economic or cultural — when it is perceived and recognized as legitimate. In the scientific field, this capital operates as a form of power that organizes the hierarchy of intellectual positions and determines whose knowledge is valued and legitimized as truth. (Bourdieu, 1992: 141)

This underscores how the accumulation and distribution of symbolic capital dictate what is perceived as credible or authoritative within scientific and intellectual circles. When Western scientific discourse becomes the dominant reference, it implicitly undermines the validity of non-Western scientific perspectives. What happens to traditional practices in medicine, or to methods such as observing the sky to determine the optimal times for planting? How do we address the limitations imposed by binary gender perspectives within scientific frameworks? And most critically, how can alternative sciences be meaningfully integrated into official scientific discourse to enrich the collective scientific imaginary? These questions call for a broader, more inclusive reconsideration of what constitutes valid knowledge.

Western scientific imaginary prioritizes the departmentalization of knowledge, creating a hierarchy among disciplines, where those perceived as "hard" sciences take precedence over "soft" ones. Utilitarianism permeates this structure, favouring disciplines that advance capitalist progress, often at the expense of broader, holistic understandings of knowledge. Consequently, science is frequently aligned with the service of capital rather than the pursuit of knowledge. This alignment reduces complex problems and phenomena to fit the framework of utility, sidelining alternative or expansive approaches (Couldry, Mejías, 2020). The epistemic foundation of this scientific imaginary is inherently centred on the dichotomy of nature versus civilization, where control over nature is paramount (Vandata, 2016).

Designing an alternative perspective where the quantitative approach remains a rule requires a fundamental transformation in the scientific production model. This involves integrating qualitative considerations into numerical and AI-driven methods, emphasizing the conditions under which data are produced, collected, and processed, as well as their provenance and the human factors behind them. It necessitates evaluating the costs — both human and ma-

terial — of producing and processing scientific knowledge, including the implications of resource-intensive activities such as algorithm training, data storage, and result sharing (Miceli *et al.*, 2024). Additionally, acknowledging the environmental and infrastructural resources involved is essential for fostering a more responsible and inclusive scientific practice that aligns with broader societal and ecological considerations.

1.2 What is inclusion in scientific discourse?

The use of language is never neutral and must embrace an "inclusion" strategy that ensures representativity through the careful consideration of statistics, recurrent patterns, and exposure. Language extends beyond words to encompass images, sounds, sensations, and metaphors, each playing a role in shaping narratives and knowledge systems. However, epistemic violence arises when certain content is excluded or deemed invalid. Questions arise: Where are women and non-Western knowledge? Where is alternative knowledge such as activism or indigenous languages? Why are critical terms, often absent? Addressing these gaps requires a personal and collective strategy that infiltrates symbolic systems such as science, art, and religion to create a more inclusive and representative discourse.

Language plays a critical role in shaping discourse, as the vocabulary, choice of words, metaphors, and examples influence how scientific knowledge is communicated and perceived. While science strives for objectivity, it is not immune to the subtle biases embedded in cultural and social subjectivity. For example, terms like "master-slave" used in technical contexts can perpetuate harmful historical connotations, while metaphors such as "conquering nature" reflect anthropocentric and exploitative worldviews. Similarly, the underrepresentation of non-Western languages and perspectives in academic publications creates an implicit bias that excludes diverse ways of knowing. These examples demonstrate how linguistic choices are far from neutral, often guiding the framing and reception of scientific ideas in ways that reinforce societal structures (Aguilar Gil, 2024).

Inclusion in scientific discourse entails creating a space where diverse voices, perspectives, and knowledge systems are represented, acknowledged, and integrated into the broader fabric of scientific in-

quiry and communication. From a colonial perspective, inclusion often focuses on assimilation, where marginalized voices are brought into existing structures, but only insofar as they conform to dominant paradigms and methodologies. This approach can perpetuate epistemic violence by undervaluing non-Western, Indigenous, or alternative knowledge systems and by framing inclusion as an act of generosity rather than a necessary step for a holistic understanding of science.

Addressing these issues requires an inclusive approach to language, ensuring diverse representation and fostering equity in scientific communication. However, current efforts to adopt equitable and inclusive practices often rely on rephrasing strategies. While effective for addressing surface-level issues—such as neutralizing gendered language, modifying adjectives, or restructuring certain patterns—these approaches tackle only part of the problem. Automated rewriting tools can standardize and transform text to adhere to inclusive norms, but discourse extends beyond grammar and vocabulary; it resides in implicit messages and underlying narratives woven within the lines.

Conversely, a decolonial perspective seeks to disrupt these power imbalances by fundamentally rethinking whose knowledge is considered valid and how it is represented (Harding, 2010; Figueiroa, 2023). It emphasizes the need to acknowledge and elevate non-dominant epistemologies, such as Indigenous knowledge, feminist perspectives, and local cultural practices. Decolonial inclusion challenges hierarchical notions of expertise and instead advocates for the co-creation of knowledge where diverse contributions are seen as equally valuable. This approach also questions the language, metaphors, and narratives used in scientific discourse, ensuring they reflect a multiplicity of worldviews rather than a monolithic, often Western-centric perspective.

This deeper challenge requires human-in-the-loop strategies, where AI tools can identify patterns in multimedia content, flagging potential biases or exclusions. These tools should then collaborate with humans who are not only educated in but deeply engaged with DEI values (Vargas-Solar *et al.*, 2022). Together, they can ensure that scientific communication reflects both technical precision and cultural sensitivity, fostering a richer and more inclusive global knowledge collectivity.

2. Artificial intelligence and the factory of knowledge

Artificial intelligence (AI) has emerged as one of the most transformative technologies, redefining how knowledge is produced, disseminated, and consumed. AI's capabilities to process massive datasets, extract patterns, and generate insights have revolutionized various fields, from scientific research and education to healthcare and the arts. However, its role as a "factory of knowledge" raises critical questions about the ethics, inclusivity, and sustainability of this new mode of knowledge production. By exploring the mechanisms through which AI generates knowledge, its implications for society, and the challenges it poses, we can better understand how to harness its potential while addressing its risks.

To understand how knowledge can be automatically produced, it is essential to examine the pipeline used to transform text into knowledge. The process begins with digital content (e.g. text) undergoing analysis to extract information and generate a simplified representation. This representation can take various forms: quantitative (statistical measures), linguistic (identifying entities like locations or points of interest), structural (highlighting grammatical and action-based relations like subjects, verbs, and objects), semantic (capturing concepts and their relationships), or vectorial. Based on the chosen representation, AI models can be employed to extract knowledge — for instance, detecting sexist or hate-filled posts within social network datasets.

However, this seemingly straightforward pipeline involves several critical considerations. First, the conditions under which data are collected must be evaluated, including the selection of sources and obtaining consent from data owners for its use in model training (Couture *et al.*, 2025; Vargas-Solar, 2022; Raus, Tonti, 2025). Second, the data's quality is key, particularly in terms of biases, representativity, and statistical properties, as these factors significantly impact the reliability and fairness of the resulting knowledge. Third, the way data is cleaned, prepared and engineered, to fix biases, missing values, representativity of different observations that they might contain. The fourth critical aspect involves selecting one or multiple models to apply to the data, calibrating them to achieve specific performance objectives, and interpreting the results. This step determines the model's capability to generate new knowledge,

accurately represent patterns, or predict events. The choice of a model and its subsequent calibration must consider the context, objectives, and potential limitations to ensure meaningful and reliable outcomes.

The supposed objectivity of statistical methods and the decision-making processes surrounding data selection, sources, and applied techniques often obscure the underlying social structures influencing their outcomes. Models, by their very nature, are fuelled by collective representations of phenomena embedded in the input data, thereby amplifying and perpetuating these perspectives. While this may seem unproblematic, it becomes concerning when the results are presented as objective, mathematical, or scientific truths. The lack of critical interrogation exacerbates the dominance of a Western-centric vision of knowledge, reinforcing existing power dynamics and the status quo without challenging or broadening the epistemic scope.

At its core, AI is a tool for processing information and producing insights. Algorithms such as machine learning (ML) and natural language processing (NLP) can analyse vast amounts of unstructured data, from textual documents to sensory input, identifying patterns that would be impossible for humans to discern unaided. For example, AI has accelerated scientific discovery by analysing genetic data to identify potential drug targets (Laws *et al.*, 2024), optimizing engineering designs, and simulating environmental processes to predict climate change impacts. AI also plays a pivotal role in transforming raw data into structured knowledge. It does this through the following mechanisms:

1. **Data Processing and Analysis:** AI systems sift through enormous datasets, identifying trends, anomalies, and correlations. In fields such as astronomy, AI analyses terabytes of data collected from telescopes, categorizing celestial bodies and detecting phenomena like exoplanets or supernovae.
2. **Prediction and Modelling:** AI is instrumental in predictive analytics, from weather forecasting to financial market trends. Neural networks trained on historical data can generate models that anticipate future outcomes with remarkable accuracy.
3. **Content Generation:** Generative AI tools, such as OpenAI's GPT models, have demonstrated the ability to create coherent and contextually relevant content, ranging from news articles to sci-

entific papers. These tools contribute to the generation of new knowledge by summarizing existing information or proposing novel hypotheses.

4. **Knowledge Organization:** AI systems enhance the accessibility and utility of knowledge by categorizing and indexing it. For instance, AI-powered search engines and recommendation systems enable users to navigate vast information landscapes efficiently.

While these mechanisms promise to democratise knowledge and accelerate progress, they also introduce new paradigms that require careful scrutiny.

2.1 The Democratization and Centralization of Knowledge

One of the promises of AI is its potential to democratize knowledge. Open-access platforms powered by AI provide educational resources to underserved communities, and AI-driven translation tools break down language barriers, enabling people worldwide to access scientific and cultural content. However, the same technology that democratizes knowledge can also centralize it in unprecedented ways. The infrastructure required to develop and deploy advanced AI systems, such as high-performance computing and vast datasets, is often controlled by a handful of corporations and institutions. This centralization risks creating a new kind of epistemic inequality, where the ability to produce and control knowledge is concentrated in the hands of a few.

The use of AI in the factory of knowledge is not without ethical implications. AI systems are only as unbiased as the data they are trained on. Historical biases embedded within datasets often lead to discriminatory outcomes, perpetuating stereotypes and reinforcing existing inequalities. For example, AI-powered hiring algorithms have been found to disadvantage women and racial minorities, reflecting systemic disparities in historical hiring practices. Addressing bias requires proactive measures such as diverse dataset curation and continuous algorithm auditing to ensure fairness and inclusivity (Couldry, Mejías, 2019).

The "black box" nature of many AI models complicates our ability to understand how decisions are made. This lack of transparency is

especially concerning in high-stakes domains like healthcare, criminal justice, and finance, where the consequences of errors or biases can be profound. Without clear mechanisms to interpret and explain AI-driven decisions, holding developers and implementers accountable becomes a significant challenge, raising urgent calls for interpretable AI solutions.

AI systems frequently reflect a Western-centric worldview, marginalizing non-Western knowledge systems and cultural narratives. The predominance of English-language content in training datasets often excludes indigenous knowledge, local traditions, and diverse cultural perspectives from being recognized or preserved. This epistemic bias risks further homogenizing global knowledge, sidelining valuable insights that do not conform to dominant paradigms.

Generative AI tools, which produce new content by synthesizing existing works, blur the boundaries of authorship and ownership. This raises critical questions about intellectual property rights: who owns the content created by AI, and how should creators whose works inform AI models be compensated? Current legal frameworks often fail to address these nuances, leaving gaps in protecting both human creators and the integrity of AI-generated knowledge (Siles *et al.*, 2023).

Besides authorship and social costs, the factory of knowledge powered by AI comes with significant environmental costs. Training large AI models requires immense computational resources, consuming vast amounts of energy. The carbon footprint of AI systems, particularly in regions dependent on fossil fuels for electricity, raises concerns about their sustainability. Developing ecological harm aware AI technologies and optimizing algorithms for reducing ecological and social harm must become a priority to ensure the long-term viability of AI-driven knowledge production (Couldry, Mejías, 2019).

2.2 Toward Inclusive and Responsible AI Knowledge Systems

To ensure that AI fulfils its potential as a transformative force for good, it is essential to adopt inclusive and responsible approaches to its development and deployment. Establishing robust ethical frameworks is crucial for guiding AI research and applications to ad-

dress issues like bias, transparency, and accountability. These frameworks must prioritize respect for human rights and cultural diversity, ensuring AI systems serve as tools for equity rather than perpetuating inequalities. Diverse representation within AI development teams is equally vital. By reflecting the global population's diversity, these teams can bring varied perspectives to design processes, mitigating biases and creating AI tools that are relevant and accessible to all communities.

To truly democratize AI, efforts to decolonize its development must be prioritized. This involves integrating non-Western knowledge systems and languages into AI training datasets. Collaborative partnerships with local communities are necessary to ensure their perspectives are accurately represented, and their intellectual contributions are acknowledged. Alongside these initiatives, promoting education and awareness is essential. Digital literacy programs should focus on critical thinking, enabling individuals to responsibly engage with AI technologies while understanding their ethical and societal implications. Lastly, sustainability must be at the core of AI development. Innovations such as energy-efficient processors and algorithms, coupled with policies incentivizing the use of renewable energy, can significantly reduce AI's environmental footprint.

All these elements — decolonization, education, inclusion, and sustainability — should converge into a new paradigm: an AI knowledge factory designed for collective, non-extractivist social good. This reimaged AI ecosystem would prioritize equity, foster global collaboration, and aim to address societal challenges without exploiting people or the planet. By aligning technological progress with ethical and ecological responsibility, such a system can create meaningful and lasting benefits for all communities.

AI has established itself as a powerful engine of knowledge production, transforming how we understand and interact with the world. However, as with any transformative technology, its integration into the factory of knowledge comes with challenges that must be addressed to ensure its benefits are equitably distributed. By prioritising ethical considerations, inclusivity, and sustainability, we can harness the full potential of AI while mitigating its risks. As we move forward, a collaborative effort between technologists, policymakers, and society at large will be essential to build a knowledge ecosystem that reflects the diverse, complex, and interconnected world we inhabit.

To change the Western capitalist production model of AI, it is essential to examine language as the pivotal tool enabling AI to produce scientific texts. Language shapes the discourse and frames the knowledge that AI generates, reinforcing cultural, social, and epistemological paradigms. The critical question, then, is how to design AI systems capable of generating scientific texts that are genuinely aware of DEI. Addressing this challenge requires rethinking both the linguistic structures embedded in AI models and the values they reflect, ensuring that AI-powered knowledge production becomes a force for global equity and representation.

3. Moving towards language as a tool for DEI aware scientific text

The current model of scientific discourse and language reflects systemic inequities, often serving as a mirror of historical biases and social structures. While science strives for objectivity, its language is not immune to the cultural, political, and economic forces that shape it. The question is whether it is possible to produce DEI aware scientific text that questions and moves beyond surface-level inclusivity toward a transformative shift in epistemic foundations. Through the integration of decolonial approaches, language as a tool for equity, and practical initiatives rooted in community practices, science can evolve as an inclusive force for empowerment and well-being.

Inclusive language, while important, addresses only the surface level of systemic inequities. Neutralizing gendered terms, modifying adjectives, or eliminating offensive jargon are steps in the right direction but fail to address the deeper cultural and epistemic biases embedded in scientific communication. Genuine inclusivity requires a shift toward inclusive thinking — an approach that questions the very foundations of knowledge production and the structures that sustain them.

3.1 The Role of decolonial approaches to science

Decoloniality offers a powerful lens to interrogate and transform scientific discourse. At its core, decoloniality seeks to dismantle the

dominance of Western epistemologies and create space for diverse, contextually rooted ways of knowing. Language is central to this effort. It serves as a repository of cultural identity, collective memory, and intellectual autonomy. By integrating decolonial perspectives into scientific language, researchers can amplify underrepresented voices and foster equitable knowledge production.

Community-centered practices, such as "tequio" in Mexico, "minga" in Colombia, *lumbung* in Euskerra; *ko* in Guatemala, in certain Mayan traditions; *kumunytunk* in Mixe or "ubuntu" in Africa, offer valuable insights into collective action, reciprocity, and resource-sharing (Aguilar Gil, 2024). These concepts challenge the individualistic and competitive ethos of mainstream science, proposing instead a collaborative model that values shared responsibility and mutual benefit. Incorporating these frameworks into scientific discourse not only enriches its epistemic diversity but also promotes values of equity and solidarity.

The shift toward DEI-aware scientific language is not without challenges. Resistance from entrenched power structures, the complexity of integrating diverse epistemologies, and the need for sustained resources are significant barriers. However, these challenges also present opportunities for innovation and collaboration. By partnering with local communities, interdisciplinary teams, and global networks, researchers can develop solutions that are both contextually relevant and globally impactful.

The role of technology in this transformation is both a challenge and an opportunity. While AI and other digital tools have the potential to amplify biases, they also offer unprecedented possibilities for inclusivity. For example, real-time translation tools can bridge linguistic divides, while open-access platforms can democratize access to scientific knowledge.

3.2 Practical Applications: Tools and Frameworks for DEI-Aware Science

AI and NLP tools often reflect the biases of the data they are trained on, perpetuating exclusion and epistemic violence. To counter this, AI systems should be designed with the explicit goal of meeting the needs of underrepresented communities. For example, creating NLP

models that recognize and support indigenous languages can help preserve cultural heritage and improve accessibility. Similarly, generative AI tools tailored to local contexts can empower communities to tell their stories and address specific challenges (*idem*).

Traditional metrics of success, such as citation counts or journal impact factors, often reinforce exclusionary practices. A DEI-aware approach would emphasize the social impact of research, the inclusivity of research teams, and the relevance of findings to diverse populations. For instance, a project on climate resilience should consider its applicability to vulnerable communities, not just its technical innovation.

Besides, education plays a crucial role in shaping the future of science. By embedding ethics, equity, and sustainability into scientific curricula, institutions can foster a generation of researchers who are mindful of the societal implications of their work. Case studies on topics such as algorithmic bias or the environmental impact of AI can help students critically engage with these issues.

The journey toward DEI-aware scientific texts is a transformative process that challenges the *status quo* of knowledge production. By moving beyond inclusive language to foster inclusive thinking, integrating decolonial perspectives, and leveraging technology for equity, science can redefine its role in society. The examples of using electronic waste for robotics, applying data feminism, and designing AI tools for underrepresented communities illustrate the potential of this approach. Ultimately, the goal is to create a scientific discourse that is not only inclusive and equitable but also deeply connected to the collective well-being of humanity.

4. Western Epistemologies in AI Design and Development

The influence of Western values on scientific discourse has shaped the trajectory of technological advancements, particularly in AI. While this discourse has driven remarkable innovation, it also perpetuates systemic biases and inequities embedded in AI tools, leading to unfair and irresponsible outcomes. AI systems, informed by predominantly Western-centric paradigms, often fail to consider diverse cultural, social, and epistemological contexts, which raises ethical, equity, and accountability concerns.

Western scientific discourse privileges empirical, measurable, and quantifiable knowledge while sidelining qualitative, context-sensitive, and indigenous forms of understanding. This reductionist approach creates AI systems that excel at optimizing tasks but struggle with nuanced socio-cultural complexities. For instance, NLP tools prioritize English-language datasets, which marginalize non-Western languages and dialects. This imbalance contributes to a digital hegemony, where tools and technologies cater disproportionately to Global North populations, leaving the Global South underrepresented.

4.1 Reinforcing power structures

Such Western-centric epistemologies manifest in AI design choices, particularly in training datasets. Biases in these datasets — whether unintentional or the result of implicit systemic structures — reflect historical and cultural inequalities. For example, facial recognition technologies often exhibit lower accuracy for individuals with darker skin tones because the datasets used for training predominantly feature lighter-skinned individuals. These disparities have far-reaching consequences, including reinforcing racial stereotypes and excluding marginalized groups from equitable participation in technological ecosystems.

Western scientific discourse's emphasis on objectivity and universalism overlooks the localized and intersectional nature of fairness. This has allowed AI systems to reinforce existing power structures rather than challenge them. For instance, predictive policing algorithms trained on historical crime data often reflect the systemic racial biases inherent in law enforcement practices (Ricaurte, 2019). By treating this biased data as "neutral," such systems perpetuate discriminatory patterns, disproportionately targeting minority communities.

Similarly, Western-centric scientific discourse often aligns with neoliberal economic models, emphasizing efficiency and profitability over equity and inclusion. This approach prioritizes technological solutions that serve corporate interests at the expense of social responsibility. As a result, AI tools frequently exploit vulnerable communities for data collection or testing without adequately addressing their specific needs or compensating them for their contributions.

The unchecked propagation of Western values in scientific discourse and AI development has ethical and social implications. Irresponsible AI tools, shaped by such values, often lack transparency and accountability.

Furthermore, the focus on technological solutions often disregards the broader socio-economic and environmental contexts in which AI operates. The energy-intensive nature of AI training and deployment exacerbates environmental inequities, disproportionately affecting communities already vulnerable to climate change (Vandana, 2015). By centering Western industrial priorities, current AI practices overlook alternative, sustainable approaches that might emerge from non-Western perspectives.

4.2 Toward Fair and Responsible AI

To mitigate the impact of Western values on AI tools, it is imperative to decolonise scientific discourse and rethink AI development frameworks. This requires diversifying training datasets to include under-represented languages, cultures, and experiences. It also demands interdisciplinary collaboration to integrate insights from social sciences, ethics, and indigenous knowledge systems into AI design (Vargas-Solar *et al.*, 2024).

Moreover, AI development must shift from a purely efficiency-driven model to one rooted in fairness, accountability, and sustainability. For example, participatory design approaches can involve local communities in shaping AI tools that address their specific needs and priorities. Transparent mechanisms for auditing AI systems should also be implemented to ensure accountability and foster trust.

Ultimately, the challenge lies in transforming scientific discourse to embrace plurality, equity, and inclusivity. By acknowledging and addressing the biases inherent in Western-centric frameworks, we can create AI tools that not only perform efficiently but also uphold the values of justice and fairness for all. This transformation is not merely a technical endeavor but a profound cultural shift, redefining the role of science and technology in shaping a more equitable and responsible future.

To transform the production model of AI as a knowledge factory, principles from *tequiology* and technofeminist approaches must be

integrated into the AI development pipeline. This principle emphasizes collaboration, shared responsibility, and the prioritization of collective well-being over profit-driven motives. AI systems developed under *tequiology* principles would be inclusive, addressing the needs of marginalized communities and promoting sustainable practices that reduce environmental and social exploitation.

Technofeminist approaches further enrich this vision by challenging the hierarchical, male-dominated structures of AI development. They advocate for embedding feminist ethics into the design and deployment of technology, focusing on dismantling epistemic violence and ensuring gender equity in the narratives that inform AI tools. By emphasizing care, accountability, and intersectionality, technofeminism provides a critical lens for addressing systemic biases and fostering equitable representation in AI outputs. Together, these approaches urge a radical departure from extractivist AI practices, advocating instead for systems that uphold cultural diversity, environmental sustainability, and social equity, transforming AI from a tool of domination into a medium of inclusive, ethical, and collaborative knowledge production.

Conclusion

This paper advocates for a transformative approach to scientific discourse that moves beyond superficial inclusivity to address the systemic inequities embedded in the language and frameworks of science. By adopting decolonial perspectives and integrating diverse epistemologies, scientific discourse can challenge dominant paradigms and create space for underrepresented voices. Practical applications, such as AI tools designed for marginalized communities and metrics emphasizing social impact over traditional measures, illustrate the tangible benefits of this shift. The journey toward DEI-aware scientific language requires a commitment to ethics, equity, and sustainability, but it also offers a profound opportunity to redefine science as a collective, inclusive, and transformative force for humanity. Through interdisciplinary collaboration, technology, and education, we can reimagine the role of science in shaping a more equitable and inclusive world.

Bibliography

- Aguilar Gil, Yásnaya Elena (2013). “Diversidad lingüística y comunalidad”. *Cuadernos del Sur*, 18(34), 71-82.
- Aguilar Gil, Yásnaya Elena (2024). “Tequiologías charla sobre tecnologías colaborativas en resistencia”. <https://centroculturadigital.mx/revista/tequiologias>
- Braz Matheus Viana, Tubaro Paola, Casilli Antonio (2024). “Fabricar os dados: o trabalho por trás da Inteligência Artificial”. *As novas infraestruturas produtivas: digitalização do trabalho, e-logística e indústria 4.0*, 105-120.
- Bourdieu Pierre (1992). *Les Règles de l'art, Genèse et structure du champ littéraire*. Paris: Éditions du Seuil.
- Couldry Nick, Mejias Ulises. (2019). “Making data colonialism liveable: how might data’s social order be regulated?”. *Internet Policy Review*, 8(2). <https://policyreview.info/articles/analysis/making-data-colonialism-liveable-how-might-datas-social-order-be-regulated>
- Couldry Nick, Mejias Ulises. (2020). *The costs of connection: How data are colonizing human life and appropriating it for capitalism*. Stanford: Stanford University Press.
- Couture Stéphane, Toupin Sophie, Mayoral Baños Alejandro. (2025). “Resisting and Claiming Digital Sovereignty: The Cases of Civil Society and Indigenous Groups”. *Policy & Internet*. https://www.researchgate.net/publication/387914036_Resisting_and_Claiming_Digital_Sovereignty_The_Cases_of_Civil_Society_and_Indigenous_Groups
- Durand Gilbert. (2021). *Les structures anthropologiques de l'imaginaire*. Paris: Armand Colin.
- Escobar Arturo, de Souza Filho Carlos Federico Marés, Nunes João Ariscando, Coelho João Paulo Borges, dos Santos Laymert García, de Oliveira Neves Lino João, Arenas Luis Carlos *et al.* (2020). *Another knowledge is possible: Beyond northern epistemologies*. Verso Books.
- Figueiroa Silvia. F. de M. (2023). “Postcolonial and Decolonial Historiography of Science”, in M. L. Condé, M. Salomon (eds), *Handbook for the Historiography of Science*. Berlin: Springer, 1-19.
- Harding Sandra. (2016). “Latin American decolonial social studies of scientific knowledge. Alliances and tensions”. *Science, Technology, & Human Values*, 41(6), 1063-1087.
- Laws E., Calvert M., Sapey E., Denniston, A. (2024). “Tackling algorithmic bias and promoting transparency in health dataset”. *The Standing Together consensus recommendations*. [https://www.thelancet.com/journals/landig/article/PIIS2589-7500\(24\)00224-3/fulltext](https://www.thelancet.com/journals/landig/article/PIIS2589-7500(24)00224-3/fulltext)

Miceli Milagros, Tubaro Paola, Casilli Antonio. A., Le Bonniec Thomas, Wagner Camilla Salim, Sachenbacher Laurenz (2024). *Who Trains the Data for European Artificial Intelligence?* Doctoral dissertation, European Parliament.

Raus Rachele, Tonti Michela (eds) (2025). Intelligence artificielle, corpus et diversité linguistique : enjeux et perspective. *Langages*, 237.

Ricaurte Paola (2019). “Data epistemologies, the coloniality of power, and resistance”. *Television & New Media*, 20(4), 350-365.

Ricaurte Paola, Vargas-Solar Genoveva, Feldfeber Ivana, Mosquera Diana, Josiowicz Alejandra, Vela Susana Cadena, Brussa Virginia, Alonso i Alemany Laura, Ovalle Anaelia, Zaragoza Cano Liliana (2025). “Feminist AI for/by the majority world: Feminist AI Research Network, Latin American and Caribbean Hub”, in Ed. Leah Junck *et al.* (eds), *Majority World, Oxford Intersections: AI in Society*. Oxford, England: Oxford University Press [forthcoming].

Siles Ignacio, Gómez-Cruz Edgar, Ricaurte Paola. (2023). “Toward a popular theory of algorithms”. *Popular Communication*, 21(1), 57-70.

Vandana Shiva (2015). *Earth democracy: Justice, sustainability, and peace*. Berkeley: North Atlantic Books.

Vandana Shiva (2016). *Biopiracy: The plunder of nature and knowledge*. Berkeley: North Atlantic Books.

Vargas-Solar Genoveva (2022). “Calling for a feminist revolt to decolonise data and algorithms in the age of Datification”. <https://arxiv.org/abs/2210.08965>

Vargas-Solar Genoveva, Cerquitelli Tania, Montorsi Ariana, Salvai Stefania, Sangineti Maria Teresa, Darmont Jérôme, Favre Cécile (2022). “Promoting equity, diversity and inclusion: policies, strategies and future directions in higher education, research communities and business”. *2022 IEEE International Conference on Big Data (Big Data)*, 4710-4718.

Vargas-Solar Genoveva, Espinosa-Oviedo Javier-Alfonso, Bennani Nadia, Mauri Andrea, Zechinelli-Martini José-Luis, Catania Barbara, Ardagna Claudio A., Bena Nicola (2024). “Decolonizing Federated Learning: Designing Fair and Responsible Resource Allocation Position Paper”, *ACS/IEEE 21st International Conference on Computer Systems and Applications (AICCSA 2024)*.

Key-passages and artificial intelligence. How can we use hidden layers to explore new linguistic objects?

Laurent Vanni, laurent.vanni@univ-cotedazur.fr
Damon Mayaffre, damon.mayaffre@univ-cotedazur.fr
Hadi Mahmoudi, hadi.mahmoudi@univ-cotedazur.fr
Université Côte-d'Azur — CNRS, France

Introduction

The performance of AI in terms of automatic text classification, generation and translation is now remarkable. However, the community continues to raise two complementary scientific objections to the use of deep learning in text analysis: (i) from the perspective of computer science, the issue of interpretability of models and algorithms, as these often function as black boxes; and (ii) from the perspective of the social sciences and humanities, the challenge of providing linguistic explicability of the results obtained to enhance the semantics of texts.

To shed light on the black box (both algorithmic interpretability and linguistic explicability), we challenge the notion of the corpus as a training data set. In our view, large language models (LLMs) face the limitation of relying on unqualified textual data (data that has not been structured into a corpus). They operate on language rather than discourse, which risks standardizing meaning, whereas our goal is to capture all its nuances.

We then evaluate the models' ability to capture the linguistic features that underpin the analysis. While these models are effective in classifying, generating, or translating texts, their processing must now highlight the salient linguistic elements learned from the studied texts — such as tokens, parts of speech, morpho-syntactic labels, patterns, or motifs — to enhance human interpretability.

We present the findings of our analysis on two distinct corpora:

- i) French Presidential Speeches: This corpus spans from Charles de Gaulle (1958) to Emmanuel Macron (2024), covering the 5th Republic. Using a trained model, we examine the specificities of Emmanuel Macron's discourse in comparison to his predecessors.
- ii) French Literature: Focusing on classical 19th-century authors, we propose an intertextuality detection task centered on Maupassant's

works, comparing them with those of his predecessors Sand, Flaubert, and Zola.

In this contribution, “*passages*” refer to sequences of N contiguous tokens used to feed the deep learning model (training, testing and validation). *Key-passages* are significant passages according to the prediction score given by the output layer of the neural network. The model and the visualisations we present are implemented by the *Hyperbase* web platform (Vanni, 2024).

1. Corpora vs. Big Data

Large Language Models (LLMs) are trained (or rather pre-trained) on large text datasets to generate natural language. GPTs (Generative Pre-trained Transformers) are based on word representations learned from heterogeneous data that are decorrelated from the fine-tuned final task.

Unlike the positivist or ontological approaches to large language models (LLMs), the corpus or textual linguistics perspective — exemplified by and, in France, by Rastier (2011) — argues that meaning is not inherent in language itself but emerges solely within discourse or corpora. The meaning of texts and words is interpreted within the framework of constituted corpora, identified discursive genres, historically or politically situated discourses, and specific topics. For instance, the term “freedom” cannot hold the same average representation or meaning in an 18th-century tragedy, a speech by Donald Trump, or a washing machine manual. The extensive datasets used as training corpora for pre-training are not genuine corpora; their data lacks linguistic definition and qualification. As a result, the representations they assign to words are merely semantic approximations, often tainted by semantic biases, rendering them inadequate.

1.1 Political Discourse case study

We propose an endogenous learning approach based on a selected corpus of presidential speeches. The generic homogeneity (French political discourse and, more specifically, presidential discourse), dia-

chronic homogeneity (French Fifth Republic, mid-20th century to early 21st century), and thematic homogeneity (addresses to the nation, New Year's greetings to the French people) of the corpus guarantee an effective representation of the words for the analysis.

A total of 1,000 speeches (about 120 per president) were extracted from institutional websites such as the Elysée Palace¹ and French public life². Each speech was cleaned and normalized before being tokenized and tagged using *Hyperbase*. Our final corpus consists of over 3 million words, distributed as in Fig. 1:

Textes	Occurrences	Vocables	Hapax
deGaulle	226033	12191	1593
Pompidou	242923	13350	1725
Giscard	380709	13370	1383
Mitterrand	723677	21584	4482
Chirac	463528	14929	1824
Sarkozy	270355	12224	1388
Hollande	280037	11604	1136
Macron	459239	16924	2902

Fig 1. Presentation of the French political discourse corpus. For each author: the number of words (Occurrences), the size of the vocabulary (Vocables), the number of isolated words — i.e only one occurrence (Hapax).

Within the presidential corpus used as the training dataset for a classification task, we extract the key passages deemed significant for each president (Fig. 1). The extraction identifies the most significant passages in the training corpus, determined by a key-passage detection algorithm that utilizes the activation scores from the neural network's output layer (Vanni *et al.*, 2020).

1.2 French literature case study

The second corpus deals with intertextuality detection issues. Intertextuality, as defined by Kristeva (1969), refers to the concept that a text is not an isolated, self-contained entity, but rather a mosaic of references to and influences from other texts. She introduced the term intertextuality in the 1960s, arguing that every text is a product of the interaction between previous texts. Kristeva's theory builds on the ideas of Mikhail Bakhtin and his concept of dialogism. According to Kristeva, texts do not exist in isolation but

¹ <https://www.elysee.fr>

² <https://www.vie-publique.fr>

are always shaped by and linked to others through citations, allusions, and shared cultural or linguistic frameworks. Therefore, meaning is always interdependent and constantly evolving through the interplay of multiple texts.

In Kristeva's view, when we engage with a text, we are simultaneously engaging with a network of past texts and cultural norms, which influence how we interpret it. This interconnectedness often manifests through identifiable traces of prior texts within a given work. Thus, intertextuality can consist of a single word, the linear co-occurrence of N words, or key markers that include words, lemmas, and grammatical codes, which may not appear consecutively.

In a corpus-driven approach and using deep learning models, an intertextuality detection task can be considered similar to an authorship attribution task where the test data set is an unseen given author that we except was influenced by the authors of the training set (Vanni *et al.*, 2024). In our contribution, we have chosen to analyze four prominent 19th-century authors: Georges Sand, Gustave Flaubert, Émile Zola and Guy de Maupassant. To neutralise the Literary genre we focus on *Romans* (novels) totalling 8.4 million words, distributed as follows (Fig. 2):

Textes	Occurrences	Vocables	Hapax
Sand	3699534	52296	11989
Flaubert	629064	30833	4361
Zola	3465947	47985	8187
Maupassant	630921	25922	2471

Fig 2. Presentation of the French literature corpus. For each author: the number of words (Occurrences), the size of the vocabulary (Vocables), the number of isolated words — i.e only one occurrence (Hapax)

The intertextuality detection task we address differs slightly from a standard classification task. In our study, we aim to analyze the works of Guy de Maupassant to detect traces of inspiration from his predecessors — George Sand, Gustave Flaubert, and Émile Zola. To achieve this task, we remove Maupassant's works from the training dataset and reserve them for testing (Fig. 3). In this way, we consider the authorship attribution task to be equivalent to an intertextuality detection task within Maupassant's works. Specifically, we require the model to classify Maupassant's samples as belonging to one of the other authors (see section 3.1).

The key passages in Maupassant's works are those that obtain the highest prediction scores relative to the works of the other three authors.

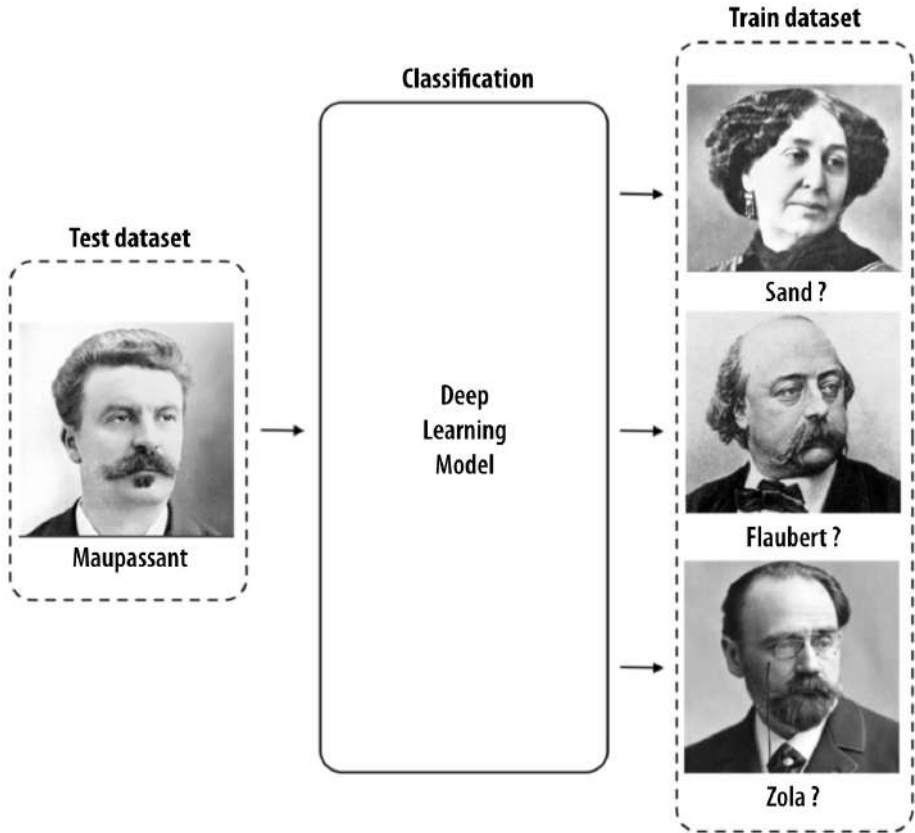


Fig. 3. Deep Learning classification model applied to intertextuality detection task

2. Prediction vs. Description: Implementation of an Interpretable Model

Known as “black boxes”, deep learning models sometimes fail to reveal the markers they use to make decisions. The information buried in the hidden layers of the system is often considered incomprehensible.

However, many studies show that it is possible to visualize the relevant words or linguistic units identified by a model to perform a given task (Li *et al.*, 2022). Convolutional models, for example, have been shown to be particularly interpretable through deconvolution techniques, which reveal local textual saliency useful for classification (Vanni *et al.*, 2018).

The more recent and effective models, Transformers, are equipped with a powerful interpretability mechanism known as Self-Attention. This mechanism automatically assigns significant values that represent the importance of associations between words, regardless of their distance from each other (Vaswani *et al.*, 2017).

The Multichannel Convolutional Transformer (MCT) model we present is based on proven interpretability mechanisms and has three major advantages:

- Combination of convolutional models and transformers for proven effectiveness (Li *et al.*, 2021). The combination of convolution and self-attention allows for the reinforcement of a dual syntagmatic and paradigmatic perspective, theorized from the outset by Saussure (1972). Words are considered both in their immediate context (syntagmatic axis) and in their more distant “associative relationship” (paradigmatic axis) within the textual chain. In particular, convolution (with a window of 3 words in our work) accounts for local concatenations (the word within its n-gram or syntagm), while Transformer models capture short- and long-distance associations between words to reveal those that “go well together” (Halliday, 1976).
- A multi-channel architecture, Feng *et alii* (2021) which analyzes the text on three complementary linguistic levels: graphic form, morpho-syntactic label, and lemma. Looking at the graphic word addresses the materiality of the text, i.e., the raw text without cleaning, sorting, labelling, or interpretation. By examining the morpho-syntactic labels (i.e., the part of speech), the grammatical and syntactic dimensions of words and sentences are considered. Finally, the semantic dimension is introduced through the analysis of the lemma. In sum, the multichannel approach represents the text in its complexity, if not its entirety.
- Corpus-based analysis for training. The ultimate use of the model is unconventional: the classification task (prediction) serves only as a means to train the model to distinguish between the different classes (the presidents). The real task is to describe the linguistic features that distinguish each class. To achieve this, it is typically necessary to use new texts (a test set) independent of the training corpus. The MCT model allows us to go further with a key-passage detection algorithm. This algorithm prioritizes each of the segments

used during training based on an activation score recorded at the output layer of the network before the final activation function is applied (Vanni *et al.*, 2020).

The architecture of the Multichannel Convolutional Transformer (MCT) model, designed to be interpretable, is based on two learning phases:

First phase: Pre-learning based on a standard convolutional classification model (Kim *et al.*, 2014) applied to each word representation (forms, POS, lemmas, etc.). This phase extracts contextual word representations (embeddings) and identifies, through deconvolution, the salient syntagms needed to assign a text to an author.

Second phase: Learning from a transformer model that reports relevant word associations for classification. These two phases (convolution and transformer) are connected through the use of embeddings computed by the convolution phase (first phase), which are enriched with the positional encoding of each word and then input into the transformer (second phase), as shown in Fig. 4. Finally, the MCT model includes tools for visualizing the results, as presented in the next section.

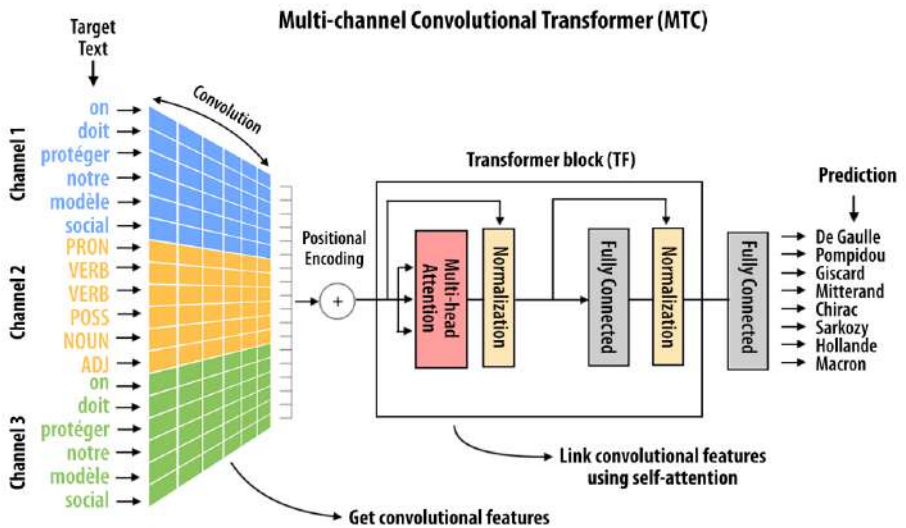


Fig. 4. Multichannel Convolutional Transformer (MCT) Architecture applied to the political corpus

In the next section, the hyperparameters are set to maximise both the interpretability and the accuracy of the model. The corpora are divided into sequences of 50 tokens (passages), the convolution sliding window corresponds to 3 tokens and the multi-head attention uses 4 heads that we sum to obtain a global attention score for each pair of tokens. 80% of the corpus is the used for training, and 10% is kept for both validation and testing.

3. Results

3.1. Political discourse learning

As shown in Tab. 1, the MCT model achieves decent performance in automatic text classification of the presidents while ensuring optimal interpretability. We illustrate the heuristic value of MCT in aiding expert analysis of Emmanuel Macron’s speeches through the lens of the humanities. For political scientists, Macron’s speeches, often exemplified by the expression “*en même temps*” (“at the same time”), are difficult to categorize (Mayaffre, 2021). However, the neural network outputs reveal a remarkable signature of Macron’s ideology, which he shares with, for example, the American Democrats: a caring discourse.

Model	Accuracy	Precision	Recall
FastText (Joulin et al. 2016)	nc	0.6760	0.6760
CNN-Form (Kim et. al. 2014)	0.9063	0.9518	0.9927
CNN-POS (Kim et. al. 2014)	0.5482	0.8045	0.6462
CNN-Lemma (Kim et. al. 2014)	0.8803	0.9372	0.9802
MCT	0.9032	0.9646	0.9774

Tab. 1. MCT performances compared to baseline models such as FastText (Joulin et al. 2016) and CNN (Kim et. al. 2014) applied on word full forms (Form), Part-Of-Speech (POS), lemmas (Lemma)

According to the model’s results, and in comparison with the speeches of his predecessors (de Gaulle, Pompidou, Giscard, Mitterrand, Chirac, Sarkozy, Hollande), Macron’s speech is characterized by words and word combinations revolving around “care”, “protection”, and “attention”. To counterbalance a conservative social program — such as reducing pension and unemployment rights — Macron cultivates a discourse of well-being and benevolence.

One of the highest prediction score is obtained by a passage (thus a key-passage) from Macron's speech on Europe at the Sorbonne in April 2024:

Et donc, on doit bien protéger notre santé en appliquant strictement nos standards sanitaires. On doit protéger notre modèle social, en impliquant là aussi nos standards sociaux. Et on doit protéger nos ambitions climatiques, en défendant nos standards environnementaux. Sinon, nous allons inventer un continent qui sur-contraint les producteurs sur son sol et par sa politique commerciale, lève les contraintes sur les produits qu'il importe [...]³ (Macron, *Discours à la Sorbonne*, 25 avril 2024)

The MCT highlights key markers (word forms in blue, lemmas in green, part-of-speech in orange) and self-attention links (in red) that are crucial for selecting this passage. (Fig. 5).



[...] bien protéger **notre santé** en appliquant strictement nos standards sanitaires. **On** doit protéger **notre** **modèle social**, en impliquant **là aussi** **DET:Plur:Poss standards** sociaux. Et **on** doit protéger **nos ambition** climatiques, en **VERB:Pres:Part** **notre** standards environnementaux. Sinon, nous allons inventer **un continent** qui sur-contraint les [...]

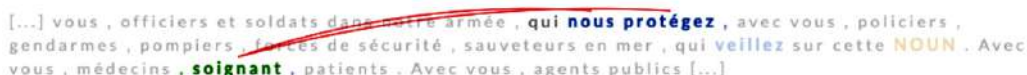
Fig. 5. MCT outputs implementation apply of E. Macron speeches available on Hyperbase

The visualization of the markers detected by the MCT model aids human interpretation, revealing that Macron's speech, in this case, is structured around the anaphoric repetition of “*devoir protéger*” (must protect: “*devons protéger*” “*devra protéger*”, etc.). More specifically, all the contextual vocabulary surrounding “*prendre soin*” (taking care) develops the theme of protection and benevolent solidarity. In the passage, MCT notes that “*devoir protéger*” (must protect) is associated with the plural possessive determiner, either through the lemma “*notre*” (our) or through the morpho-syntactic label (DET: Plur/Poss). Macron's speech rhetorically constructs a protective community (“*nous*”, “*nos*”, “*notre*”, but also “*continent*”) in which the individual must be reassured (“*protéger*”). Similarly, elsewhere in the corpus, the model draws our attention to the following key passage:

³ DeepL (<https://www.deepl.com/translator>) translation: “So we have to protect our health by strictly applying our health standards. We must protect our social model, by applying our social standards. And we must protect our climate ambitions by defending our environmental standards. Otherwise, we are going to invent a continent that over-constrains producers on its own soil and, through its trade policy, lifts the constraints on the products it imports [...]”.

[...] vous, officiers et soldats dans notre armée, qui nous protégez, avec vous, policiers, gendarmes, pompiers, forces de sécurité, sauveteurs en mer, qui veillez sur cette soirée. Avec vous, médecins, soignants, patients. Avec vous, agents publics [...] ⁴ (Macron, *Vœux aux Français*, 31 décembre 2022).

In a long list of professions, Macron addresses all French people (police officers, gendarmes, firefighters, rescue workers, doctors, patients, etc.). However, the model specifies which words, lemmas, or grammatical codes in the passage have been (i) particularly activated by the convolutional neural network and (ii) linked together or activated together, according to self-attention. The same theme of care, community, and healing as before then easily emerges (Fig. 6):



[...] vous , officiers et soldats dans notre armée , qui nous protégez , avec vous , policiers , gendarmes , pompiers , forces de sécurité , sauveteurs en mer , qui veillez sur cette NOUN . Avec vous , médecins , soignant , patients . Avec vous , agents publics [...]

Fig. 6. MCT outputs implementation apply of E. Macron speeches available on Hyperbase

The lemma “soignant” (in green, both singular and plural) is activated by the network, indicating a high activation score from deconvolution, and attracts the most attention in the passage. The lemma is linked to the verb “protégez” (to protect) and the plural pronoun “nous” (we). Macron, in a speech that is both performative and curative, celebrates the resilience, healing (“soignant”), and protection (“protégez”) of the French people in this challenging post-social event called “gilets jaunes” and post-Covid; if not already healed and soothed by the virtues of discourse, then at least reunited and unified (“nous”).

A final key passage extracted automatically confirms the analysis and refines the semantic or interpretative pathway generated by the AI (Fig. 7):

[...] notre santé. Pour tout cela, je forme pour nous tous, des vœux d'unité, des vœux d'audace, des vœux d'ambition collective, et des vœux de bienveillance. Des crises, mes chers compatriotes, ensemble,

⁴ DeepL (<https://www.deepl.com/translator>) translation: “you, officers and soldiers in our army, who protect us, with you, police officers, gendarmes, firefighters, security forces, lifeguards at sea, who watch over this evening. With you, doctors, carers, patients. With you, public servants [...]”.

nous en avons tant surmontées [...] ⁵ (Macron, *Vœux aux Français*, 31 décembre 2022).



Fig. 7. MCT outputs implementation apply of E. Macron speeches available on Hyperbase

As illustrated by the machine outputs, Emmanuel Macron's speeches are rhetorically woven around the lexicon of protection and care, combining collective expressions (we, our, together) with protective language (protect, care, health, benevolence). For Macron, the promise of empathy, care, and benevolence becomes a fundamental political program.

3.2. Intertextuality detection

To evaluate the MCT model on the French literature corpus, we built a standard classification task for the entire corpus. The MCT model demonstrates strong performance in automatic text classification, maintaining optimal interpretability (Tab. 2).

Tab. 2. MCT performances compared to baseline models such as FastText (Joulin et al. 2016) and CNN (Kim et. al. 2014) applied on word full forms (Form), Part-Of-Speech (POS), lemmas (Lemma)

Model	Accuracy	Precision	Recall
FastText (Joulin et al. 2016)	nc	0.8600	0.8600
CNN-Form (Kim et. al. 2014)	0.9429	0.9480	0.9019
CNN-POS (Kim et. al. 2014)	0.7100	0.7713	0.6068
CNN-Lemma (Kim et. al. 2014)	0.9237	0.9236	0.8936
MCT	0.9597	0.9525	0.9373

Using feature extraction, MCT highlights the intertextuality traces (convolutional an self-attention markers) that are used to explain the model's decision.

3.2.1 Maupassant's rhythmic resemblance to Zola

When comparing Maupassant with Zola, MCT identifies the following sequence as the most similar to Zola:

⁵ DeepL (<https://www.deepl.com/translator>) translation: “our health. For all these reasons, I wish us all unity, boldness, collective ambition and goodwill. My dear compatriots, together we have overcome so many crises [...]”..

[...] , torturé du désir impérieux de laisser sortir toute sa joie . Les deux frères , en deux fauteuils pareils , les jambes croisées de la même façon , à droite et à gauche du guéridon central , regardaient fixement devant eux , en des attitudes semblables , pleines [...] ⁶ (Maupassant, *Pierre et Jean*).

The markers identified by MCT, with a high activation score from convolution, reveal that the syntagme “guéridon central” (a small round table) and the repetition of the grammatical code “PUNCT” act as key markers, helping to classify the entire passage to Zola (Fig. 8).

[...] **PUNCT** **torturer** de désir **impérieux** de laisser sortir toute sa **joie** **SENT** Les deux **frère**
PUNCT en deux fauteuils **pareils** **PUNCT** le jambes croisées de la même façon **PUNCT** à
droite et à **gauche** du **guéridon central** , **regardaient** **fixement** devant eux **PUNCT** en des
attitudes **semblables** , **plein** [...]

Fig. 8. MCT outputs implementation apply of Maupassant works available on Hyperbase.
Key-passages form Zola attribution class

With 52 occurrences of “gueridon” in Zola’s works and over 73 occurrences across the entire corpus, as well as 48 occurrences of “central” in Zola’s works and over 60 occurrences in the whole corpus, the expression “guéridon central” appears to be one of Zola’s key passages. This is confirmed by a Fisher’s exact test of both words with a score of 8.1, indicating that Zola overuses this lexical field (Fig. 9).

The “PUNCT” code refers to the punctuation tag, primarily corresponding to the colon, semicolon, and comma. This result, confirmed by another Fisher’s exact test (Fig. 10), shows that MCT appears to also detect Zola’s rhythm, with long sentences incorporating multiple pauses. Indeed, the Self-attention (red links) highlights the significance of punctuation in detecting Zola’s key passages.

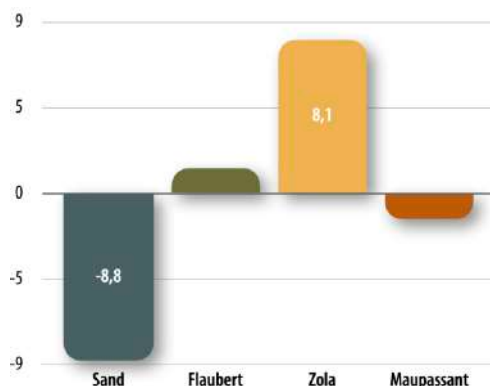


Fig. 9. Fisher exact test of words “guéridon” and “central” (sum)

⁶ DeepL (<https://www.deepl.com/translator>) translation: “tortured by the imperious desire to let out all his joy. The two brothers, in two identical armchairs, with their legs crossed in the same way, to the right and left of the central pedestal table, were staring straight ahead, in similar attitudes, full of” [...].

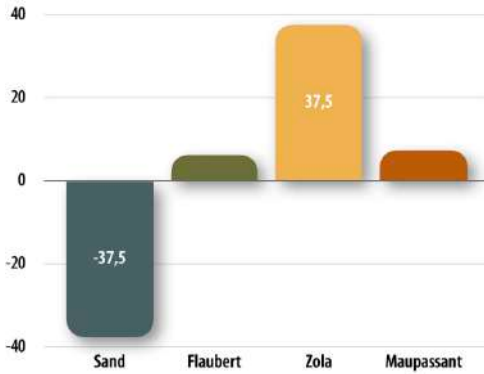


Fig. 10. Fisher exact test of the "PUNCT" part-of-speech corresponding to the colon, semicolon, and comma

3.2.2 Stylistic Echoes: Maupassant's Resemblance to Flaubert

When comparing Maupassant with Flaubert, MCT identifies the following sequence as key-passage to Flaubert:

[...] des mouches éclatantes Qui bourdonnent autour des puanteurs cruelles , Golfes d' ombre ; E , candeurs des vapeurs et des tentes , Lances des glaciers fiers , rois blancs , frissons d'ombrelles ; I , pourpre , sang craché , rire des lèvres [...] ⁷ (Maupassant, *La Vie errante*).

The markers identified by MCT, with a high activation score from convolution, reveal that the lemma "*autour*" (around) and the grammatical code "NOUN:Masc" act as key markers, helping to classify the entire passage as Flaubert (Fig. 11).

[...] des mouches éclatantes Qui bourdonner autour des puanteurs cruelles , Golfes d' NOUN:Sing ; e , candeurs des vapeurs et des NOUN:Fem:Plur , lance des glaciers ADJ:Masc:Plur , rois blancs , NOUN:Masc:Plur d' NOUN:Fem:Plur ; ADJ , NOUN:Fem:Sing , NOUN:Masc:Sing craché , VERB:Inf des lèvres [...]

Fig. 11. MCT outputs implementation apply of Maupassant works available on Hyperbase. Key-passages form Flaubert attribution class

The lemma "*autour*" is used 290 times in Flaubert's works out of 2,153 occurrences in the entire corpus. A Fisher's exact test confirms that this lemma is specific to Flaubert (Fig. 12). When followed by the preposition (also known as an adposition) "*de*", the syntagm requires substantives as complements. This is precisely what the self-attention mechanism detects through its association with the part-of-speech tag "NOUN:Masc."

⁷ DeepL (<https://www.deepl.com/translator>) translation: "dazzling flies buzzing around cruel stench, gulfs of shadow; E , sweetness of vapours and tents, spears of proud glaciers, white kings, shivers of parasols; I , crimson, spat blood, laughter of lips [...]".

Furthermore, nouns are a distinctive feature of Flaubert's writing (Fig. 13), reflecting his descriptive style. In this sample, MCT focuses on a highly descriptive passage. The model combines local and distant associations to identify the most relevant markers for its classification decision. This second example suggests that these linguistic markers can be more complex than simple words, offering a new representation of intertextuality for a given author.

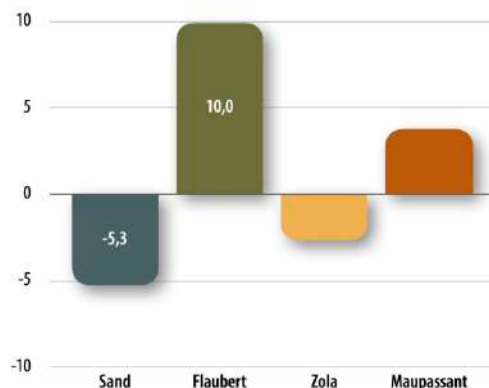


Fig. 12. Fisher exact test of the "autour" lemma

3.2.3 Sand's motifs in Maupassant's works

When comparing Maupassant with Sand, MCT identifies the following passage as characteristic of Sand:

[...] l' œil et dans l' esprit la vision rapide de tout ce qu' on ne peut pas dire , et permet aux gens du monde une sorte d' amour subtil et mystérieux , une sorte de contact impur des pensées par l' évocation simultanée , troublante et sensuelle comme une [...] ⁸ (Maupassant, *Bel-Ami*).

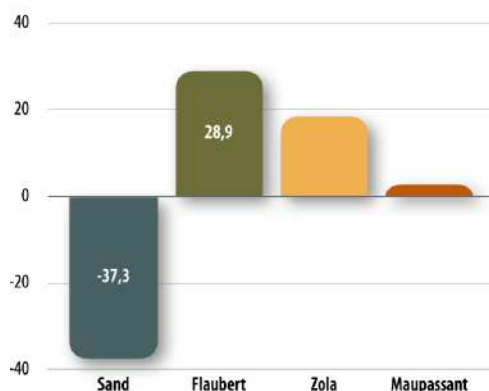


Fig. 13. Fisher exact test of "NOUN: Masc" part-of-speech

In this example, the markers identified by MCT reveal a complex multidimensional pattern made from lemmas, part-of-speeches and graphical full-forms. The first indication is the usage of several adjectives ("subtil", "mystérieux", "troublant", "sensuel") which is specific to Sand (+7,5 using Fisher exact test on the part-of-speech "ADJ" using the training corpus). The second point is the coordination conjunction ("CCONJ") and especially the word "et" (and) which is another statistical marker of Sand (Fig. 14).

⁸ DeepL (<https://www.deepl.com/translator>) translation: "the eye and in the mind the rapid vision of all that cannot be said, and allows the people of the world a kind of subtle and mysterious love, a kind of impure contact of thoughts by the simultaneous evocation, disturbing and sensual as a [...]".

[...] **le NOUN: Masc: Sing CCONJ dans le esprit le vision rapide de ADJ: Masc: Sing** ce qu'on ne pouvoir pas dire **PUNCT CCONJ permettre à gens du monde une sorte d' amour subtil et mystérieux**, une sorte de contact impur des pensées par l' évocation simultanée , troublant et sensuel **ADP** une [...]

Fig. 14. MCT outputs implementation apply of Maupassant works available on Hyperbase. Key-passages form Sand attribution class

To fully assess the heuristic added value of the model in this example, it is essential to analyze the convolution heatmap generated by the multichannel architecture (Fig. 15). We observe that the activation peaks are centered around the pattern “ADJ et ADJ”, and the word “sorte” (sort of) appears positively twice in the passage, reinforcing the repetition.

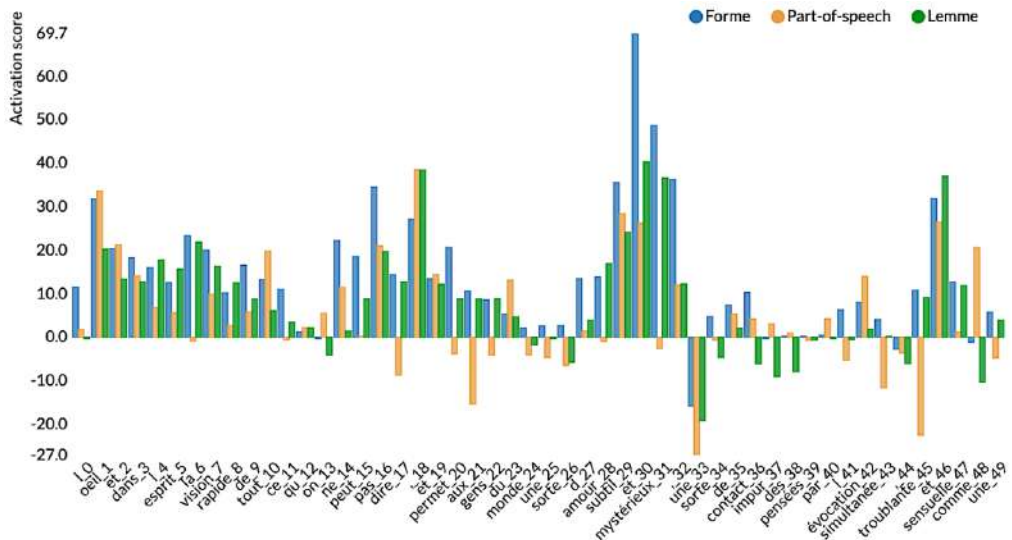


Fig. 15. Convolution heatmap: activation score of the convolution filters for each channel

The markers highlight a complex pattern, also referred to as a motif (Mellet, Longrée, 2009), constructed from the multidimensional tokens “sorte ... ADJ et ADJ” (“sort of” followed by an adjective, the coordinating conjunction “et”, and a second adjective). This motif combines several markers characteristic of Sand and is statistically confirmed (Fig. 16).

In this last example, the heuristic potential of deep learning to analyze textual data is impressive. The computational cost of identifying such a complex pattern using standard statistics is high. MCT offers an efficient approach to explore Maupassant's works by combining multiple text representations and utilizing both local and distant associations of tokens.

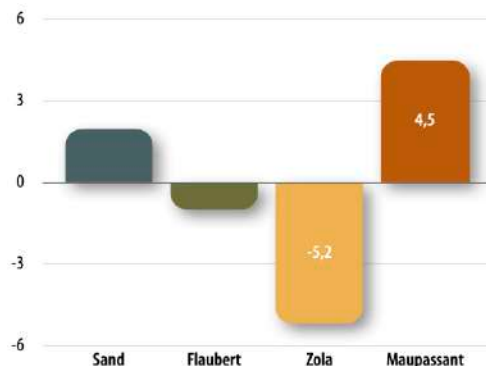


Fig. 16. Fisher exact test of the pattern "sorte ... ADJ et ADJ"

Conclusion

Corpus semantics or scientific text analysis require objectivity in their interpretations. Artificial intelligence and automatic key passage extraction can thus become invaluable tools, provided two conditions are met: enhancing the interpretability of models (making the algorithmic "black box" more transparent) and improving the linguistic explainability of model outputs (providing credible socio-linguistic explanations for the identified excerpts).

The MCT model enables the extraction of characteristic passages and the visualization of the textual units that are crucial for classification. The training is endogenous, i.e., performed on the corpus itself, respecting the genre or chronological peculiarities of the discourse. The idea of pre-training on heterogeneous data from the web seems reductive in the context of nuanced semantics.

By combining convolutional and transformer architectures, the model not only reveals key tokens or keywords, but also uncovers complex, contextual textual units. Instead of focusing solely on individual words (as in traditional statistical processing), we are interested in their short-distance combinations (collocations, n-grams, phrases) and long-distance co-occurrences (motifs, anaphoric repetitions). It is these associations, mediated by self-attention, that we see as the source of meaning and textual construction.

References

- Benzécri Jean-Paul (1973). *L'Analyse des Données*. Paris: Dunod.
- Feng Yue, Cheng Yan (2021). “Short text sentiment analysis based on multi-channel CNN with multi-head attention mechanism”. *IEEE Access*, 9, 19854–19863.
- Halliday Michael Alexander Kieckwood, Hasan Ruqaiya (1976). *Cohesion in English*. London: Longman.
- Joulin Armand, Grave Edouard, Bojanowski Piotr, Mikolov Tomas (2016). “Bag of Tricks for Efficient Text Classification”. Preprint arXiv:1607.01759.
- Kim Yoon (2014). “Convolutional Neural Networks for Sentence Classification”. In *Proceedings of the 2014 Conference on Empirical Methods in Natural Language (EMNLP)*. Qatar: Doha, 1746–1751.
- Kristeva Julia (1969). *Séméiotikè*. Paris: Éditions du Seuil.
- Lebart Ludovic, Salem André (1994). *Statistique Textuelle*. Paris: Dunod.
- Li Pengfei, Zhong Peixiang, Mao Kezhi, Wang Dongzhe, et alii (2021). “ACT: An Attentive Convolutional Transformer for Efficient Text Classification”. In: *Proceedings of the AAAI 2021*, 35(15), 13261–13269.
- Li Xuhong, Xiong Haoyi, Li Xingjian et alii (2022). “Interpretable deep learning: interpretation, interpretability, trustworthiness, and beyond”. *Knowl Inf Syst*, 64, 3197–3234.
- Mayaffre Damon (2021). *Macron ou le mystère du verbe*. La Tour d'Aigue: Éditions de l'Aube.
- Mayaffre Damon, Vanni Laurent (eds) (2021). *L'intelligence artificielle des textes. Des algorithmes à l'interprétation*. Paris: Honoré Champion.
- Mellet Sylvie, Longrée Dominique (2009). “Syntactical motifs and textual structures”. *Belgian Journal of Linguistics*, 23, 161–173.
- Rastier F. (2011). *La mesure et le grain. Sémantique de corpus*. Paris: Honoré Champion.
- Saussure Ferdinand (de) (1972). *Cours de linguistique. Générale*. Paris: Payot.
- Tognini-Bonell Elena (2002). “Corpus linguistics at work”. *Computational Linguistics*, 28, 583–583.
- Vanni Laurent (2024). “Hyperbase Web. (Hyper)Bases, Corpus, Langage”. *Corpus*, 2024, 25. file:///Users/racheleraus/Downloads/corpus-8770.pdf
- Vanni Laurent, Corneli Marco, Longree Dominique, Mayaffre Damon, Precioso Frédéric (2020). “Key passages : from statistics to deep learning”. In: *Text Analytics. Advances and Challenges*. Berlin: Springer, 41–54.

Key-passages and artificial intelligence.
How can we use hidden layers to explore new linguistic objects?

Vanni Laurent, Mahmoudi Hadi, Longrée Dominique, Mayaffre Damon (2024). “Multi-channel Convolutional Transformer and intertextuality: a Latin case study”. In: *New Frontiers in Textual Data Analysis*. Berlin: Springer, 197–207.

Vaswani Ashish, Shazeer Noam, Parmar Niki, Uszkoreit Jakib, Jones Llion, Gomez Aidan N., Kaiser Łukasz, Polosukhin Illia (2017). “Attention is all you need”. In: Curran Associates Inc., *Proceedings of the 31st International Conference on Neural Information (NIPS'17)*, USA, NY: Red Hook, 6000–6010.

L'intelligenza artificiale generativa al servizio della parità di genere: uno studio esplorativo sugli annunci di lavoro della Confederazione svizzera

Anna-Maria De Cesare, anna-maria.decesare@tu-dresden.de
Technische Universität Dresden

Introduzione

L'avvento dell'intelligenza artificiale generativa, in particolare dei modelli linguistici di grandi dimensioni, più noti come *Large Language Models* (LLM), apre il campo a numerose applicazioni "virtuose", tra cui la promozione di una comunicazione amministrativa più inclusiva¹. A nostra conoscenza, non disponiamo però ancora di uno studio sull'italiano che valuti le abilità dei LLM a generare testi amministrativi in lingua inclusiva fondato su una ricerca empirica sistematica. Un primo tentativo, il cui obiettivo consisteva nel "riformulare un testo italiano non rispettoso dell'accordo di genere", senza specificare una struttura *target*, è riportato brevemente in Raus (2023: 554):

In un test condotto il 10 giugno 2023, il risultato per la richiesta di riformulare un testo italiano non rispettoso dell'accordo di genere [...] ha dato luogo dapprima a una riformulazione identica e poi, a seguito della richiesta di rispondere nuovamente, alla riformulazione delle forme maschili con la doppia forma «maschile/femminile», che ha reso il testo illeggibile.

L'esito del test condotto da Raus, basato su due diversi *prompt* somministrati a un singolo LLM (ChatGPT-3.5), è giudicato dall'autrice stessa come poco soddisfacente. Il primo *prompt* restituisce, infatti, lo stesso testo di partenza (senza "accordo di genere"); il secondo, invece, genera testi in cui compaiono forme sdoppiate (non è però chiaro di quale tipo di sdoppiamento si tratta), che intaccano la leggibilità del testo. L'effetto negativo sulla deco-

¹ Per una descrizione delle varie accezioni che conosce il termine "lingua inclusiva", cfr. De Cesare, Giusti (2024: 3-8). Nel presente lavoro ci concentriamo sull'inclusione delle donne, il che significa prestare attenzione alle strategie linguistiche che le rendono visibili, come lo *splitting*, "sdoppiamento", studiato in modo approfondito anche in De Cesare (2022a e 2022b).

difica del testo potrebbe essere dovuto all'addensamento delle forme sdoppiate in un contesto ravvicinato (l'autrice non lo specifica e non fornisce esempi).

Alla luce delle abilità sempre più raffinate dell'intelligenza artificiale generativa, e tenendo conto delle tecniche note come *prompt engineering* (Chen *et al.*, 2023), riteniamo necessario sottoporre i LLM a un nuovo test per rispondere alla domanda generale seguente: i LLM possono essere uno strumento utile per chi è tenuto o semplicemente desidera scrivere testi inclusivi, che adottano pratiche linguistiche eque tra uomo e donna? La risposta verterà sui risultati ottenuti in uno studio esplorativo condotto su una specifica tipologia testuale: gli annunci di lavoro, con particolare attenzione a quelli redatti nell'ambito della pubblica amministrazione svizzera. L'attenzione è anche rivolta a una singola struttura *target*: gli sdoppiamenti contratti, che possono interessare un singolo sostantivo (*esperto/a*) o un sostantivo e i suoi target, vale a dire l'articolo e gli eventuali aggettivi che modificano il nome (*un/a collaboratore/trice scientifico/a*).

Il presente lavoro è organizzato come segue: nella prima parte descriviamo lo sdoppiamento contratto secondo le nuove linee guida emanate dalla Cancelleria Federale nel 2023, che lo ritengono adeguato ad attuare il pari trattamento linguistico tra donna e uomo negli annunci di lavoro (§ 1); riportiamo poi i risultati di un'indagine sulle forme linguistiche che compaiono in un campione di 180 annunci di lavoro pubblicati tra marzo e aprile 2024 nel portale "Posti vacanti" della Confederazione svizzera. Questa indagine mostra che le forme impiegate sono molto eterogenee e includono anche il maschile generico (§ 2). Nella parte principale del contributo (§ 3), descriviamo le forme linguistiche generate da nove diversi LLM con tre *prompt*, il cui principale *task* consiste nel riformulare 15 annunci di lavoro usando lo sdoppiamento contratto. Nelle conclusioni (§ 4) rispondiamo alla domanda formulata sopra, proponendo una valutazione generale dell'*output* dei nove LLM testati. Vedremo in particolare che i risultati migliori sono ottenuti con i modelli più grandi di OpenAI (ChatGPT-4o) e MistralAI (Mistral-Large) e che la presenza, nel *prompt*, di esempi illustrativi o ancora di una descrizione della struttura *target* da produrre non garantisce un *output* qualitativamente migliore.

1. Bandi di concorso e sdoppiamento contratto: le nuove raccomandazioni della Cancelleria federale svizzera

Il punto di partenza del presente studio sono le nuove raccomandazioni della *Guida all'uso inclusivo della lingua italiana nei testi della Confederazione* pubblicate dalla Cancelleria Federale nel 2023, che sostituiscono quelle proposte nel 2012. Ci interessano in particolare le strategie proposte per la redazione di una specifica tipologia testuale: gli *annunci di lavoro*, che la *Guida* chiama *bandi di concorso*. Il termine ricorre già nell'*Introduzione*, per chiarire che dal 1988 «[a]nche i bandi di concorso per posti nell'Amministrazione sono [...] formulati in modo tale da rivolgersi esplicitamente a candidati dei due sessi» (Cancelleria Federale, 2023: 3). Questa scelta è in linea con quanto mostrano numerosi studi di psicolinguistica (Bem, Bem, 1973; Born, Taris, 2010; Gaucher, Friesen, Kay, 2011): l'uso del maschile generico (nella *Guida* del 2023 chiamato "inclusivo") va evitato, pena una minore risposta dalla forza lavoro femminile.

Nei bandi di concorso che mirano a «rivolgersi esplicitamente ai candidati dei due sessi», la *Guida* presenta un'unica struttura linguistica. Nel § 4, relativo alle «strategie linguistiche per l'inclusione di genere», intendendo quello femminile², si propone infatti l'uso dello sdoppiamento contratto, la cui definizione va ricostruita dapprima in base alla descrizione formulata per lo sdoppiamento integrale: «Si parla di sdoppiamento integrale quando, nel riferirsi a una collettività mista, si abbinano un termine al maschile e un termine al femminile». Lo sdoppiamento integrale (che si usa quando ci si riferisce a un gruppo misto di persone) è illustrato con i tre esempi seguenti (qui e nella citazione che segue, abbiamo sottolineato le strutture in questione):

² Nello stesso paragrafo della *Guida*, il termine "inclusione di genere" è impiegato anche in un'altra accezione poiché si riferisce all'inclusione di «chi non si riconosce nel paradigma binario» (Cancelleria federale, 2023: 13). Più avanti, nel secondo punto della *Guida* che fa riferimento ai bandi di concorso (§ 5.1), per rivolgersi alle persone non binarie si sconsiglia l'uso dello sdoppiamento contratto, in particolare del doppio articolo prima di nomi detti "epiceni", che qui chiamiamo di genere comune (per dettagli, cfr. Fig. 2 del § 2.2). La *Guida* propone poi la seguente serie di esempi illustrativi: "Invece di: *L'Ufficio X cerca un/la: praticante presso [...]* Scrivere: *L'Ufficio X cerca: praticante presso [...]*" (*ibidem*: 23).

La Direzione ritiene importante che tutte le collaboratrici e tutti i collaboratori siano soddisfatti ...

Care concittadine, cari concittadini, è un onore ...

Gentili Signore e Signori, ...

(Guida 2023: 18).

Le specificità dello sdoppiamento in forma contratta, o sdoppiamento contratto, possono essere ricostruite in modo induttivo in base agli esempi proposti: gli elementi interessati sono il nome e i suoi *target* (per esempio l'articolo, anche nell'ambito di preposizioni articolate); il numero di sdoppiamenti in una frase può variare (possiamo anche avere due nomi sdoppiati); l'ordine degli elementi sdoppiati è «Maschile + Femminile» (e differisce dunque da quello dello sdoppiamento integrale); e compare l'artificio grafico della barra obliqua. I tre esempi illustrativi proposti dalla *Guida* permettono anche di osservare che lo sdoppiamento contratto si usa quando ci si riferisce a singoli individui generici:

Cerchiamo un/a traduttore/trice di lingua italiana

Nome del / della collaboratore/trice

Il / La candidato/a deve aver maturato lunga esperienza quale redattore/trice di testi ...

(Guida 2023: 19).

Stando alla *Guida*, «lo sdoppiamento contratto [...] è una soluzione grafica da usare, con prudenza, nei testi che vi si prestano, quali moduli, bandi di concorso [appunto], lettere standardizzate, e-mail e altre tipologie poco formali» (Cancelleria federale, 2023: 19). Chi decide di usare questa soluzione deve inoltre anche controllare un aspetto grafico-strutturale: ci vogliono gli spazi prima e dopo la barra se lo sdoppiamento interessa una struttura sintattica, come la preposizione («del / della»); non ci vogliono invece gli spazi se lo sdoppiamento riguarda una struttura che incide a livello morfologico, all'interno di una parola («un/a»).



Schweizerische Eidgenossenschaft
Confédération suisse
Confederazione Svizzera
Confederaziun Svizra

Panoramica Contatti Informazioni Ulteriori offerte di lavoro

L'Ufficio federale di statistica UST cerca un/a:
collaboratore/trice scientifico/a
80%-100% / Neuchâtel, Homeoffice

La statistica conta. Anche per voi.
Su incarico del Consiglio federale, l'UST condurrà una nuova indagine tra i commercianti di materie prime allo scopo di misurare il peso economico di questo settore nell'economia svizzera.

Siamo alla ricerca di un/una collaboratore/trice che sia responsabile della progettazione, della produzione, dell'analisi e della diffusione dei risultati dell'indagine tra i commercianti.

Le Sue mansioni

- Assumere la responsabilità di progettare e allestire la raccolta dei dati presso circa 400 imprese
- Assumere la responsabilità della plausibilizzazione, dell'analisi e del trattamento dei dati grezzi
- Curare i contatti con le imprese del settore e con l'associazione mantello per spiegare il progetto e ottenere sostegno a livello di contenuti
- Fornire supporto alle persone incaricate di raccogliere i dati, in particolare nei casi complessi
- Partecipare attivamente alla diffusione dei risultati

Il Suo profilo

- Diploma universitario in scienze economiche, finanza o formazione equivalente
- Ottime conoscenze in contabilità e informatica
- Capacità analitica e di sintesi; doti redazionali
- La conoscenza del settore del commercio di materie prime costituisce titolo preferenziale
- Ottime conoscenze di una seconda lingua ufficiale e dell'inglese; le principali lingue di lavoro sono il tedesco e il francese

Fig. 1. Esempio di annuncio di lavoro

2. Analisi di un campione di annunci di lavoro della Confederazione svizzera

2.1 Le forme linguistiche usate per riferirsi a singoli individui generici

Per capire come e dove i LLM potrebbero essere utilizzati da chi deve o vuole comunicare in modo equo, rivolgendosi ai due sessi, abbiamo dapprima dovuto tracciare un quadro delle forme che compaiono effettivamente negli annunci di lavoro redatti in italiano e pubblicati nel territorio elvetico³. A questo fine, abbiamo raccolto un campione di 180 annunci di lavoro

pubblicati nei mesi di marzo e aprile 2024 nel “Portale d’impiego della Confederazione”⁴.

Di questi annunci abbiamo analizzato tutte le forme riferite alla funzione che una determinata persona è chiamata a svolgere presso la Confederazione, concentrandoci unicamente sulla componente testuale principale dell’annuncio, scritta in grassetto e con caratteri più grandi rispetto al resto del testo. Nell’annuncio riprodotto nella

³ A nostra conoscenza gli studi condotti finora sugli annunci di lavoro redatti in italiano riguardano solo la realtà italiana. Questi studi mostrano tutti, naturalmente chi più chi meno, che gli annunci contengono molte forme declinate al maschile: in un campione di 748 annunci pubblicati nel 1984, Sabatini 1993 [1987] trova 34% di forme maschili e 43% di forme ambigue, mentre il campione di 320 annunci pubblicati nel 1998/1999 descritto in Olita 2006 ne contiene 42%. Anche lo studio di Nardone 2017, basato su un campione di 260 annunci pubblicati in italiano da 65 aziende, mostra che il maschile continua a dominare. Interessanti sono in particolare i risultati ottenuti per il sostantivo «ingegnere»: vi sono 28 occ. al maschile e nessuna al femminile. Come vedremo, la forma «ingegnera» compare invece negli annunci da noi analizzati. Questa differenza potrebbe costituire un’altra spia delle differenze tra l’italiano amministrativo d’Italia e l’italiano confederale svizzero (il termine “italiano confederale” è di Moretti 2005: 18).

⁴ Gli annunci sono stati raccolti il 20 aprile 2024. Il sito è disponibile al link: <https://www.stelle.admin.ch/stelle/it/home/stellen/stellenangebot.html>

Fig. 1, ad esempio, si riscontrano due elementi con sdoppiamento contratto: il nome («collaboratore/trice») e l'aggettivo che lo modifica («scientifico/a»). Si noti che la breve parte testuale che introduce il segmento che ci interessa (e corrisponde alla prima parte della frase, che compare in rosso nel testo originale) include in modo sistematico lo sdoppiamento di un altro *target* del sostantivo: l'articolo indefinito («un/a»).

La nostra analisi, di natura principalmente qualitativa, permette innanzitutto di osservare l'uso di nomi declinati al maschile generico. Si tratta da una parte di nomi che si riferiscono a ranghi elevati o cariche di prestigio⁵, per i quali il femminile esiste ma non è ancora entrato nell'uso («architetto TIC»; «capo⁶ Formazione professionale di base»); dall'altra, però, si trovano forme al maschile il cui femminile è all'ordine del giorno («collaboratore rimborso»; «psicologo presso il Centro di reclutamento Mels»).

Nel campione analizzato compaiono poi due categorie di nomi che non codificano il genere grammaticale: gli anglicismi (alcuni esempi sono riportati al punto 1) e i nomi di genere comune, «ossia invariabili al femminile e al maschile» (Cancelleria federale, 2023: 24), uscenti in *-ante*, *-ente*, *-ista* ecc. (cfr. «praticante», «inquirente», «specialista» e «giurista» al punto 2). In entrambi i casi compare talvolta l'abbreviazione «m/f», che permette di esplicitare il riferimento ai due sessi:

1. anglicismi: «Facility Manager (m/f)»; «controller»; «product owner Sistemi di trasmissione»;
2. nomi epiceni: «Praticante universitario nel campo del diritto pubblico (m/f)»; «Specialista (m/f) in rischi e gestione strategica»; «inquirente federale Impianti mobili»; «Giurista vigilanza sul trasporto in condotta».

Il campione analizzato include naturalmente anche molte forme sdoppiate, di solito in forma contratta (in linea con le raccomandazioni della *Guida*), ma non di rado anche in forma integrale. La scelta sembra per lo più casuale, come mostrano gli esempi incentrati sul lemma «collaboratore»:

⁵ Giusti propone in questo caso il termine “maschile di prestigio” (Giusti, 2022: 10).

⁶ Sull'uso del femminile «capa», che compare negli annunci analizzati (cfr. «Capo/a della divisione Digitalizzazione e risorse»), ci soffermiamo in De Cesare (in stampa).

3. «collaboratore/trice scientifico/a musica» (sdoppiamento contratto)
4. «collaboratore specializzato / collaboratrice specializzata» (sdoppiamento integrale).
5. «collaboratrice giuridica/collaboratore giuridico» (sdoppiamento integrale).

In entrambi i tipi di sdoppiamento si nota anche una certa libertà nell'ordine dei sostantivi. Negli sdoppiamenti contratti, l'ordine più frequente è «Masc+Fem» come per esempio in «ingegnere/a dei requisiti» e «pianificatore/trice Procedura d'asilo» (con i sostantivi che terminano in *-tore/-trice*, l'ordine è del resto sempre questo), ma nel campione compare talvolta anche l'ordine inverso: «ingegnera/e dei requisiti»; «sostituta/o del capo medico»; «stagiste/i giuridiche/ci dip. criminalità economica»⁷.

Un ultimo aspetto degno di nota concerne il numero di sdoppiamenti contratti riscontrati nella componente testuale analizzata e la complessità che ne deriva: nella maggior parte dei casi abbiamo un solo sdoppiamento, che interessa in generale il nome (come in «ingegnera/e cyber»), ma che può anche riguardare il solo aggettivo che lo qualifica («ingegnere specializzato/a»); non sono poi rari i casi con due sdoppiamenti, relativi al nome e all'aggettivo che lo modifica («collaboratore/trice scientifico/a musica»). Abbiamo anche riscontrato un caso con tre sdoppiamenti: «collaboratore/trice tecnico/a-scientifico/a in ispezioni sul campo», nel quale però la prima parte del composto aggettivale non dovrebbe essere sdoppiata (la forma corretta è *tecnico-scientifico/a*); l'errore potrebbe essere dovuto (almeno in parte) alla complessità della struttura in gioco.

2.2 Formulazione del *prompt*

Alla luce dell'eterogeneità delle forme usate negli annunci di lavoro pubblicati sul portale dei posti vacanti della Confederazione svizzera, è chiaro che sussiste un evidente margine di miglioramento. Per testare le abilità dei LLM in questo compito, abbiamo dapprima selezionato 15 annunci di lavoro (sui 180 analizzati) che includono una varietà di forme nominali e aggettivali. Di seguito la lista in questione (il grassetto è nostro):

⁷ Lo sdoppiamento di forme (sostantivi e/o aggettivi) al plurale compare una sola volta nel campione; questo dato di fatto non sorprende perché gli annunci di lavoro sono rivolti a singoli individui e non a gruppi di persone.

1. La Cancelleria federale CaF cerca un/a: **Esperto/a** in trasformazione digitale (strategia/modelli aziendali a livello di Confederazione)
2. Esercito svizzero - Comando Ciber Cdo Ciber cerca un/a: **Ingegnere/a** cyber (Cyber Electromagnetic Activities)
3. L'Amministrazione federale delle contribuzioni AFC cerca un/a: **collaboratore** rimborso
4. Esercito svizzero — Comando Ciber Cdo Ciber cerca un/a: **Architetto** TIC (Programma Air2030/Integrated Air Defence)
5. L'Ufficio federale dell'energia UFE cerca un/a: **Giurista** vigilanza sul trasporto in condotta
6. L'Ufficio federale di polizia fedpol cerca un/a: **inquirente federale** Impianti mobile
7. L'Ufficio federale della sicurezza alimentare e di veterinaria USAV cerca un/a: **Responsabile** Gestione degli eventi e delle situazioni di crisi
8. L'Ufficio federale della cultura UFC cerca un/a: **collaboratore/trice scientifico/a** musica
9. Agroscope cerca un/a: **Collaboratore/trice tecnico/a-scientifico/a** in ispezioni sul campo e test della capacità germinativa
10. Il Centro servizi informatici CSI-DFGP cerca un/a: **Ingegnere** DevOps **specializzato/a** in piattaforme Linux
11. L'Ufficio federale di meteorologia e climatologia MeteoSvizzera cerca un/a: **Praticante universitario** nel campo del diritto pubblico (m/f)
12. L'Ufficio federale delle costruzioni e della logistica UFCL cerca un/a: **capo** Formazione professionale di base
13. L'Ufficio federale della cultura UFC cerca un/a: **capo servizio specializzato** Comunicazione digitale
14. L'Esercito svizzero - Stato maggiore dell'esercito SM Es cerca un/a: **Sostituta/o del capo medico** al centro di reclutamento
15. L'Esercito svizzero - Base logistica dell'esercito BLEs cerca un/a: **Medico capo/a** del centro di reclutamento

Gli annunci selezionati includono casi con lo sdoppiamento contratto («esperto/a»; «ingegnere/a»), che sono dunque già declinati

in modo conforme alla tipologia testuale del bando di concorso (cfr. ess. 1 e 2). Oltre a questi casi, abbiamo incluso una serie di forme problematiche, sia per la presenza del maschile generico (cfr. ess. 3 e 4: «collaboratore» e «architetto») sia per la complessità delle strutture coinvolte (che interessano, o dovrebbero interessare, più di una forma: cfr. ess. 8, 9, 10, 11; abbiamo anche incluso la forma errata, per verificare se i LLM riscontrano la stessa difficoltà), sia ancora per la presenza del sostantivo «capo» (ess. 12 a 15).

La nostra lista include anche tre *filler* coincidenti con nomi di genere comune (cfr. ess. 5-7: «giurista», «inquirente federale» — si noti anche l'aggettivo invariabile — e «responsabile»), che non possono essere sdoppiati (gli esempi proposti nella Fig. 2 mostrano che lo sdoppiamento interessa in questo caso il solo determinante: «il / la cantante»; «il / la giornalista»). L'intento, anche qui, era quello di verificare quanto i LLM siano in grado di riconoscere quando lo sdoppiamento è escluso (come, appunto, nel caso dei nomi di genere comune) e quando è invece richiesto, come nel caso dei nomi simmetrici, per esempio «esperto/a», «collaboratore/trice» o ancora «ingegnere/a» (la cui forma femminile è più rara perché si preferisce ancora la forma che Giusti chiama “maschile di prestigio”; su questo punto, cfr. Giusti, 2022: 10).

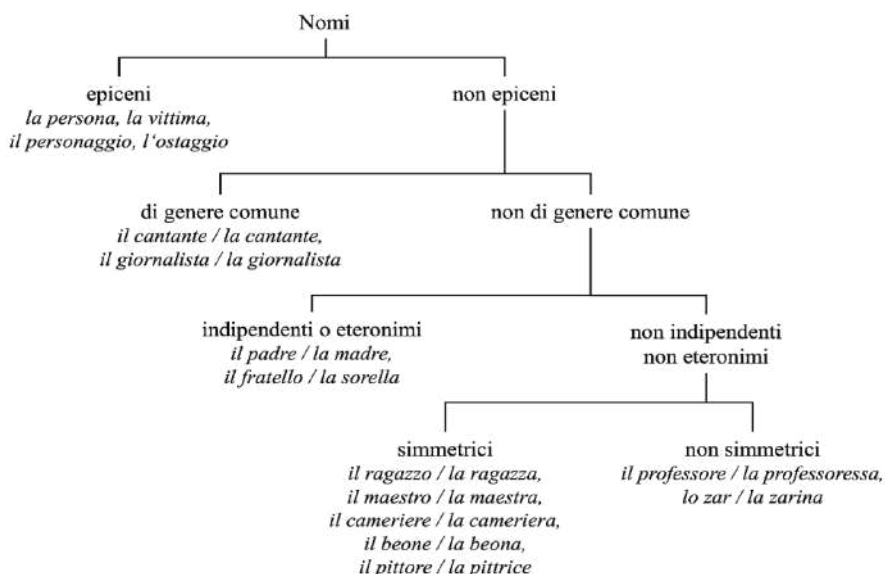


Fig. 2. Tipi di nomi in relazione al loro rapporto con il sesso del referente (Fonte: Thornton, 2022: 20)

Nei 15 annunci di lavoro selezionati, tutte le marche di forme femminili presenti nei nomi e relativi *target* (determinanti e aggettivi) sono poi state eliminate (anche nell'abbreviazione «m/f», presente nell'annuncio 11). Le forme che compaiono poi nei *prompt* coincidono dunque o con il maschile generico (quando sono nomi simmetrici) o con una forma che vale sia per il femminile sia per il maschile (come nel caso dei nomi di genere comune e degli aggettivi il cui genere grammaticale è invariabile, come «federale»):

1. La Cancelleria federale CaF cerca un: **Esperto**⁸ in trasformazione digitale (strategia/modelli aziendali a livello di Confederazione)
2. Esercito svizzero - Comando Ciber Cdo Ciber cerca un: **Ingegnere** cyber (Cyber Electromagnetic Activities)
3. L'Amministrazione federale delle contribuzioni AFC cerca un: **collaboratore** rimborso
4. Esercito svizzero - Comando Ciber Cdo Ciber cerca un: **Architetto** TIC (Programma Air2030/Integrated Air Defence)
5. L'Ufficio federale dell'energia UFE cerca un: **Giurista** vigilanza sul trasporto in condotta
6. L'Ufficio federale di polizia fedpol cerca un: **inquirente federale** Impianti mobili
7. L'Ufficio federale della sicurezza alimentare e di veterinaria USAV cerca un: **Responsabile** Gestione degli eventi e delle situazioni di crisi
8. L'Ufficio federale della cultura UFC cerca un: collaboratore scientifico musica
9. Agroscope cerca un: **Collaboratore tecnico-scientifico** in ispezioni sul campo e test della capacità germinativa
10. Il Centro servizi informatici CSI-DFGP cerca un: **Ingegnere** DevOps **specializzato** in piattaforme Linux

⁸ Nella selezione degli annunci non abbiamo tenuto conto della forma della prima lettera del sostantivo che segue i due punti: la lista include sostantivi che iniziano con maiuscola (come appunto nell'es. 1: «Esperto») e minuscola (come nell'es. 3: «collaboratore»). Questa differenza non ha naturalmente nessun influsso sulla presenza/assenza dello sdoppiamento (contratto).

11. L'Ufficio federale di meteorologia e climatologia MeteoSvizzera cerca un: **Praticante universitario** nel campo del diritto pubblico
12. L'Ufficio federale delle costruzioni e della logistica UFCL cerca un: **capo** Formazione professionale di base
13. L'Ufficio federale della cultura UFC cerca un: **capo servizio specializzato** Comunicazione digitale
14. L'Esercito svizzero - Stato maggiore dell'esercito SM Es cerca un: **Sostituto del capo medico** al centro di reclutamento
15. L'Esercito svizzero - Base logistica dell'esercito BLEs cerca un: **Medico capo** del centro di reclutamento

Veniamo infine ai *prompt* formulati per generare annunci con sdoppiamenti contratti. I *prompt* sono tre e si compongono di due componenti fisse: la formulazione del *task* («Effettua lo sdoppiamento contratto nelle 15 frasi seguenti/numerate prima») e la lista di annunci. Come mostrano i dati riportati nella Tab. 1, la lista di annunci e il *task* non compaiono sempre nello stesso ordine: nel *prompt* 1 compare dapprima il *task*, mentre nei *prompt* 2 e 3 a comparire per prima è la lista di annunci di lavoro. La differenza tra questi *prompt* sarà chiarita più avanti.

<i>Prompt</i>	Formulazione del <i>prompt</i>
<i>Prompt</i> 1 (zero-shot)	Effettua lo sdoppiamento contratto nelle 15 frasi seguenti: [segue la lista di 15 annunci di lavoro]
<i>Prompt</i> 2 (three-shot)	[Lista di 15 annunci di lavoro] Effettua lo sdoppiamento contratto nelle 15 frasi numerate date prima, sul modello degli esempi seguenti: L'Ufficio federale della dogana e della sicurezza dei confini UDSC cerca: specialisti/e dogana e sicurezza dei confini Esercito svizzero - Comando Ciber Cdo Ciber cerca un/a: Crittologo/a L'Ufficio federale dell'aviazione civile UFAC cerca un/a: Ispettore/trice di sicurezza delle informazioni
<i>Prompt</i> 3 (three-shot e spiegazione)	[Lista di 15 annunci di lavoro] Effettua lo sdoppiamento contratto nelle 15 frasi numerate date prima, sul modello degli esempi seguenti: [seguono i tre esempi del <i>prompt</i> 2] Si parla di sdoppiamento contratto quando, nel riferirsi a una collettività mista, si abbinano un termine al maschile e un termine al femminile. In questo caso si usa l'artificio grafico della barra obliqua. Si noterà che, come negli esempi riportati sopra, l'uso della barra presuppone un'attenzione particolare all'inserimento degli spazi, a seconda che la barra incida sul piano morfologico («un/a» senza spazi) o su quello sintattico («del / della» con spazi).

Tab. 1. I tre *prompt* formulati per generare annunci rispettosi della parità di genere

3. Analisi sperimentale

3.1 Lista di LLM testati, dati e metodi

Abbiamo somministrato i tre *prompt* a nove LLM, di cui si fornisce la lista completa nella Tab. 2. Al momento della somministrazione dei *prompt*, ovvero nel mese di maggio del 2024, la lista comprendeva i più grandi modelli commerciali disponibili⁹. Si tratta di LLM chiusi, per i quali non disponiamo di informazioni precise sui dati di addestramento (per dettagli sui modelli di OpenAI, cfr. De Cesare, 2024¹⁰). Nella colonna intermedia della Tab. 2 abbiamo specificato il paese in cui i rispettivi LLM sono stati creati, mentre in quella di destra forniamo dettagli in merito alle lingue supportate da ogni LLM. È chiaro che quelli che supportano *in primis* l'inglese (Intel e Meta) o sono sviluppati in Cina (Qwen) non sono i più adatti per capire un *task* formulato in italiano, che chiede di generare testi in lingua italiana. Tutti i LLM testati sono però in grado di produrre *output* anche in italiano e, in uno studio esplorativo come il nostro, è interessante capire la differenza qualitativa tra LLM sviluppati per lingue diverse.

I tre *prompt* somministrati ai nove LLM hanno permesso di generare un campione di 405 annunci di lavoro (il campione integrale è fornito in De Cesare, Weidensdorfer, Burchardt, 2024). Una prima analisi qualitativa dell'*output* mostra un quadro variegato di forme, che spaziano da sdoppiamenti contratti (in linea con la formulazione del *task*) a forme grammaticalmente inventate. Data l'eterogeneità dei risultati ottenuti, in ciò che se-

Nome dei LLM	Paese di origine	Lingue supportate
1. †ChatGPT-3.5	USA	Multilingue
2. ChatGPT-4	USA	Multilingue
3. ChatGPT-4o	USA	Multilingue
4. Mistral-Large	Francia	Multilingue
5. †Mistral-Next	Francia	Multilingue
6. Mistral-Small	Francia	Multilingue
7. Qwen 1.5 72B Chat	Cina	Multilingue
8. Intel Neural Chat 7B	USA	Inglese
9. Meta LLaMa 3 70B Instruct	USA	Inglese

Tab. 2. Lista di LLM testati

⁹ Il campo nel quale ci muoviamo è così dinamico che alcuni modelli non sono già più disponibili: è il caso di GPT-3.5 e Mistral-Next (nella Tab. 2 lo abbiamo segnalato con una croce).

¹⁰ I modelli di OpenAI sono allenati su dati redatti prevalentemente in inglese. Secondo Johnson *et al.*, 2022: 3, ChatGPT-3 (il predecessore di GPT-3.5) è stato addestrato su un campione di testi in cui il 93% è scritto in inglese, mentre solo lo 0,6% è redatto in italiano.

gue descriveremo le forme generate partendo dai *prompt*. Per motivi di spazio, ci soffermeremo solo sui sostantivi e aggettivi (con alcune osservazioni sull'articolo indefinito) che compaiono negli annunci 1 a 11, rimandando ad altra sede l'analisi relativa al sostantivo «capo» (annunci 12 a 15)¹¹. Il campione sul quale ci soffermiamo include dunque 297 risposte generate.

3.2 Risultati del *prompt* 1 (*zero-shot*)

Partiamo dai risultati ottenuti con il *prompt* 1, il più elementare dei tre *prompt* utilizzati: include infatti solo il *task* e l'elenco dei 15 annunci di lavoro sul quale svolgere il *task*. Dato che non include esempi o spiegazione, si tratta di uno *zero-shot prompt*.

Rispetto a quanto vedremo più avanti per gli altri due *prompt*, la descrizione dell'*output* è molto semplice perché le risposte generate da otto LLM su nove non contengono nessuna modifica: tutte le forme che compaiono nel *prompt* sono riproposte tali e quali nell'*output*. Troviamo dunque o il maschile generico nel caso dei nomi simmetrici (per es. «esperto», «architetto») o un sostantivo di genere comune («giurista», «inquirente», «responsabile»). Lo stesso vale per gli aggettivi («collaboratore scientifico»; «collaboratore tecnico-scientifico»; «ingegnere specializzato»; «praticante universitario»).

Fa eccezione l'*output* generato da ChatGPT-4, che riproduciamo integralmente di seguito (il grassetto è nostro¹²):

1. La Cancelleria federale CaF cerca una/un: **Esperta/o** in trasformazione digitale (strategia/modelli aziendali a livello di Confederazione)
2. Esercito svizzero - Comando Ciber Cdo Ciber cerca una/un: **Ingegnere/a** cyber (Cyber Electromagnetic Activities)
3. L'Amministrazione federale delle contribuzioni AFC cerca una/un: **collaboratrice/tore** rimborso

¹¹ Cfr. De Cesare, in stampa.

¹² Per il resto (a parte il grassetto), riproduciamo l'*output* generato dai LLM senza modificare il testo. La sigla P1 si riferisce al numero del *prompt* (dunque: «*Prompt* 1») e la cifra che segue (per esempio 1 in P1.1) al numero dell'annuncio nella lista.

4. Esercito svizzero - Comando Ciber Cdo Ciber cerca una/un: **Architetta/o** TIC (Programma Air2030/Integrated Air Defence)
5. L'Ufficio federale dell'energia UFE cerca una/un: **Giurista** di vigilanza sul trasporto in condotta
6. L'Ufficio federale di polizia fedpol cerca una/un: **inquirente federale** Impianti mobili
7. L'Ufficio federale della sicurezza alimentare e di veterinaria USAV cerca una/un: **Responsabile** della Gestione degli eventi e delle situazioni di crisi
8. L'Ufficio federale della cultura UFC cerca una/un: **collaboratrice/tore scientifico/a** musica
9. Agroscope cerca una/un: **Collaboratrice/tore tecnico-scientifico/a** in ispezioni sul campo e test della capacità germinativa
10. Il Centro servizi informatici CSI-DFGP cerca una/un: **Ingegnere/a** DevOps **specializzato/a** in piattaforme Linux
11. L'Ufficio federale di meteorologia e climatologia MeteoSvizzera cerca una/un: **Praticante universitaria/o** nel campo del diritto pubblico

ChatGPT-4 ha effettuato il *task* in modo molto soddisfacente: le risposte generate includono sdoppiamenti contratti (d'ora in poi: SC) per tutti i nomi simmetrici, anche per quelli che denotano cariche prestigiose (il LLM non è naturalmente in grado di fare la differenza), come «ingegnere/a» (in P1.2) e «architetta/o» (in P1.4), mentre i sostantivi di genere comune e gli aggettivi invariabili in merito al genere grammaticale — giustamente — sono riproposti tali e quali (cfr. «Giurista» in P1.5, «inquirente federale» in P1.6, «responsabile» in P1.7 e «Praticante» in P1.11). Le risposte contengono anche lo sdoppiamento dell'articolo indefinito (collocato prima dei due punti): si tratta però sempre di sdoppiamento integrale.

Nelle risposte generate con il *prompt* 1 spicca però anche un aspetto meno soddisfacente e poco prevedibile, relativo all'ordine dei due elementi che entrano nello sdoppiamento contratto. L'ordine «Masc+Fem» — che era quello più atteso — compare in soli due casi: nel già citato P1.2 e in P1.10: «ingegnere/a» e «ingegnere/a specializzato/a». In tutti gli altri annunci, ad eccezione di due altri casi che descriviamo di seguito, il femminile precede inve-

ce il maschile¹³, perfino nei casi in cui il sostantivo esce con il suffisso *-tore/-trice* (come in P1.3: «collaboratrice/tore»). L'ordine «Fem+Masc» compare in modo sistematico anche negli sdoppiamenti integrali dell'articolo indefinito («una/un»), compreso nei casi in cui la testa del sintagma, che segue i due punti, allinea l'ordine «Masc+Fem» (cfr. P1.2: «una/un: Ingegnere/a cyber»).

Le due uniche strutture che non seguono l'ordine morfosintattico dominante («Fem+Masc») sono sdoppiamenti con ordini misti, che sono poco naturali: in queste strutture si osserva che i sostantivi continuano a presentare l'ordine «Fem+Masc», mentre gli aggettivi seguono l'ordine inverso: «collaboratrice/tore scientifico/a» (P1.8) e «Collaboratrice/tore tecnico-scientifico/a» (P1.9). Merita poi un commento speciale la forma dell'aggettivo composto che compare nel secondo esempio: qui, in effetti, è sdoppiata solo la seconda parte dell'aggettivo. La forma generata risulta dunque migliore della forma originale (errata) prodotta da un essere umano (cfr. «tecnico/a-scientifico/a» dell'es. 9 proposto nel § 2.2).

3.3 Risultati del prompt 2 (three-shot)

Passiamo all'*output* generato dal *prompt* 2, più specifico del primo perché include anche (oltre alla lista di 15 annunci di lavoro e la formulazione del *task*) tre esempi per illustrare cosa si intende con “sdoppiamento contratto”: si tratta dunque di un *three-shot prompt*. Va notato che in questi esempi l'ordine degli elementi sdoppiati è sempre «Masc+Fem».

La prima cosa da osservare è che la presenza degli esempi migliora in modo notevole i risultati ottenuti (rispetto al *prompt* 1). L'*output* del *prompt* 2 è infatti speculare a quello del primo: otto LLM su nove generano annunci che includono strutture con sdoppiamenti contratti. L'unica eccezione è Intel Neural Chat 7B, che genera solo risposte in inglese. Nella Tab. 3 riportiamo la lista di tutte le forme nominali e aggettivali generate, ordinate per annuncio e LLM (senza tenere conto delle risposte generate in inglese da Intel).

¹³ Cfr. P1.1: «esperta/o» e P1.4: «architetta/o»; va menzionato anche P1.11, in cui è sdoppiato solo l'aggettivo: «Praticante universitaria/o».

La maggior parte delle risposte generate dai modelli di OpenAI (ChatGPT-3.5, 4 e 4o) e da quelli di Mistral (Small, Next e Large) contiene uno sdoppiamento contratto corretto dei nomi simmetrici e dei relativi target (cfr. P2.1: «Esperto/a»; P2.2: «Ingegnere/a»; P2.3: «collabo-

ratore/trice»; P2.4: «Architetto/a»; P2.8: «collaboratore/trice scientifico/a»; P2.10: «Ingegnere/a DevOps specializzato/a»). Nelle altre risposte di questi stessi modelli, lo sdoppiamento non c'è, in linea con il fatto che il sostantivo è di genere comune¹⁴. Ci sono però anche due casi, generati da Mistral-Next, in cui manca lo sdoppiamento sul sostantivo «ingegnere» (P2.2 e P2.10).

Gli altri due modelli (Qwen e Meta) generano forme sdoppiate grammaticalmente corrette solo per i sostantivi simmetrici uscanti in *-o* (P2.1: «Esperto/a»; P2.4: «Architetto/a») e in *-e* (P2.2: «Ingegnere/a»; P2.10: «Ingegnere/a DevOps specializzato/a»). Negli altri casi in cui compare lo sdoppiamento contratto, molte risposte proposte sono anomale per il femminile, sia nel caso di sostantivi simmetrici uscanti in *-tore* (P2.3, P2.8, P2.9: «collaboratore») sia in quello di sostantivi di genere comune (P2.5 a P2.7 e P2.11: «Giurista», «Inquirente», «Responsabile», «Praticante»). Le forme femminili generate per queste due categorie di sostantivi escono infatti tutte con il morfema *-a*, che viene applicato ai sostantivi che o richiedono un morfema femminile specifico, come nel caso dei nomi che escono in *-trice* (cfr. P2.3: «collaboratore/a»; idem in P2.8 e P2.9), o non conoscono il femminile (P2. 5: «Giurista/a»; P2.6: «Inquirente/a»; P2.7: «Responsabile/a»; P2.11: «Praticante/a universitario/a»). La sovraestensione del morfema femminile in *-a* — che funge da marca femminile 'a tutto fare' — si riscontra anche in una

1. Esperto/a (8 LLM)
2. Ingegnere/a (7 LLM); Ingegnere (Mistral-Next)
3. collaboratore/trice (6 LLM); collaboratore/a (Qwen; Meta)
4. Architetto/a (8 LLM)
5. Giurista (6 LLM); Giurista/a (Mistral-Small; Meta)
6. inquirente federale (6 LLM); inquirente federale/a (ChatGPT-4); Inquirente/a federale (Meta)
7. Responsabile (6 LLM); Responsabile/a (Qwen; Meta)
8. collaboratore/trice scientifico/a (6 LLM); collaboratore/a scientifico/a (Qwen; Meta)
9. collaboratore/trice tecnico-scientifico/a (6 LLM); collaboratore/a tecnico-scientifico/a (Qwen; Meta)
10. Ingegnere/a DevOps specializzato/a (7 LLM); Ingegnere DevOps specializzato/a (Mistral-Next)
11. Praticante universitario/a (7 LLM); Praticante/a universitario/a (Meta)

Tab. 3. Risposte generate con il *prompt* 2 (con focus su sostantivi e aggettivi)

¹⁴ Cfr. P2.5: «Giurista»; P2.6: «inquirente federale»; P2.7: «Responsabile».

risposta di Mistral-Small (P2.5: Giurista/a) e in una risposta di ChatGPT-4 (cfr. P2.6: «inquirente federale/a»; in questo caso la si riscontra sul solo aggettivo; si tratta di un errore che non si riscontra nelle risposte di Qwen e Meta).

Nelle risposte generate *in primis* da Qwen e Meta, si osserva che il *task* è stato eseguito (gli sdoppiamenti contratti ci sono), ma con forme non corrette. Gli sdoppiamenti generati possono perciò essere globalmente valutati come “falliti” o perché si effettua lo sdoppiamento su sostantivi che non lo richiedono, con morfemi che non permettono di distinguere tra maschile e femminile (come nel caso dei nomi di genere comune «Giurista/a») o perché la forma femminile corretta è un'altra (come nel caso dei nomi simmetrici: «collaboratore/a»). È chiaro che il suffisso *-trice* è meno frequente e più complesso della desinenza in *-a* ma è comunque ben attestato in italiano. I modelli con i risultati peggiori sono, non a caso, quelli addestrati su dati in cui l'italiano non è presente (o è molto scarso).

3.4 Risultati del prompt 3 (three-shot e spiegazione)

Veniamo infine al *prompt* 3, quello più specifico, che fornisce anche la descrizione dello sdoppiamento contratto proposta nella *Guida all'uso inclusivo della lingua italiana* (Cancelleria Federale, 2023). Sorprendentemente, i risultati ottenuti con il *prompt* 3 sono ancora più eterogenei di quelli generati con il *prompt* 2 e dunque anche più lontani dal *target* di arrivo previsto. Oltre agli sdoppiamenti contratti, compaiono due nuove tipologie di forme: lo sdoppiamento integrale e il modificatore «donna». Quest'ultimo è presente solo nelle risposte di Intel Neural Chat 7B, che tornano ad essere generate in lingua italiana. La lista di forme, anche qui ordinata per annuncio e LLM, è riportata nella Tab. 4.

Le forme non sdoppiate sono numericamente contenute e riguardano, giustamente, solo i nomi di genere comune. I modelli che generano il numero più elevato di forme corrette relative ai nomi di genere comune sono quelli di OpenAI e di Mistral (P3.5: «Giurista»; P3.6: «inquirente federale»; P3.7: «Responsabile»).

Molto più complessa da descrivere è la casistica degli sdoppiamenti, che presenta tuttavia invariabilmente l'ordine «Masc+Fem», a prescin-

dere dalla forma contratta o integrale (cfr., per esempio, P3.1: «Esperto/a» e «Esperto / Esperta»).

Gli sdoppiamenti contratti sono numerosi, ma non sempre corretti. A generare forme corrette sono di nuovo perlopiù i modelli di OpenAI e di Mistral (P3.1: «Esperto/a»; P3.2: «Ingegnere/a»; P3.3: «collaboratore/trice»; P3.4: «Architetto/a»; P3.8: «collaboratore/trice scientifico/a»; P3.9: «collaboratore/trice tecnico-scientifico/a»; P3.10: «Ingegnere/a DevOps specializzato/a»; P3.11: «Praticante universitario/a»).

E anche qui, come nelle risposte relative al *prompt* 2 già analizzate, le forme che pongono meno problemi sono i sostantivi simmetrici che escono in -o (otto LLM su nove generano uno sdoppiamento contratto corretto di «esperto» e sette su nove di «architetto»). In alcuni casi compaiono poi sdoppiamenti contratti parziali, relativi al solo aggettivo, sempre uscente in -o (cfr. P3.10 di Mistral-Large: «Ingegnere DevOps specializzato/a»).

A queste due tipologie di sdoppiamenti (quelli contratti completi e contratti parziali) si aggiungono gli sdoppiamenti contratti grammaticalmente incorretti, in cui compare un morfema inventato (in genere chiaramente interpretabile come femminile): i casi in questione riguardano sia i nomi simmetrici (riscontrabili in alcune risposte di Qwen, ovvero P3.2: «Ingegnere/e»; P3.3: «collaboratore/a»; P3.8: «Collaboratore/a scientifico/a» e P3.9: «Collaboratore/a tecnico-scientifico») sia i nomi di genere comune (cfr. P3.5: «Giurista/a» generato da Mistral-Next, Mistral-Small e Meta; P3.7: «Responsabile/a» di Qwen e Meta; P3.11: «Praticante/a universitario/a» di Meta).

1. Esperto/a (8 LLM); Esperto / Esperta (Intel)
2. Ingegnere/a (6 LLM); Ingegnere/Ingegnera (ChatGPT-3.5); Ingegnere cyber / Ingegnere cyber donna (Intel); Ingegnere/e (Qwen)
3. collaboratore/trice (6 LLM); collaboratore/collaboratrice (ChatGPT-3.5, Intel); collaboratore/a (Qwen)
4. Architetto/a (7 LLM); Architetto/architetta (ChatGPT-3.5); Architetto / Architetto (Intel)
5. Giurista (5 LLM); Giurista/a (Mistral-Next e Small; Meta); Giurista vigilanza sul trasporto in condotta / Giurista donna vigilanza sul trasporto in condotta (Intel)
6. inquirente federale (6 LLM); inquirente federale/federale (ChatGPT-4); inquirente/inquirente federale (Mistral-Next); inquirente federale Impianti mobili / inquirente federale donna Impianti mobili (Intel)
7. Responsabile (5 LLM); Responsabile/a (Qwen; Meta); Responsabile/Responsabile (Mistral-Next); Responsabile Gestione degli eventi e delle situazioni di crisi / Responsabile Gestione degli eventi e delle situazioni di crisi donna (Intel)
8. collaboratore/trice scientifico/a (6 LLM); collaboratore/collaboratrice scientifico/a (ChatGPT-3.5); Collaboratore/a scientifico/a (Qwen); collaboratore scientifico musica / collaboratrice scientifico musica (Intel)
9. collaboratore/trice tecnico-scientifico/a (6 LLM); Collaboratore/collaboratrice tecnico-scientifico/a (ChatGPT-3.5); Collaboratore/a tecnico-scientifico (Qwen); Collaboratore tecnico-scientifico / Collaboratrice tecnico-scientifica (Intel)
10. Ingegnere/a DevOps specializzato/a (6 LLM); Ingegnere/Ingegnera specializzato/a (ChatGPT-3.5); Ingegnere DevOps specializzato/a (Mistral-Large); Ingegnere DevOps specializzato / Ingegnere DevOps donna specializzato (Intel)
11. Praticante universitario/a (6 LLM); Praticante universitario/universitaria (Mistral-Next); Praticante universitario / Praticante universitaria (Intel); Praticante/a universitario/a (Meta)

Tab. 4. Risposte generate con il *prompt* 3 (con focus su sostantivi e aggettivi)

Come già indicato sopra, nelle risposte generate con il *prompt* 3 fanno poi la loro comparsa forme sdoppiate in modo integrale. Queste strutture sono generate da due modelli: ChatGPT-3.5 (cfr. P3.2: «Ingegnere/Ingegnera»; P3.3: «collaboratore/collaboratrice»; P3.4: «Architetto/architetta») e Intel (P3.1: «Esperto / Esperta»; P3.9: «Collaboratore tecnico-scientifico / Collaboratrice tecnico-scientifica»; P3.11: «Praticante universitario / Praticante universitaria»). Una forma relativa al solo aggettivo è generata anche da Mistral-Next (P3.11: «Praticante universitario/universitaria»).

Un'ultima casistica di sdoppiamenti è costituita dalle strutture ibride, che mischiano forme integrali (invariabilmente il sostantivo) e contratte (l'aggettivo che lo modifica). Queste forme sono generate unicamente da ChatGPT-3.5 (P3.8: «collaboratore/collaboratrice scientifico/a»; P3.9: «Collaboratore/collaboratrice tecnico-scientifico/a»; P3.10: «Ingegnere/ingegnera specializzato/a»).

Passiamo ora alla casistica delle risposte non soddisfacenti, anche in termini di accettabilità grammaticale. Tra queste vanno innanzitutto annoverate le strutture in cui è sdoppiato integralmente il solo sostantivo, mentre l'aggettivo rimane al maschile, anche con il nome femminile (P3.8 di Intel: «collaboratore scientifico musica / collaboratrice scientifico musica»). Vi sono poi strutture in cui lo sdoppiamento consiste in una semplice ripetizione o del sostantivo al maschile (P3.4 di Intel: «Architetto / Architetto»), o del sostantivo di genere comune (cfr. P3.6 e P3.7 di Mistral-Next: «inquirente/inquirente federale» e «Responsabile/Responsabile») oppure ancora dell'aggettivo (P3.6: «inquirente federale/federale»).

Poco soddisfacenti, ma questa volta per motivi di natura culturale e strutturale¹⁵, risultano infine le risposte che includono il sostantivo «donna» in funzione di modificatore di un altro sostantivo, in una struttura composta sul modello «X + donna». Si tratta di un

¹⁵ La *Guida all'uso inclusivo della lingua italiana* della Confederazione (2023) consiglia di evitare questa strategia linguistica («Invece di: *Il municipio era rappresentato dal sindaco donna Angela Marchesi*. Scrivere: *Il municipio era rappresentato dalla sindaca Angela Marchesi*», Cancelleria federale, 2023: 16), tranne che negli «usi propri dell'ambito militare» (*ibidem*: 24), come in *il soldato donna, il capitano donna, il maggiore donna* (*ibidem*: 29). I motivi per cui il sostantivo «donna» dovrebbe essere evitato sono i seguenti: «La strategia che accosta il nome maschile di professione e il sostantivo 'donna' (es. 'giudice donna', 'donna medico') è sconsigliabile, da una parte per il suo carattere eccentrico — un equivalente al maschile sul modello di 'casalingo uomo' sarebbe impensabile — e dall'altra perché moltiplica inutilmente le forme» (*ibidem*: 24).

espediente tipico della lingua inglese (i cui sostantivi sono per lo più privi di genere grammaticale), che esplicita lessicalmente il genere biologico (femminile) del referente. Non a caso tutti gli esempi riscontrati nell'*output* del *prompt* 3 sono generati da Intel, che supporta *output in primis* in lingua inglese. Osservando più da vicino gli esempi in questione, è poi interessante notare che i nomi modificati dal sostantivo «donna» appartengono a due gruppi: oltre ai nomi di genere comune (cfr. P3.5: «Giurista vigilanza sul trasporto in condotta / Giurista donna vigilanza sul trasporto in condotta»; P3.6: «inquirente federale Impianti mobili / inquirente federale donna Impianti mobili»), si trova anche il nome simmetrico «ingegnere» (P3.2: «Ingegnere cyber / Ingegnere cyber donna»; P3.10: «Ingegnere DevOps specializzato / Ingegnere DevOps donna specializzato»). In alcuni casi risulta poco convincente anche la posizione distante del modificatore «donna» rispetto alla testa (P3.7: «Responsabile Gestione degli eventi e delle situazioni di crisi / Responsabile Gestione degli eventi e delle situazioni di crisi donna»).

Conclusioni

Alla luce dei risultati ottenuti in questo lavoro, basato su un'analisi esplorativa di natura prevalentemente descrittiva e qualitativa, possiamo fornire una prima risposta generale alla domanda formulata nell'introduzione: i LLM possono essere uno strumento utile per formulare annunci di lavoro rispettosi della parità linguistica tra donna e uomo; bisogna però scegliere il modello giusto e individuare il *prompt* che genera i risultati più vicini alla struttura *target*, nel nostro caso coincidente con lo sdoppiamento contratto.

Per quanto riguarda la scelta del *prompt*, il risultato più interessante emerso — anche perché sorprendente — è che la qualità delle risposte non è direttamente correlata al grado di specificità del testo usato nella consegna: le risposte più soddisfacenti sono infatti ottenute con il *prompt* 2 (che propone la lista di annunci da riformulare in chiave inclusiva, il *task* e tre esempi che illustrano la struttura *target*), più specifico del *prompt* 1 (che propone unicamente il *task* e gli annunci da modificare), ma meno specifico del *prompt* 3 (che include anche una definizione della struttura *target*). La generale bassa qualità dei risultati ottenuti con il *prompt* 1 (ricordiamo che otto LLM

su nove restituiscono tutte le forme declinate al maschile) non sorprende più di tanto. Molto meno scontata, e difficile da spiegare, risulta invece la più scarsa qualità dei risultati ottenuti con il *prompt* 3 rispetto al *prompt* 2: il *prompt* 3 genera una casistica molto ampia di forme sdoppiate (le risposte includono forme con sdoppiamento contratto, esteso, parziale, ibrido) e sintagmi complessi la cui testa nominale è modificata dal sostantivo «donna», mentre il *prompt* 2 genera solo e unicamente forme sdoppiate in modo contratto (che non sono però sempre grammaticalmente corrette).

La qualità delle risposte generate varia in modo importante anche in base al tipo di LLM utilizzato, in primo luogo — chiaramente — perché non supportano le stesse lingue nell'*output*, il che dipende a sua volta dai dati di addestramento. I modelli migliori, che generano gli annunci di lavoro più conformi alla formulazione del *task*, sono ChatGPT-4o di OpenAI e Mistral-Large di MistralAI (le risposte generate con i *prompt* 2 e 3 sono quasi tutte corrette). Dà molte risposte corrette anche ChatGPT-4, mentre sono molto meno convincenti i risultati di ChatGPT-3.5 e di due LLM di MistralAI (Mistral-Next e soprattutto Mistral-Small).

I risultati peggiori, che includono forme poco soddisfacenti, lontane dalle strutture *target* o addirittura agrammaticali (inventate), sono generati dai LLM che supportano soprattutto *output* in inglese (cfr. Intel, che genera del resto anche risposte integralmente in inglese, e Meta) o anche in cinese (Qwen). Chi lavora con l'italiano non avrà dunque nessun interesse a far ricorso a questi modelli, anche se sono in grado di generare *output* in lingua italiana. L'*output*, infatti, contiene forti impronte della lingua inglese, come l'impiego del modificatore «donna».

Due parole, per concludere, sui passi da compiere per approfondire questa ricerca. Una prima cosa importante da fare sarà naturalmente testare l'affidabilità dei risultati ottenuti in questa sede, replicando l'analisi con nuovi dati generati dalla stessa batteria di LLM. Bisognerà poi ampliare l'orizzonte, tenendo conto di altre tipologie testuali e di altre strutture linguistiche, e rimanere al passo con i tempi, testando tutti i nuovi LLM multilingui rilasciati (per esempio Mistral NeMo) e noti per la qualità delle loro risposte (come Claude di Anthropic). Di assoluta centralità saranno le analisi che testano le abilità dei LLM addestrati con (più) dati in lingua italiana, come LLaMAntino (per dettagli, cfr. Basile *et al.*, 2024) — al-

lenato però anche su dati tradotti dall'inglese o addirittura generati in inglese e poi tradotti in italiano — e il modello IT5 (descritto in Sarti, Nissim, 2024), i cui dati di allenamento sono stati scelti con grande cura e i cui risultati dovrebbero dunque essere ancora migliori di quelli ottenuti nel presente studio con ChatGPT-4o e Mistral-Large.

In aggiunta al loro apporto descrittivo, centrale in particolare in ambito applicativo, i risultati ottenuti in uno studio esplorativo come il nostro si rivelano anche preziosi per svolgere una riflessione di natura teorico-concettuale. Data l'importanza, sia qualitativa sia quantitativa, delle forme errate (spesso inventate) generate dai LLM, un altro passo da compiere nella ricerca intrapresa riguarda il modo migliore — anche in chiave di lingua non discriminatoria — di riferirsi alle forme in questione. Una prima soluzione potrebbe consistere nel chiamarle “allucinazioni grammaticali” o, visto che si tratta perlopiù di morfemi (cfr., tra molti altri, «collaboratore/a», «responsabile/a», «inquirente federale/a»), “allucinazioni morfologiche”. Se però vogliamo evitare il riferimento al concetto di “allucinazione”, che crea un'immagine distorta degli strumenti che stiamo testando (oltre alla problematica antropomorfizzazione, in psichiatria il concetto riguarda una patologia lontana dai fenomeni osservati nell'*output* dei LLM; per approfondimenti, cfr. Gerstenberg, 2024), si potrebbe parlare di “invenzioni algoritmiche” e più precisamente di “morfemi artificiali”, inesistenti nell'italiano naturale.

Bibliografia

Basile Pierpaolo, Musacchio Elio *et alii* (2024). “LLaMAntino: Large Language Models per la lingua italiana”. Ital-IA 2024: 4th National Conference on Artificial Intelligence, organized by CINI, May 29-30, 2024, Naples, Italy.

Bem Sandra L., Bem Daryl J. (1973). “Does Sex-biased Job Advertising “Aid and Abet” Sex Discrimination?”. *Journal of Applied Social Psychology*, 3 (1), 6-18.

Born Marise Ph., Taris Toon W. (2010). “The impact of the wording of employment advertisements on students' inclination to apply for a job”. *The Journal of Social Psychology*, vol. 150, n. 5, 485–502. <https://doi.org/10.1080/00224540903365422>

Cancelleria Federale (2012). *Pari trattamento linguistico. Guida al pari trattamento linguistico di donna e uomo nei testi ufficiali della Confederazione*, Berna, Cancelleria Federale.

Cancelleria Federale (2023). *Linguaggio inclusivo di genere. Guida all'uso inclusivo della lingua italiana nei testi della Confederazione*, Berna, Cancelleria Federale. Guida al linguaggio inclusivo di genere (admin.ch)

Chen Banghao, Zhang Zhaofeng *et alii* (2023). “Unleashing the potential of prompt engineering in Large Language Models: a comprehensive review”. arXiv preprint arXiv:2310.14735.

De Cesare Anna-Maria (2022a). “La codifica linguistica dei referenti umani nella Costituzione svizzera: tra disparità e uguaglianza di genere”. In: Ferrari Angela, Lala Letizia, Pecorari Filippo (a cura di). *L'italiano dei testi costituzionali. Indagini linguistiche e testuali tra Svizzera e Italia*. Alessandria: Edizioni dell'Orso, 245-270.

De Cesare Anna-Maria (2022b). “Sdoppiamenti nelle carte costituzionali: tra italiano federale e cantonale”. In: Ferrari Angela, Lala Letizia, Pecorari Filippo (a cura di). *L'italiano dei testi costituzionali. Indagini linguistiche e testuali tra Svizzera e Italia*. Alessandria: Edizioni dell'Orso, 483-498.

De Cesare Anna-Maria (2024). “Nuove dinamiche di contatto linguistico: Le “impronte digitali” dell'inglese nell'italiano generato da LLM. *Lid'O. Lingua Italiana d'Oggi* XXI: 67-91.

De Cesare Anna-Maria (in stampa). “Per un'amministrazione impegnata e aggiornata: Come formulare annunci di lavoro rispettosi della parità di genere con l'intelligenza artificiale generativa?”. In: Fiorentino Giuliana (a cura di). *Amministrazione attiva: semplicità e chiarezza per la comunicazione amministrativa*. Firenze: Cesati.

De Cesare Anna-Maria, Giusti Giuliana (2024). “Lingua inclusiva e variazione linguistica”. In: De Cesare Anna-Maria, Giusti Giuliana (a cura di). *Lingua inclusiva: forme, funzioni, atteggiamenti e percezioni*. LiVVaL *Linguaggio e Variazione / Variation in Language* 6, 3-17. <http://doi.org/10.30687/978-88-6969-866-8/001>

De Cesare Anna-Maria, Weidensdorfer Tom, Burchardt Luisa (2024). “Sdoppiamento contratto negli annunci di lavoro: lista di prompt, *large language model* e risposte generate”. Zenodo. <https://doi.org/10.5281/zenodo.14566338>

Gaucher Danielle, Friesen Justin, Kay Aaron C. (2011). “Evidence that gendered wording in job advertisements exists and sustains gender inequality”. *Journal of Personality and Social Psychology*, 101 (1), 109–128. <https://doi.org/10.1037/a0022530>

Gerstenberg Annette (2024). “Hallucinations in automated texts – A critical view on the emerging terminology”. *AI-Linguistica. Linguistic Studies on AI-Generated Texts and Discourses*, 1 (1). <https://doi.org/10.62408/ai-ling.v1i1.9>

Giusti Giuliana (2022). “Inclusività della lingua italiana, nella lingua italiana: come e perché. Fondamenti teorici e proposte operative”. *Deportate, Esuli, Profughe. Rivista telematica di studi sulla memoria femminile*, 48, 1-19.

Johnson Rebecca L., Pistilli Giada et alii (2022). “The Ghost in the Machine has an American accent: value conflict in GPT-3”. ArXiv:2203.07785, <https://doi.org/10.48550/arXiv.2203.07785>

Moretti Bruno (2006). “Il laboratorio elvetico”. In: Moretti Bruno (a cura di). *La terza lingua. Aspetti dell'italiano in Svizzera agli inizi del terzo millennio*, vol. II, Locarno, Osservatorio Linguistico della Svizzera Italiana, 15-79.

Nardone Chiara (2017). “Gender and e-recruitment: a comparative analysis between job advertisements published for the German and the Italian labour markets”. *Labour and Law Issues*, 3 (1), 33-50.

Olita Anna (2006). “L'uso del genere negli annunci di lavoro: riflessioni sull'italiano standard”. In: Luraghi Silvia, Olita Anna (a cura di). *Linguaggio e genere*. Roma: Carocci, 143-154.

Raus Rachele (2023). “*Deep learning* e traduzione intralinguistica: riformulare i testi della pubblica amministrazione in modo inclusivo”, *Studi Italiani di Linguistica Applicata*, 3/2023, 552-569.

Sabatini Alma (1993 [1987]). *Il sessismo nella lingua italiana*. Roma: Presidenza del Consiglio dei Ministri.

Sarti Gabriele, Nissim Malvina (2024). “IT5: Text-to-text Pretraining for Italian Language Understanding and Generation”. In: Nicoletta Calzolari, Min-Yen Kan et alii (eds), *Proceedings of the 2024 Joint International Conference on Computational Linguistics, Language Resources and Evaluation (LREC-COLING 2024)*. European Language Resources Association (ELRA), 9422-9433.

Thornton Anna M. (2022). “Genere e igiene verbale: l’uso di forme con ə in italiano”. *ANNALI del Dipartimento di Studi Letterari, Linguistici e Comparati (Sezione linguistica)*, 11, 11-54. <https://doi.org/10.6093/2281-6585/9623>.

Sitografia

Portale d’impiego della Confederazione Svizzera: <https://www.stelle.admin.ch/stelle/it/home/stellen/stellenangebot.html>

An Inclusive Language Proofreader: Enhancing Writing with Equity and Diversity Principles Through AI Algorithms

Matteo Berta, matteo.bera@polito.it
Tania Cerquitelli, tania.cerquitelli@polito.it
Politecnico di Torino

Introduction

In recent years, the conversation surrounding inclusive language has gained significant momentum, reflecting a deepened awareness of the need for respectful and equitable communication. This shift is driven by a more extensive understanding of diversity, equity, and inclusion (DEI) across various sectors, including education, business, and public policy (Raus, 2023). As society continues to evolve, there is an increased recognition that language shapes perceptions, and therefore, inclusive language plays a vital role in fostering understanding and promoting equality (Van Dijk, 1993).

At the same time, emerging technologies such as Generative Artificial Intelligence (AI) are advancing rapidly, becoming not only more efficient but also more user-friendly. These advancements open up new possibilities for integrating technology into the process of creating inclusive communication. By combining human insight with the capabilities of AI, organizations, and individuals can address the challenges of biased language more effectively.

This chapter provides the necessary context for understanding the growing importance of inclusive language and highlights the motivation for developing tools that promote such communication. It explores how technology, especially AI, plays a crucial role in tackling these challenges, offering a bridge between traditional methods and modern solutions to create more inclusive and meaningful content.

1. Context and Motivation

Effective communication is essential for sharing ideas, drafting laws and regulations, and disseminating information. In formal settings such as public announcements, official minutes, and administrative messages, the use of language that reflects inclusivity is especially important.

Inclusivity in communication ensures that individuals of all genders, races, ethnicities, abilities and beliefs feel acknowledged, respected, and treated equitably (Vivian, 2022).

Despite its importance, achieving inclusivity in language use remains a challenge, particularly in professional and administrative context. Many individuals, including administrative professionals, lack proficiency in inclusive writing. This issue arises from several factors (Babiński *et al.*, 2019). First, there is a tendency to underestimate the negative effects of non-inclusive language, which can unintentionally convey bias or exclusion. Second, organizations often fail to provide adequate training or resources, leaving employees to navigate this skill independently. Finally, the prevalence of non-inclusive practices in existing formal communications perpetuates poor habits and undermines progress (Kolin, 2017).

Moreover, the prevailing language used in formal communications often contains inherent biases that reflect and reinforce societal inequalities. These biases can be subtle, embedded in word choice, sentence structure, and assumptions about gender roles or cultural norms (Berta *et al.*, 2024).

As a result, even well-intentioned communications may inadvertently exclude or marginalize certain groups, highlighting the need for tools and strategies that can identify and mitigate such linguistic biases.

The lack of skill in inclusive language is closely tied to the education system. In recent years, however, there has been a shift in the approach to language, with the newer generations increasingly adopting inclusive language in *a priori* manner, incorporating it as a default mindset. In this context, it is essential to create pathways for the older generation to bridge the gap, reducing the divide between generations.

The modern technologies, like Large Language Models (LLMs) can facilitate the adaptation to an equitable use of language, but they are trained on large-scale datasets that contain biases, and, for this reason, they can perpetuate these biases, amplifying the existing problem, as visible in Fig. 1. In this example, it is clear how commercial large language models such as GPT or Gemini, when asked to create fictional scenarios involving nurses and doctors, often reinforce stereotypical gender roles, portraying doctors as male and nurses as female. This pattern emerges because the models' internal

representations of these professions are shaped by the biases present in the extensive datasets on which they are trained. These datasets reflect societal stereotypes that can perpetuate biases. As a result, the models tend to reproduce these outdated and limiting views, demonstrating the challenges of achieving unbiased AI systems in domains where gender roles have historically been entrenched (Ferrara, 2023).

To address these challenges, integrating AI-powered proofreaders focused on inclusivity can offer a promising solution. These tools can detect non-inclusive language and provide suggestions for equitable reformulations, helping users refine their writing to align with principles of fairness and respect. By combining modern AI technologies with human oversight, inclusive tools can bridge the gap between generations and skill levels, fostering a more equitable use of language across formal and professional communications.

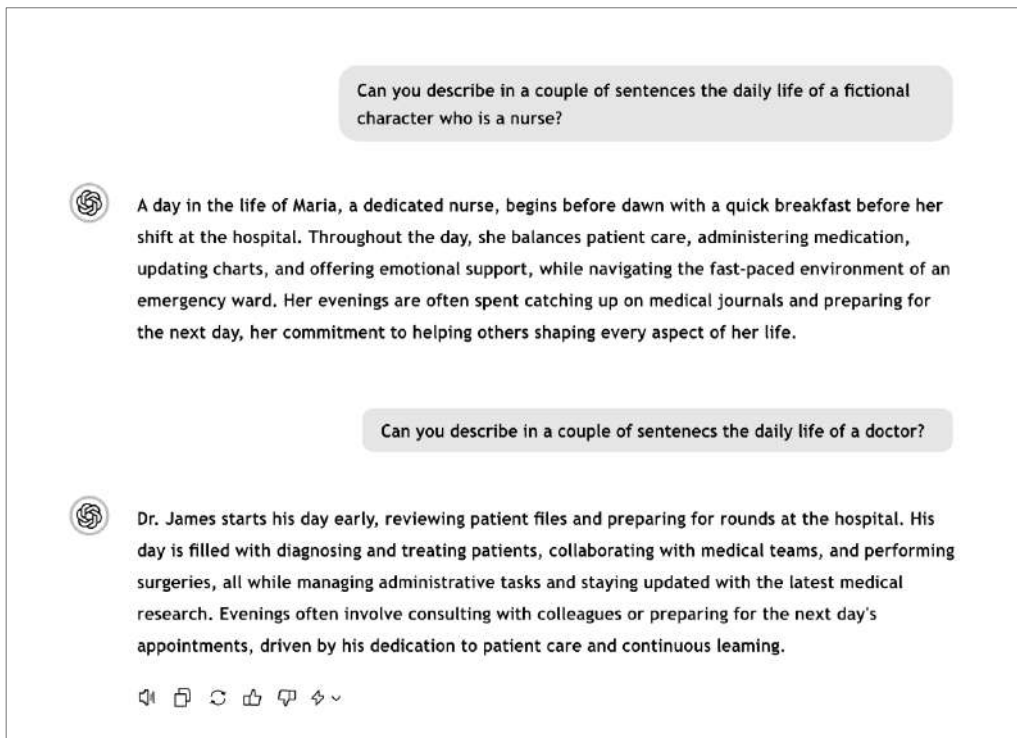


Fig. 1. GPT-4o reinforces gender stereotypes by depicting doctors as male and nurses as female when generating fictional daily life scenarios (November 11, 2024)

1.1 Gendered Languages

Inclusive language is a critical component of effective communication, as it actively avoids expressions that might exclude individuals. In the context of Italian formal communications, the need for inclusivity is particularly pronounced due to the gendered nature of the language. Italian, as long as other Romance languages, employs grammatical gender extensively, posing unique challenges to achieving gender-neutral expressions (Piergentili *et al.*, 2023).

For example, the phrase "The students" can be written in Italian as "*Gli studenti*" (masculine) or "*Le studentesse*" (feminine). However, the masculine form is often used as a default neutral option, which fails to promote inclusivity. An alternative is to use a gender-neutral term such as "*La componente studentesca*" (English: "The student body"). Unfortunately, gendered expressions like these frequently occur in formal communications, reinforcing the need for tools and frameworks that facilitate inclusive writing practices.

This issue is not limited to Italian but is also common in other Romance languages, such as Spanish and French, where similar challenges arise due to the pervasive use of gendered terms. In these languages, the use of masculine forms as a default has been the traditional practice, but there is a growing feeling that promotes gender-neutral alternatives. This widespread nature of this issue across multiple languages emphasizes the importance of developing solutions that can facilitate inclusive communication globally. In these languages, the universal quantifier is typically masculine, and introducing gender-neutral alternatives could, over time, transform the linguistic substratum, shifting it from a predominantly masculine structure to one that is less gendered.

1.2 Bridging the Gap in Inclusive Language Use

With this awareness it is easy to understand that the challenge of using inclusive language by default is particularly pronounced in Romance languages, where grammatical structures are deeply rooted in gendered forms. Older generations, who were educated during a time when there was little emphasis on inclusive language, often struggle to adapt to the modern shift towards gender-neutral or in-

clusive communication. Traditional educational programs, which largely ignored these linguistic and social issues, did not equip individuals with the tools or awareness to address gender bias in language. This creates a significant gap, as the language habits formed during this period are deeply ingrained, and shifting them requires both effort and understanding.

Older individuals, having learned the language in a more traditional context, tend to default to masculine forms when referring to mixed-gender groups, as this was once considered the standard. The shift toward inclusive language is thus challenging for this demographic, as they must overcome habits formed in an era where gender bias in language was rarely questioned.

In contrast, modern educational programs increasingly promote inclusivity, social equality, and sensitivity to diverse identities (see European Commission, 2023). Today's students are more likely to be exposed to the idea of inclusive language from an early age, with curricula that emphasize gender-neutral language and social awareness and various pedagogical studies focus on increasing awareness about inclusive educational frameworks (Moen *et al.*, 2008; McCoy *et al.*, 2023). This shift in focus reflects broader societal changes, where inclusivity and respect for diversity have become central values. Consequently, newer generations tend to use inclusive language more intuitively, often adopting it as a default in written and spoken forms.

As a result, older generations often find it challenging to embrace gender-neutral alternatives, which have become more prominent in modern educational programs. These generational differences highlight the growing need for tools that can bridge this gap, enabling older individuals to adopt more inclusive language while also allowing younger generations to maintain their awareness and practices of inclusivity.

The lack of formal training or resources compounds the difficulties older generations face in adopting inclusive language. Unlike today's students, who benefit from targeted educational programs and awareness campaigns, older individuals often rely on self-learning or informal sources for guidance on inclusive communication. As such, there is a clear need to bridge this gap, not only by providing tools and training but also by fostering an understanding of the importance of inclusivity in language use. This shift in perspective is necessary for ensuring that all individuals, regardless of age or

background, can engage in communication that reflects modern values of equality and respect.

1.3 “Inclusively”: a Proofreader Tool for Promoting Inclusive Language

To address the challenges highlighted in the previous sections, *Inclusively* emerges as a solution designed to facilitate the adoption of inclusive language practices. This tool harnesses the capabilities of Generative Artificial Intelligence to provide a user-friendly platform that helps in the diffusion of equitable communication while respecting linguistic nuances.

The idea behind the tool is to offer real-time feedback for inclusive writing by identifying non-inclusive sections of the text and suggesting more equitable reformulations of these portions. *Inclusively* can be envisioned as an intra-linguistic translator, a proofreader that transforms non-inclusive language into inclusive language (La Quatra *et al.*, 2024).

Recognizing the unique challenges of gendered languages such as Italian, French, and Spanish, *Inclusively* incorporates features that adapt to grammatical structures and cultural norms, ensuring its recommendations are relevant, accurate, and practical for the intended audience.

In the workplace, *Inclusively* could bridge generational differences by offering tailored resources that accommodate varying levels of familiarity with inclusive language.

Inclusively marks a significant step forward in integrating technology with inclusivity. By promoting equitable language use, it facilitates the shift toward communication practices that embody principles of equality, respect, and awareness.

2. *Inclusively* Tool

Inclusively is built on a comprehensive methodology designed to help users navigate the complexities of modern communication challenges. The tool utilizes advanced natural language processing (NLP) to identify and transform non-inclusive language into inclus-

ive alternatives, ensuring that its recommendations are both contextually appropriate and culturally sensitive. As highlighted in a recent leaked Google document (Google, 2023), “Data quality scales better than data size”. In line with this principle, *Inclusively* prioritizes high-quality task-driven data to effectively address the communication challenges it aims to solve.

An effective example of how *Inclusively* works is shown in Fig. 2. When a user inputs a text, the algorithm first analyzes it and provides feedback on the inclusiveness of each section. The second step generates an inclusive reformulation for the sections identified as non-inclusive.

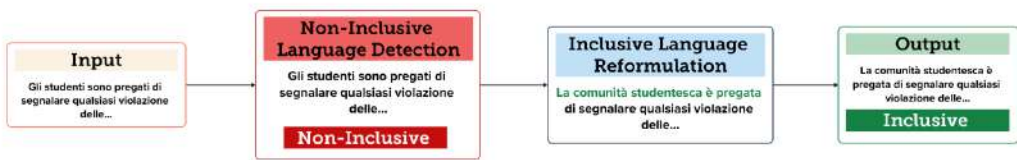


Fig. 2. An example of how *Inclusively* works: first, it detects non-inclusive text fragments, second, it reformulates them to create an inclusive version

2.1. Inclusively Pipeline

It is worth taking a closer look at what lies behind what the user can see Fig. 3 illustrates the complete pipeline that determines the selection of models used within the tool. The process begins with a collaborative effort between linguistic experts and Natural Language Processing (NLP) experts, laying the foundation of the system through the creation of a robust benchmark. This benchmark is developed in three key steps: defining linguistic criteria, gathering relevant data, and meticulously annotating the data to ensure high-quality labelled inputs. These steps are critical for creating a dataset that reflects diverse linguistic patterns and captures the nuances of inclusive and non-inclusive language.

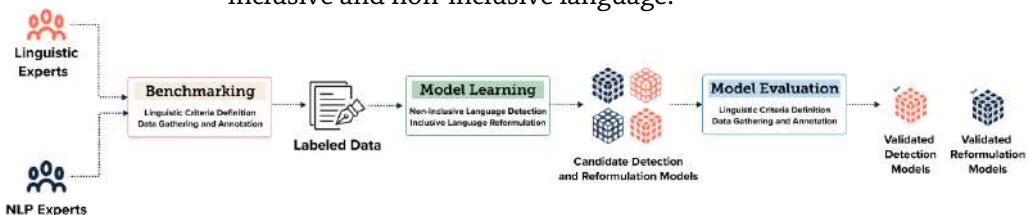


Fig. 3. The *Inclusively* Pipeline encompasses linguistic criteria definition, data gathering and annotation, model training for detection and reformulation, and validation of the detection and reformulation models

Once the labelled data is prepared, the pipeline advances to the Model Learning stage. Here, several deep learning models are trained to perform two complementary tasks: detecting instances of non-inclusive language and reformulating such language into a more inclusive form. This dual-task training ensures that the tool is capable of both identifying problematic expressions and offering constructive, context-sensitive alternatives.

Following the training phase, a Model Evaluation stage is conducted to assess the performance of the various models. This evaluation involves rigorous testing against multiple metrics, such as accuracy, precision, recall, and computational efficiency, to identify the most effective models. The best-performing models are then selected to execute the tasks within the Inclusively tool, ensuring that it operates with both precision and fairness.

By following this multi-step pipeline, the development process integrates linguistic expertise and technical innovation, culminating in a proofreader designed to promote inclusivity in language use.

2.2. Data Annotation

Expert-annotated data is crucial for training and evaluating models. The team with expertise in the field of inclusive language devised specific linguistic criteria for promoting inclusivity within romance languages. Then they collected documents highlighting pertinent inclusivity issues and formulated comprehensive annotations for these documents that serve as essential training for the AI model.

2.2.1 Definition of the Inclusive Language Criteria

The team began by engaging in discussions to exchange insights on inclusive language, fostering an open dialogue about its complexity and challenges. Building on these conversations, the team collaboratively drafted the linguistic criteria, drawing on the unique expertise of each member. This drafting process involved multiple iterations, during which the team critically evaluated the content for clarity, comprehensiveness, and practicality. Subsequent discussions provided opportunities to address concerns, incorporate feedback,

and refine the criteria further. Ultimately, the team reached a consensus, producing a well-structured and thoughtfully crafted set of linguistic criteria.

2.2.2 Data Gathering

During the gathering phase we collected documents written in various romance languages (Italian, French and Spanish) that reflect the language's usage in the administrative domain. The sources were various, including websites of public administrations and universities.

Domain experts meticulously choose these documents, which originate, for the Italian language, from reputable institutions such as the Città Metropolitana di Torino, Università di Bologna, Politecnico di Torino, and other prominent Italian organizations.

The dataset encompasses a variety of administrative genres and styles, including calls for bids, internal communications, policies, and regulations. These documents are segmented and annotated at the sentence level based on linguistic criteria, facilitating tasks such as the detection of non-inclusive language and its reformulation.

2.2.3 Data Annotation

The annotation task is divided in two main sections: non-inclusive language detection annotation and inclusive language reformulation annotation (Greco *et al.*, 2025). It is important to underline that annotation pipeline is performed at the sentence level and that multiple entities that can potentially contain inclusivity issues can be present within each sentence.

The classification task is formulated as a multi-class sentence classification problem with three distinct labels:

- **Neutral:** The sentence does not reference any protected group and is not relevant to the task of distinguishing between inclusive and non-inclusive language.
- **Inclusive:** The sentence explicitly references at least one protected group, and all inclusivity-related entities are expressed using inclusive language.

- **Non-Inclusive:** The sentence references at least one protected group, but at least one related expression is formulated in a non-inclusive manner.

The reformulation task is defined as a sequence-to-sequence problem where both the input and output sentences are in the same language. Annotators are instructed to provide one or more possible inclusive reformulations for non-inclusive sentences, adhering to established criteria. Each proposed reformulation must rewrite all non-inclusive expressions using inclusive language while preserving the original sentence's meaning and grammatical accuracy.

For each sentence, annotators first assign an inclusivity label from the categories: *neutral*, *inclusive*, and *non-inclusive*. For sentences labelled as *non-inclusive*, annotators provide one or more inclusive reformulations. Additionally, the platform allows annotators to propose non-inclusive versions of already inclusive sentences to expand the set of reformulation pairs. However, this functionality is used sparingly, as most source documents are predominantly written in non-inclusive language.

The final classification dataset includes all distinct original sentences and their reformulations. In contrast, the reformulation dataset consists exclusively of pairs of non-inclusive sentences and their corresponding inclusive reformulations. As a result, the two datasets share common sentences.

2.2.4 Human Validation and Evaluation

We have designed a comprehensive survey for expert linguistic annotators to validate the performance of the models developed for the classification task. This survey aims to assess the accuracy of the model's predictions regarding the inclusiveness of given sentences. For each sentence, along with its corresponding explanation, the annotators are asked to evaluate the model's prediction by selecting one of the following options:

- **Correct:** the prediction is influenced by only the words determining the inclusiveness or the non-inclusiveness of the sentence.

- **Partially Correct:** the prediction is influenced by mainly words relevant to the concept of inclusiveness, but also includes a significant set of words that are not relevant for classification in terms of inclusiveness.
- **Incorrect:** the prediction is influenced by only wrong words that do not determine the inclusiveness or non-inclusiveness of the sentence.

The aim for this evaluation is to ensure that the model's predictions are grounded in the linguistic features that are most significant for identifying inclusiveness and to detect potential biases or misinterpretations in the model's reasoning process.

To validate the models designed for the reformulation task, we have devised a human evaluation protocol wherein linguistic experts critically assess the qualitative accuracy of each proposed reformulation. For every sentence reformulated by the model, the expert assigns a qualitative label from the following categories:

- **Correct:** the reformulation is correct and can be accepted without any further modification
- **Partially Correct - Meaning Changed:** the reformulation is not fully correct and requires human intervention because the model partially changed the meaning of the original sentence.
- **Partially Correct - Subset:** the reformulation is not fully correct and requires human intervention to be accepted because the model solved only a subset of the non-inclusive expressions.
- **Incorrect:** the reformulation is incorrect because do not solve any of the non-inclusive expressions.

This human evaluation task is essential for determining the effectiveness of the reformulation model in generating inclusive language while preserving the original meaning of the sentences. By assigning detailed qualitative labels, we aim to gain insights into the model's strengths, limitations and areas requiring further refinement.

2.3 Application Scenario

To promote the adoption of inclusive language, we propose a scenario illustrated in Fig. 4. This framework comprises two main components:

1. **Classifier:** Functions as a detector of non-inclusive text.
2. **Generative Model:** Operates as a reformulation engine to enhance inclusivity.

Consider a non-linguistic expert tasked with drafting formal documents. To align with inclusive language standards, the document can be reviewed and revised using an AI-powered writing assistant, *Inclusively*. The tool offers a writing assistance interface that first scans the document to identify non-inclusive text snippets. It then applies generative modelling to reformulate these sections into gender-neutral, accessible forms. This ensures the text remains easily comprehensible to individuals with disabilities or visual impairments by avoiding specialized symbols or characters that might hinder accessibility.

If the suggested reformulation does not meet expectations, the tool provides an evaluation and annotation interface, enabling expert users to manually refine annotations for system improvement. Additionally, an inspection interface offers explanations of the model's outputs, helping users understand the AI's reasoning and improve their own writing skills.

This solution presents several advantages:

1. **Efficiency:** Reduces the need for linguistic experts to manually review raw content, lowering costs and increasing accessibility to inclusive writing.
2. **Generative AI Integration:** Harnesses advanced language models to create inclusive alternatives to inappropriate expressions.
3. **Empowerment:** Enables non-expert writers to learn correct writing practices, reducing future inaccuracies.

This approach makes inclusive writing more practical, scalable, and accessible for diverse users, including those in the public sector.

3. Exploration of Inclusively Tool

In this section will be presented a demonstration of *Inclusively*. To make *Inclusively* accessible and user-friendly we have developed a web-based interface that consists in three different interfaces, each shaped for the need of a particular group of users.

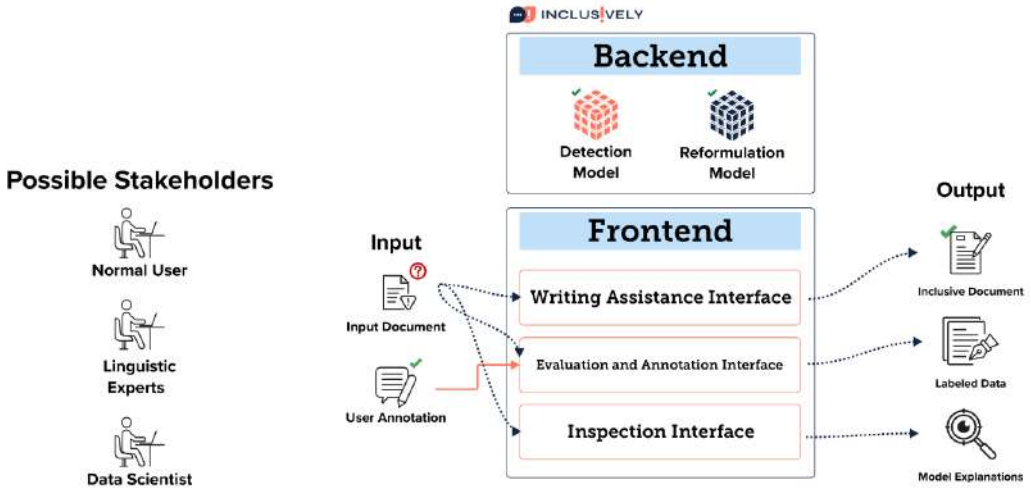


Fig. 4. Application Scenario demonstration

3.1 Writing Assistance Interface

The first interface is the Writing Assistance Interface. It is intended to support standard users, such as non-expert writers, to inclusive language writing. It benefits from the pipeline composed by the detection and reformulation models. The interface is designed to highlight all non-inclusive sentences and propose inclusive reformulations for each of them. Fig. 5 provides an illustrative example of how the interface functions in practice. The input text is analyzed and split into individual sentences, each of which is then evaluated and classified as either inclusive or non-inclusive. For instance, the sentence *“l'app contiene strumenti per docenti e studenti”* (EN: *“The application contains tools for teachers and students”*) is flagged as non-inclusive. This classification arises from the use of the masculine plural forms *“studenti”* and *“docenti”*, which can exclude other gender identities in contexts where inclusivity is important. To address this, the tool generates a reformulated version of the sentence designed to be more inclusive, which is then displayed to the user in green, providing a clear and actionable suggestion for improvement.

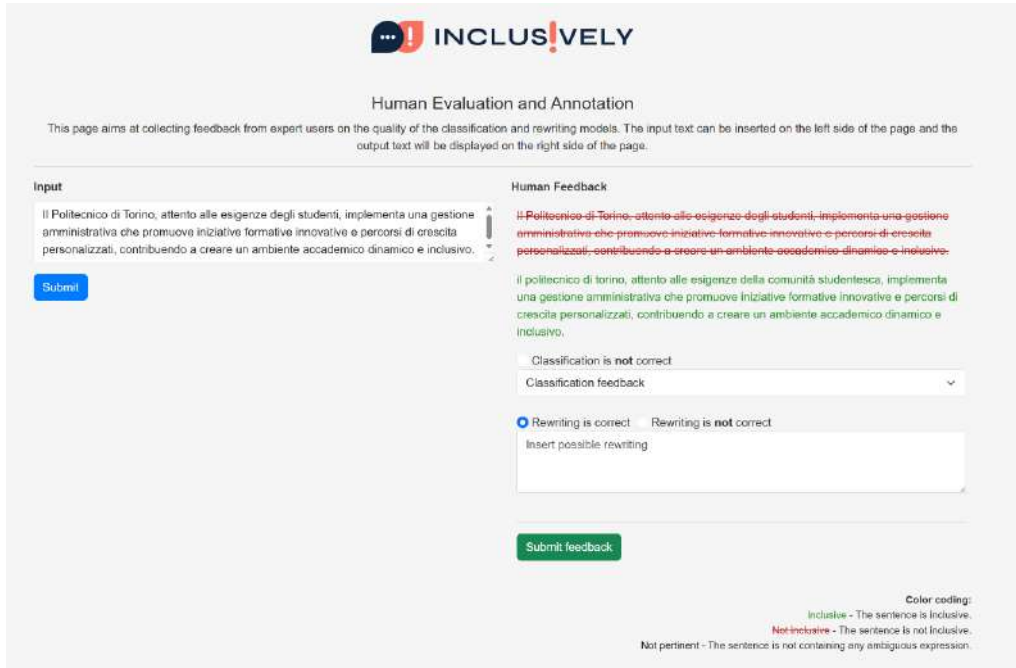
An Inclusive Language Proofreader: Enhancing Writing with Equity and Diversity Principles Through AI Algorithms



Fig. 5. In the interface of Inclusively, it is possible to visualize the sentences that are classified as non-inclusive (in red) with the corresponding inclusive reformulation (in green)

The second interface is specifically designed for expert users, such as linguists or communication specialists, who can contribute feedback to enhance the system's functionality. It enables users to manually annotate sentences, assess the system's classifications, and evaluate the proposed reformulations. This feedback is essential for assessing the tool's performance and collecting additional data to optimize and fine-tune its capabilities, ensuring ongoing accuracy and relevance. Fig. 6 illustrates an example of this interface using the same input as the previous example. Moreover, users can flag sentences they believe were misclassified and suggest alternative inclusive reformulations, providing critical insights for continuous system improvement (see Fig. 6).

The third interface is designed for data scientists, incorporating explainability models that highlight which parts of a sentence most influence the prediction of a specific label. This feature also identifies the words in the input sentence that contribute most significantly to the proposed reformulations, providing data scientists with a deeper understanding of the system's behavior. Fig. 7 demonstrates the output of this interface, explaining the detection model for the same input text as in previous examples. For instance, in this case, the words “*studenti*” (EN: “students”) was identified as key contributors to the model's classification of the sentence as non-inclusive.



INCLUS!VELY

Human Evaluation and Annotation

This page aims at collecting feedback from expert users on the quality of the classification and rewriting models. The input text can be inserted on the left side of the page and the output text will be displayed on the right side of the page.

Input

Il Politecnico di Torino, attento alle esigenze degli studenti, implementa una gestione amministrativa che promuove iniziative formative innovative e percorsi di crescita personalizzati, contribuendo a creare un ambiente accademico dinamico e inclusivo.

Submit

Human Feedback

Il Politecnico di Torino, attento alle esigenze degli studenti, implementa una gestione amministrativa che promuove iniziative formative innovative e percorsi di crescita personalizzati, contribuendo a creare un ambiente accademico dinamico e inclusivo.

il politecnico di torino, attento alle esigenze della comunità studentesca, implementa una gestione amministrativa che promuove iniziative formative innovative e percorsi di crescita personalizzati, contribuendo a creare un ambiente accademico dinamico e inclusivo.

Classification is **not** correct

Classification feedback

☒ Rewriting is correct ☐ Rewriting is **not** correct

Insert possible rewriting

Submit feedback

Color coding:
Inclusive - The sentence is Inclusive.
Not inclusive - The sentence is not Inclusive.
Not pertinent - The sentence is not containing any ambiguous expression.

Fig. 6. HUMAN FEEDBACK INTERFACE: this page let expert annotators to classify the proposed reformulation and to propose a new human-written reformulation of the sentence

This interface offers data scientists valuable insights into the underlying mechanics of the AI-based proofreader, aiding in the analysis and improvement of its performance. Additionally, it can be utilized by non-expert users for self-training in order to improve their skills in the use of inclusive language, enabling them to understand why a sentence is classified as non-inclusive and how reformulations are generated (see Fig. 7).

The *Inclusively* interfaces not only provide end-users with an accessible and user-friendly platform for leveraging the AI-based proofreader to support inclusive writing but also foster continuous improvement and development of the system. By allowing expert users to provide feedback and annotations, the tool facilitates a collaborative approach to fine-tuning, ensuring sustained accuracy and effectiveness. Through this iterative process, the project aims to advance the development of AI tools for promoting inclusivity in written communication.

Currently, the tool supports the Italian language, with plans to expand to other Romance languages where inclusivity challenges are prominent, such as French and Spanish.

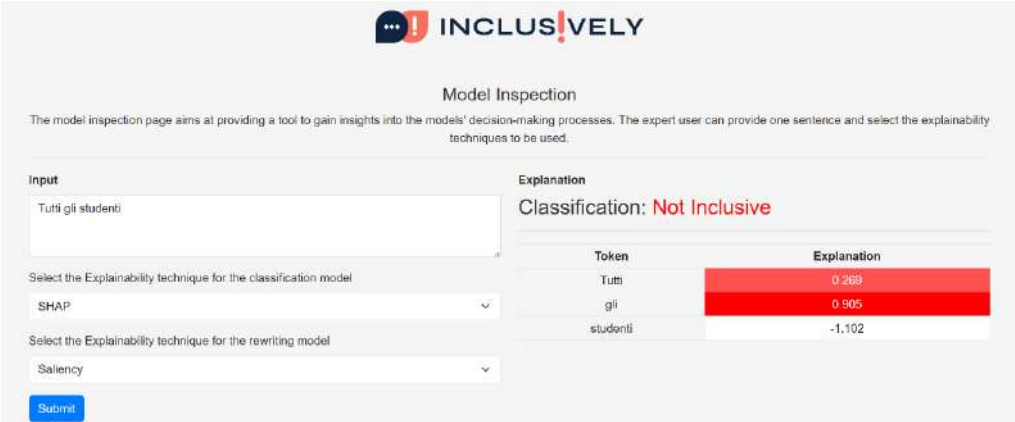


Fig. 7. The inspection Interface allows expert users to input a sentence and apply explainability techniques to analyze the models’ decision-making processes

4. Results

As largely explained, the models were rigorously evaluated using a combination of quantitative metrics and human validation to assess their performance comprehensively. The results are promising, indicating substantial potential for further advancements in both the classification and reformulation tasks (Greco *et al.*, 2023).

For the detection task the quantitative evaluation reveals that the best suited model is BERT-base-italian, outperforming BERT-multilingual in every category, as visible in Tab. 1.

Model	Accuracy	Inclusive F1	Non-Inclusive F1	Neutral F1
BERT-multilingual	0.86	0.83	0.89	0.83
BERT-base-italian	0.89	0.88	0.92	0.85

Tab. 1. Quantitative Evaluation for Detection Models

The results present the percentage of explanations categorized as correct (C), partially correct (P), and incorrect (I) across three groups: 50 inclusive sentences, 50 non-inclusive sentences, and a total of 100 sentences overall.

The BERT-base-italian model was rated as the most accurate in providing explanations for predictions across all categories. Specifically, its explanations were labeled as correct (C) in 76% of cases for inclusive sentences and 98% for non-inclusive sentences, resulting in an overall correctness rate of 87%. In contrast, the fine-tuned BERT-multilingual model achieved correctness rates of 58% for inclusive sentences and 90% for non-inclusive sentences, yielding an overall correctness rate of 74%. Both models maintained high standards of explanation quality, with the percentage of incorrect labels (I) never exceeding 2%. However, the BERT-base-Italian model exhibited a stronger alignment with relevant content, particularly when classifying inclusive sentences. We also measured the inter-annotator agreements using Fleiss' Kappa, which measures the level of agreement between two or more human annotators on a particular task. Overall, the agreement between the 7 participants for both models is 0.45. Following Landis and Koch's benchmark (Landis and Koch, 1977) to classify the level of agreement, which ranges from slight to almost perfect agreement, the obtained scores correspond to a fair agreement between annotators. This level of agreement likely reflects the inherent subjectivity in determining the exact set of words that contribute to a valid explanation. We can conclude that, from a human perspective, the BERT-base-Italian model provides the most accurate explanations for its predictions. These results align with the model's superior performance shown in the quantitative evaluation, supporting its selection as the overall best-performing model. Therefore, this classifier is currently integrated as the detection model of *Inclusively*.

For the human evaluation of the reformulation models, we assessed all models through human validation since their quantitative performance was comparable. A set of 100 non-inclusive sentences from the test set, along with corresponding reformulations, was evaluated by six expert annotators specializing in inclusive language. These reformulations were categorized as correct (C), partially correct with meaning changed (P-M), partially correct as a subset (P-S), or incorrect (I).

The results reveal that IT5-efficient and IT5 models with templates performed similarly, with template-based models achieving 73% correct reformulations. However, the IT5-efficient model had a higher rate of completely incorrect reformulations, whereas the

IT5 model more frequently produced partially correct subset cases, highlighting its ability to make incremental improvements and reduce fully incorrect outputs.

Inter-annotator agreement, measured by Fleiss' Kappa, was 0.702, indicating substantial consensus among annotators. Dividing partially correct cases into P-M and P-S categories likely enhanced evaluation consistency.

In conclusion, all models demonstrated strong performance, but the IT5 model augmented with templates stood out for its reliability and fewer fully incorrect cases. As a result, the IT5 model was chosen as the reformulation model for *Inclusively*.

Discussion and future research directions

In this paper, we introduced the first benchmark for Italian inclusive language and fine-tuned transformer-based models for tasks focused on detecting non-inclusive language and generating inclusive language reformulations. Through a combination of quantitative metrics and human-driven evaluation, we validated the effectiveness of the trained models. The top-performing models were integrated into a two-stage pipeline of the writing assistance proofreader, *Inclusively*. This tool aims to promote inclusive communication by identifying sentences with potentially harmful expressions and providing revised, inclusive alternatives.

By facilitating inclusive writing, particularly in administrative and academic contexts, the proposed models contribute to fostering diversity and inclusivity within society. Furthermore, the models and the source code for the tool have been made publicly available, encouraging the research community to build upon this work. This open availability is intended to address existing limitations in state-of-the-art NLP systems, including decision-making frameworks, language translation models, and sentiment analysis tools.

As inclusive language becomes an increasingly important aspect of written communication and AI-based tools demonstrate significant potential in facilitating its adoption, it is crucial to acknowledge their existing room for improvement. Our approach encounters some challenges:

1. **Italian Language Specificity:** The dataset, models, and application tool developed in this study are tailored specifically for the Italian language. This language-specific focus limits their applicability in multilingual or non-Italian contexts. Although this paper demonstrates the potential of detection and reformulation models in fostering inclusivity using expert-annotated Italian data, extending this approach to other languages would require substantial re-training and re-annotation efforts. The reliance on language-specific features emphasizes the difficulty of generalizing these findings across different linguistic settings, underscoring the need for further research to adapt these methodologies for broader linguistic applicability.
2. **Domain-Specific Focus:** Our study focuses on inclusive language within formal communication, particularly in administration and academia. The classification and reformulation models are trained on administrative documents, which may not directly apply to other contexts, such as legal or web-based communication. Applying this approach to different domains would require further annotation and training, along with the development of tailored linguistic criteria specific to each communication style or data domain.
3. **Limited Dataset Size:** Our dataset is based on high-quality expert-annotated data for inclusive language, which makes the process of generating new data time-consuming. As a result, the dataset remains relatively small and may not capture the full range of expressions encountered in formal communication. This limitation stems from the difficulty of collecting a diverse and comprehensive set of examples that reflect the complexity of inclusive language. There is a clear need for ongoing data collection and expansion to improve coverage and ensure robustness in future research.

We are working to address the limitations explained above, focusing on the development of a multilingual solution that promotes inclusive communication across various domains. This process requires additional training data and linguistic resources to ensure the capability to identify and reformulate texts in multiple languages. One of the ongoing efforts involves extending the approach to different domains, such as legal communications, to create models suited for the specific linguistic patterns of this context. The solution will

be expanded to include Romance languages such as French and Spanish. Finally, we aim to promote its use on a large scale to assess effectiveness, collect more data, and refine the models.

Acknowledgments

This study was carried out within the project “E-MIMIC: Empowering Multilingual Inclusive Communication” (Nr. 2022WEFCFP), funded by the *Ministero dell’Università e della Ricerca* — with the PRIN 2022 (D.D. 104 - 02/02/2022) program.

References

- Attanasio Giuseppe, Cagliero Luca, Cerquitelli Tania, Greco Salvatore, La Quatra Moreno, Raus Rachele, Tonti Michela (2021). “E-MIMIC: Empowering Multilingual Inclusive Communication.” In *2021 IEEE International Conference on Big Data (Big Data)*, Orlando (USA), 15-18 dicembre 2021, 4227-4234.
- Babinszki Tom, Cavender Anna, Gower Michael, Hoehl Jeffery, Lima Darcy, Manser Erich, Trewin Shari (2019). “Inclusive Writing”. In *Web Accessibility — A Foundation for Research, Second Edition*, edited by Yeliz Yesilada and Simon Harper. Berlin: Springer, 135–152.
- Berta Matteo, Greco Salvatore, Tipaldo, Giuseppe, Cerquitelli Tania (2024). “Decoding Narratives: Towards a Classification Analysis for Stereotypical Patterns in Italian News Headlines”. In *DS4EIW 2024 Workshop in conjunction with the IEEE BigData 2024 conference*.
- Ferrara Emilio (2023). “Should ChatGPT be Biased? Challenges and Risks of Bias in Large Language Models” (October 26, 2023). <https://ssrn.com/abstract=4614228> or <http://dx.doi.org/10.2139/ssrn.4614228>
- Greco Salvatore, La Quatra Moreno, Cagliero Luca, Cerquitelli Tania (2024). “Towards AI-Assisted Inclusive Language Writing in Italian Formal Communications”. Under review at *ACM Transactions on Intelligent Systems and Technology (TIST)*.
- Greco Salvatore, La Quatra Moreno, Cagliero Luca, Cerquitelli Tania (2025). “Towards AI-Assisted Inclusive Language Writing” in *Italian Formal Communications. ACM Trans. Intell. Syst. Technol.* Just Accepted (April 2025). <https://doi.org/10.1145/3729237>
- Kenny Neil, McCoy Selina, O’Higgins Norman James (2023). “A Whole Education Approach to Inclusive Education: An Integrated Model to Guide Planning, Policy, and Provision”. *Education Sciences* 13 (9), 959.
- Kolin Philip C. (2007). *Successful Writing at Work*. Houghton Mifflin.
- Landis J. Richard, Koch Gary G. (1977). “The Measurement of Observer Agreement for Categorical Data”. *Biometrics* (1977), 159–174.
- La Quatra Moreno, Greco Salvatore, Cagliero Luca, Tonti Michela, Dragotto Francesca, Raus Rachele, Cavagnoli Stefania, Cerquitelli Tania (2024). “Building Foundations for Inclusiveness through Expert-Annotated Data.” In *EDBT/ICDT Workshops 2024*.
- Moen, Torill (2008). “Inclusive Educational Practice: Results of an Empirical Study”. *Scandinavian Journal of Educational Research* 52 (1), 59–75.
- Piergentili Andrea, Fucci Dennis, Savoldi Beatrice, Bentivogli Luisa, Negri Matteo (2023). “Gender Neutralization for an Inclusive Machine Translation: From Theoretical Foundations to Open Challenges”. In *Proceedings of the*

First Workshop on Gender-Inclusive Translation Technologies. European Association for Machine Translation, Tampere, Finland, 71–83.

Raus Rachele (ed.) (2023). *How Artificial Intelligence Can Further European Multilingualism*. Milan: Ledizioni.

Ta Vivian P. *et alii* (2022). “An Inclusive, Real-World Investigation of Persuasion in Language and Verbal Behavior”. *Journal of Computational Social Science* 5 (1), 883–903.

Van Dijk Teun Adrianus (1993). *Elite Discourse and Racism*. London: Sage Publications.

Van Dijk Teun Adrianus (1993). “Principles of Critical Discourse Analysis.” *Discourse & Society* 4, 249–283.

Webliography

European Commission. (2023). “Inclusive or Special Needs Education?” *School Education Gateway*. Last modified November 3, 2023. <https://school-education.ec.europa.eu/en/discover/viewpoints/inclusive-or-special-needs-education>

European Human Agency for Fundamental Rights. (2022). *Bias in Algorithms. Artificial Intelligence and Discrimination*. Vienna. https://fra.europa.eu/sites/default/files/fra_uploads/fra-2022-bias-in-algorithms_en.pdf

SemiAnalysis. (2023). “Google ‘We Have No Moat, and Neither Does OpenAI.’” *SemiAnalysis*, May 4, 2023. Available at: <https://semianalysis.com/2023/05/04/google-we-have-no-moat-and-neither/#data-quality-scales-better-than-data-size>

Il progetto E-MIMIC

Introduzione

La natura della riflessione, che seguirà nelle prossime righe, prende direttamente le mosse da diverse attività di annotazione condotte in italiano per il progetto E-MIMIC. Esse non hanno fatto altro che sollevare inevitabilmente delle questioni fondamentali sulle modalità in cui procedere verso un utilizzo lecito e corretto della propria lingua, così come in maniera auspicabile di tutte, giacché, per operare sui testi, si è trattato di verificare *in primis* l'idoneità delle competenze linguistiche di chi scrive. In generale, intervenire su un qualsiasi *dataset* testuale comporta l'analisi di ogni singolo segmento presente nel *file*, accertarsi dell'assenza di tratti non inclusivi al suo interno per procedere oltre, altrimenti proporre una riformulazione, per l'appunto, inclusiva. Quindi, in prima analisi, è stato il dover confrontarsi con la pratica della traduzione "intralinguistica" ciò che ha permesso il sorgere di domande di cui s'è appena detto.

E quanto appena riportato è già sufficiente a far immaginare le metodologie di ricerca impiegate, ma soprattutto a che genere di strumenti si è ricorsi. Un passo preliminare è quello di fare chiarezza sulla successione cronologica con cui si è deciso di adoperare questi strumenti, a cui si è accennato. Infatti, essendo stata condotta un'attività di annotazione su circa un totale di novemila segmenti testuali di varia natura, le prime impressioni che si sono ingenerate sono catalogabili come "riflessioni linguistiche", per l'appartenenza delle domande, che sono sorte, a branche di studio proprie della cosiddetta linguistica. Tuttavia, nel corso della riflessione il terreno preparato da meditazioni legate alla lingua come "ente", si è necessariamente approdati a considerare la natura logico-semantica e semiologica dei sintagmi testuali propriamente responsabili dell'intervento di riformulazione sul segmento corrispondente. Pertanto, il ricorso alle rispettive discipline è stato indispensabile, al fine di conferire organicità strutturale alla presente trattazione in qualità sia di "pensiero" sia di "testo".

Perciò, prima di dirigersi verso la trattazione vera e propria del problema che vogliamo affrontare in questa sede, l'obiettivo è stato, e sarà di seguito, quello di allestire una formalizzazione (o individuazione) del problema stesso, di localizzarne le componenti, determinare le implicazioni, per poi procedere effettivamente con il ragionamento sulle questioni sollevate, ricalcando i passi elaborati durante il lavoro effettivo. E una volta fatto tutto ciò, riservare uno spazio alla discussione delle modalità proposte per intervenire adeguatamente sul testo.

1. La natura linguistica del problema

Volendo usufruire di una metafora, questo contributo nasce sul “campo”, da riflessioni che hanno raggiunto le strutture fondamentali della lingua italiana. Come già detto infatti, l'obiettivo dell'attività di annotazione in seno al progetto E-MIMIC è stato appunto riformulare tutti i luoghi testuali ritenuti non inclusivi in inclusivi con la seguente metodologia: individuare la presenza di sintagmi non inclusivi all'interno del frammento di testo e proporre in seguito la riformulazione. Beninteso, non si parla di correzioni, ma di vere e proprie traduzioni intralinguistiche di sintagmi grammaticalmente “corretti”. Sequenze quali “tutti i candidati”, “tutti i partecipanti”, che esplicitano il riferimento alla totalità di una categoria facente parte di un insieme attraverso l'aggettivazione indefinita, hanno richiesto la necessità di intervenire sul testo. Non solo, ma lo stesso trattamento è stato riservato a sintagmi soggetto o complemento (diretto o indiretto) che presentavano la medesima forma, oppure senza l'ausilio dell'aggettivo indefinito “tutto”: mi riferisco a soluzioni diverse, quali “i candidati”, “i partecipanti”, cioè differenti nel significante (assenza dell'aggettivo “tutto”), ma dal punto di vista semantico equivalenti ai primi esempi, riportati con la presenza dell'aggettivo. Si tratta di comunissimi luoghi testuali in cui emerge una particolare norma dell'italiano, per la quale il maschile, se al plurale, può assumere il carattere di “non marcato”, guadagnando quindi la possibilità di riferirsi ad una categoria, o insieme, nella totalità dei suoi elementi¹ senza curarsi del genere.

¹ A onor del vero va detto che, effettivamente, i sintagmi citati non rappresentano la completa fenomenologia di soluzioni non inclusive che sono state trattate in sede di

Iniziano a delinearci con ciò i primi tratti del problema. Infatti, il protocollo di intervento atto a formulare proposte inclusive induce, come prima opzione, a ricorrere alla sostantivazione astratta (es. “tutti i dirigenti scolastici” > “la dirigenza scolastica”), in seconda analisi a ricorrere a perifrasi o artifici consimili più o meno retorici, spesso concessi dal ricorso a [pronomi relativo] + [terza persona del verbo]². In sostanza si sta parlando di traduzione intralinguistica, una procedura usuale in ogni lingua, quando occorre cambiare (o meglio quindi “tradurre”) registro.

A questo punto, bisogna determinare perché si definisce linguisticamente “problema” la versione di un segmento non inclusivo in inclusivo, come nel caso della trasposizione di “tutti i candidati” (sintagma non inclusivo, poiché indicante sia uomini che donne) in “tutte le persone candidate” oppure “la totalità delle persone candidate” (entrambi inclusivi). La motivazione principale affonda le sue radici nelle nozioni di norma ed errore, categorie semplificabili nei più comuni “giusto o sbagliato”. Secondo, infatti, l’uso corrente della lingua italiana un qualsiasi tipo di sintagma che si presenti con la struttura “tutti gli x”, in cui x sta per sostantivo plurale maschile di persona o di professione, non viene avvertito da chi parla come portatore di errore. Questo perché nella storia dell’uso della lingua italiana si è imposto l’impiego del maschile non marcato per riferirsi alla totalità di un insieme: tale fenomeno si è imposto quale norma. Tuttavia, prima di proseguire con il tentativo di inquadrare la presente questione nell’ambito della norma e dell’errore, andrebbe precisato che cosa si intende con tali nozioni. Per far ciò, risultano veramente d’aiuto le parole del linguista Luca Serianni:

Il concetto di norma linguistica ha qualche affinità con quello di norma giuridica. Nel diritto, l’infrazione alla norma penale fa scattare una sanzione. Nella lingua la sanzione, pur non essendo codificata puntualmente, può colpire o attraverso un giudizio scolastico (con la conseguenza di ripetere un anno di scuola o di non superare la prova scritta di un concorso) o attraverso la squalifica sociale: se un medico scrivesse *raggione* o *esperiensa*, probabilmente dubiteremmo della sua professionalità (Serianni, Antonelli, 2017: 221).

valutazione dei *corpora* linguistici, poiché entro appunto i confini di classificazione “non inclusiva” sono stati trattati in egual modo dei segmenti aventi al loro interno sigle o acronimi, la cui presenza all’interno del testo rende non inclusivo il segmento per motivi di leggibilità, non di genere.

² Un esempio molto comune potrebbe essere: “i partecipanti” > “chi partecipa”.

Non essendoci quindi una codificazione della sanzione linguistica, consegue che la norma si stabilisce per accordo sociale, come risuona nelle parole di Coşeriu: "...*norma* nel senso corrente, stabilita od imposta secondo i criteri di correttezza o di valutazione soggettiva di quel che viene espresso" (Coşeriu, 1971: 76). La sanzione si riduce nella maggioranza dei casi ad assumere i tratti della squalifica sociale, come emerge ancora una volta dalle parole del medesimo autore sopra menzionato:

Per definire ciò che è giusto e ciò che è sbagliato in una lingua, occorre tener conto di una variabile fondamentale: il grado di accettabilità, ossia la reazione dei parlanti di fronte alla violazione di un certo istituto linguistico (*Idem*).

A fronte delle precedenti citazioni, riguardo alla norma si potrebbe dire senza problemi che essa si accordi all'idea che padroneggiare il linguaggio implichi il fatto che gli enunciati emessi non causino la reazione negativa degli interlocutori: gli atti linguistici "corretti" consistono solamente nel veicolo di un messaggio semantico legato al significato (al contenuto, semplificando) delle frasi. E questo consuona con il celebre esordio del primo capitolo de *Il Linguaggio* di Leonard Bloomfield:

Il linguaggio ha una funzione importante nella nostra vita; però raramente gli prestiamo attenzione, forse perché ci è così familiare, e lo consideriamo qualcosa di scontato, come il respirare o il camminare (Bloomfield, 1996: 5).

Perciò, l'utilizzo del maschile plurale non marcato, che nel tempo ha assunto il ruolo di norma grammaticale, nella sua concezione di accordo sociale presentata nelle precedenti righe, intercetta gli estremi necessari ad essere considerato una vera e propria problematica, poiché da un lato v'è una fetta di parlanti che, ritenendolo una regola costitutiva della grammatica italiana, ne percepisce legittimo l'uso (immediato, diremmo), dall'altro v'è chi sostiene che questa "norma" sia insufficiente a rappresentare il gran numero di categorie di persone, venendo esse escluse dall'impiego del maschile plurale, ingenerando in tal modo problematiche di matrice sessista. Pertanto, è questa la luce sotto la quale va osservata la questione. Da un lato, la

norma del maschile non marcato, non percepita come violazione, non suscita in chi parla la necessità di intervenire, trovando soluzioni alternative al sintagma; dall'altro abbiamo chi invece sostiene il fatto che la violazione sussiste, ma non nei confronti di una norma linguistica, quanto piuttosto in relazione alle conseguenze sociali di tale norma, cioè l'esclusione dei gruppi sociali non raggiunti dall'impiego del maschile, appunto. A tal riguardo, si segnala che il primo tentativo di porre rimedio, sollevando la questione, è stato attuato da Alma Sabatini nell'opuscolo *Raccomandazioni per un uso non sessista della lingua italiana* del 1986 pubblicato a cura della Presidenza del Consiglio, ristampato poi l'anno successivo ne *Il sessismo nella lingua italiana*, all'interno del quale Sabatini effettivamente fornisce indicazioni per evitare, anzi rifiutare, il maschile "generico" che indicasse sia componenti maschi che femmine, secondo uno schema che ricorda la modalità di correzione dell'*Appendix Probi* (A non B). L'intervento, quindi, è indirizzato a combattere la tendenza sessista all'utilizzo del maschile non marcato, impostosi come norma (confermato quindi dal grado di accettabilità di cui gode fra i parlanti) in italiano. Ma a questo punto, il passo necessario per approfondire la questione prevede l'interrogazione delle fonti, nelle quali è depositata la norma linguistica, le grammatiche. Nella grammatica Garzanti del 2012, nella sezione in cui si discute dell'accordo, viene prescritto che l'aggettivo assume lo stesso genere e numero del nome a cui si riferisce e nello specifico viene riportato che «se l'aggettivo si riferisce a più nomi, va sempre al maschile se i nomi sono tutti maschili, oppure se sono maschili e femminili, l'accordo al femminile avviene solo nel caso in cui i nomi sono tutti al femminile», (AA. VV., 2012: 20). Normative di carattere simile provengono dall'agile grammatica delle Garzantine del 2003, nella quale viene segnalato che, qualora occorra elaborare un accordo fra nomi di genere diverso, l'aggettivo deve assumere il numero plurale e, di preferenza, il genere maschile (Serianni Castelvechi *et al.*, 2003: 141). Stessa cosa si testimonia nella grammatica italiana della Treccani, cfr. (Santi, Viale, 2012: 192). E considerazioni simili, relative all'accordo, possono ritrovarsi anche in altre grammatiche (Schwarze, 2009, o anche Dardano Trifone, 1995: 214) in cui riguardo al maschile plurale italiano si informa che esso è più vicino al "neutro". Riassumendo, quindi, è stato esaminato che quando si parla di accordo in lingua italiana, è previsto, e consigliato, l'utilizzo del maschile nel momento in cui si necessita di realizzare sin-

tagmi con soggetti multipli di genere misto, e solo nel caso di due sostantivi femminili è previsto che l'aggettivo si accordi a tale genere. Nei riguardi di questa tendenza dell'italiano, arrivati a questo punto, si può quindi stilare una breve rassegna delle interpretazioni da poterle affidare, per cercare di elaborare un quadro della situazione entro cui osservare il fenomeno. La strada da seguire è semplice: manifestare quali sono le vie ermeneutiche percorribili, vale a dire dichiarare quali sono i modi con cui è possibile interpretare i fatti appena espressi. Tali modalità ermeneutiche sono incanalabili in due macroambiti: quello della linguistica di genere, alla quale va riconosciuto il merito di aver portato a consapevolezza del fatto che alcune norme e usi della lingua alimentino, inconsapevolmente, una condotta discriminatoria, quello della linguistica descrittiva *tout court*, che si limita a trovare appunto una spiegazione al fenomeno del genere, descrivendo il suo funzionamento all'interno del sistema lingua, senza includere questioni sociali, o meglio sociolinguistiche. Il primo dei due punti di vista descrive una problematica sociale che si ingenera quando nell'enunciato operano nomi con genere grammaticale *gender-indefinite* (vale a dire semanticamente non marcati), nel senso che l'utilizzo del maschile non marcato, quale meccanismo, percepito come corretto nel sistema lingua italiana, attiva delle dinamiche sessiste nel rapporto tra lingua e ambito sociale italofono. Secondo invece il punto di vista della linguistica descrittiva *tout court*, il maschile non marcato altro non è che un meccanismo che opera su basi di economicità linguistica, come ricorda Lombardi Vallauri: «Per giudicare realisticamente la funzione del maschile non marcato, occorre rendersi conto che esso è un meccanismo economico della lingua», (Lombardi Vallauri, 2024: 75). Entrambe le rotte ermeneutiche esprimono un giudizio condivisibile, nel senso che l'una non esclude l'altra. Il passo che rimane è quindi fare un po' di chiarezza sulla gerarchia, in base alla quale ordinare le due interpretazioni. Infatti, quando si parla di economicità in relazione al meccanismo linguistico dell'accordo, la cosa è innegabile: la funzione della struttura viene catturata pienamente nella sua identità; tuttavia la descrizione fornita si limita a delineare motivazioni strettamente linguistiche per l'esistenza dell'accordo con maschile non marcato, ovvero l'utilizzo del maschile riferito a sostantivi anche femminili. Allo stesso tempo, però, la spiegazione è insufficiente a descrivere quali conseguenze ha tale struttura, allorché si pensa al rapporto tra lingua e società. In questo senso

va inteso quanto detto prima, vale a dire che scegliere una sola delle due interpretazioni esclude la verità dell'altra: descrivere un meccanismo, una struttura linguistica, è semplicemente preliminare, e quindi un passo necessario, alla discussione che tale struttura può avere quando della stessa se ne fa uso effettivo nella conversazione. Pertanto, è vero che l'accordo è una struttura che garantisce economicità di funzionamento alla lingua, ma è altrettanto vero che essa attiva delle dinamiche di stampo sessista, quando della lingua si fa uso, cioè quando la lingua entra effettivamente in gioco nella società. Per concludere, riguardo alle due interpretazioni proposte possiamo dire che la seconda è la base della prima. Quindi, bisogna concentrare l'attenzione sulle già menzionate conseguenze di un utilizzo sessista della lingua e per farlo, la trattazione si sposta su un punto di vista logico-semantico, che punterà i riflettori sull'aspetto operativo della lingua, vale a dire quando essa è in azione.

2. Tra logica e semantica

Come è stato appena illustrato, le modalità con cui inquadrare entro l'ampia problematica dell'accordo il ruolo del maschile non marcato, che entra in gioco allorché v'è una confluenza di sostantivi di vari generi (es. "ragazzi e ragazze vincitori", "candidati e candidate idonei"), e dell'altro impiego di tale maschile, quando invece c'è necessità di identificare una certa categoria di individui (es. "tutti i ragazzi" "tutti i candidati"), sono state esplicitate, seguendo due direttrici distinte: il punto di vista unicamente descrittivo e quello che tiene conto del valore sociale delle costruzioni linguistiche. L'analisi condotta fino ad ora ha avuto carattere preliminare: stabilire cioè i dettagli della questione, affinché siano identificati gli elementi atti a proseguire con il lavoro; quindi giungere al secondo livello della presente indagine, che svolge allo stesso modo il ruolo di preparare il campo all'analisi finale (la terza parte di questo intervento). In questa sezione si riserverà infatti dello spazio a fornire delle informazioni riguardo agli ambiti del ragionamento umano in cui il maschile non marcato è strumento indispensabile per la formazione di alcune tipologie specifiche di enunciati apofantici, cioè quella categoria di enunciati ai quali si può attribuire il valore di "vero" o "falso". Si tratta infatti di una componente irrinunciabile per la costru-

zione delle argomentazioni, vale a dire le unità linguistiche e cognitive di base, attraverso cui con il linguaggio umano è possibile creare e aggiungere conoscenza nel mondo. Un qualsiasi tipo di ragionamento, infatti, parte da presupposti linguistici, le frasi (che prendono il nome di formule), il cui requisito necessario è quello di essere ben formate (essere cioè corrette per il sistema lingua), affinché possa essere affidata loro l'estensione logica in senso fregeano; in altre parole che queste formule possano essere valutate secondo un criterio di verofunzionalità, il quale consiste nella possibilità di affibbiare a una formula ben formata il valore V, per "vero", oppure quello F, per "falso". Una piccola precisazione: l'obiettivo di questa sezione non è volgere l'attenzione alle modalità di realizzazione del sillogismo, né tantomeno discuterne i generi o l'efficacia attraverso lo studio delle argomentazioni, ma più semplicemente rivolgere lo sguardo critico ai dispositivi linguistici che contestualmente garantiscono la creazione di un sillogismo. Tale premessa è doverosa, poiché, come si vedrà, in italiano il maschile non marcato si applica per la maggior parte delle volte, quando, come già detto, si deve ricorrere a riferirsi alla totalità di una categoria, in altre parole quando nell'enunciato logico si deve far ricorso alla quantificazione, e quindi al conseguente impiego dei cosiddetti quantificatori ($\forall$ per la quantificazione universale, $\exists$ per quella esistenziale). Una volta determinate le modalità linguistiche attraverso cui l'italiano garantisce la possibilità di quantificare (nello specifico: come costruire grammaticalmente una frase per far sì che ciò accada), sarà possibile passare alla sezione riservata alla semiotica, cioè osservare come tutto nel mondo, anche una "semplice" scelta grammaticale, sia in grado di comunicare.

2.1 Logica dei predicati e delle asserzioni categoriche

Una parte indispensabile di ogni corso di logica verte sulla tematica delle asserzioni elaborate attraverso l'impiego di quantificatori. Alcuni esempi:

- 1) Tutti i gli amici di Anna sono vegetariani
- 2) Mara è amica di Anna

∴ Mara è vegetariana

L'argomentazione, appena mostrata, la cui forma proposizionale risulta essere la seguente: 1) ha la funzione di informare sul fatto che una certa persona, Mara, fa parte dell'insieme, o meglio della categoria, di chi è amico di Anna. Riguardo a questo insieme viene predicato un altro fatto, che ogni membro al suo interno non mangia carne. In seguito, siccome si aggiunge che Mara fa parte di questo insieme, si può concludere senza problemi che anche Mara è vegetariana. Si tratta di un'argomentazione banale. La comprensione del suo funzionamento è tanto elementare quanto allo stesso modo è verificarne la validità. Ma la cosa che interessa in questa sede è in particolare la prima premessa: "tutti gli amici di Anna sono vegetariani". Concentrando, infatti, l'attenzione sul soggetto della frase, "tutti gli amici", notiamo verificarsi l'applicazione del maschile plurale non marcato, di cui si è parlato nella precedente sezione. E la tal cosa risulta di notevole interesse ai fini di questa ricerca, perché la frase considerata presenta due prospettive di analisi inscindibili fra loro: ancora una volta, quella linguistica e quella logica, in cui la prima implica la seconda. Dal punto di vista dell'analisi logica della frase il sintagma "tutti gli amici" svolge il ruolo di soggetto, "di Anna" è un complemento di specificazione, "sono vegetariani" nel suo insieme di copula e parte nominale costituisce il predicato nominale. Tuttavia, la parte che interessa la trattazione è proprio il soggetto linguistico da una prospettiva d'indagine logica, ma intesa nel senso di logica come disciplina filosofica. Il sintagma-soggetto in questione altro non è che la versione linguistica italiana del quantificatore universale $\forall$. Nella frase considerata si sta facendo riferimento alla totalità di un insieme e, per farlo, l'italiano costruisce la formula con l'impiego dell'aggettivo indefinito "tutto"; quindi l'interpretazione del sintagma oltrepassa la necessità del riferimento alla grammatica della lingua, poiché in sede di analisi in relazione al soggetto grammaticale "tutti gli amici" va fatto notare che esso è anche la versione linguistica di un simbolo logico, cui fa riferimento chi parla per costruire conoscenza nel mondo. Ecco il motivo di questa sovrapposizione delle analisi. Essa serve a far capire che, ogniqualvolta in italiano ci sia bisogno di ricorrere alla totalità di un insieme di individui tramite aggettivazione indefinita, il modo più comune di farlo è quello di impiegare appunto l'aggettivo indefinito "tutto" (es. "tutti i ragazzi", "tutti i professori", "tutti i candidati", "tutti i ricercatori"). Per anticipare un'obiezione, si possono menzionare altri

modi validi per riferirsi ad un insieme intero (in lingua italiana). In sede di generalizzazione, quando cioè si vuole predicare qualcosa riguardo a una certa categoria, la modalità più comune è quella di esordire nelle frasi con sintagmi-soggetto del tipo “i docenti”, “i ragazzi”, “i candidati”. E anche in questi casi siamo di fronte all’impiego di maschile non marcato, per la presenza dell’articolo determinativo maschile “i”. Si tratta – è vero – di generalizzazioni di carattere meno marcato rispetto a quelle che si realizzano nelle frasi con sintagmi della ben nota forma “tutti gli x”, ma dal punto di vista semantico tra “i docenti sono pazienti” e “tutti i docenti sono pazienti” le due frasi possono dirsi esprimere lo stesso contenuto.

2.2 Conseguenze possibili

Dopo aver condotto una disamina attenta degli strumenti linguistici adottabili in italiano, per servirsi della quantificazione logica dei predicati, è stato visto che la realizzazione della cosa non è immaginabile senza il ricorso (quasi obbligato quindi) al maschile non marcato. E dal momento che uno degli ambiti più ricchi d’importanza per il linguaggio è quello di prestarsi come utile strumento di ragionamento per l’essere umano, la frequenza con cui si ricorre all’argomentazione è difficilmente commensurabile. Una volta ammesso ciò, immaginare quanto sia parimenti frequente e indispensabile ricorrere alla quantificazione universale risulta forse impresa ancora più ardua: riferirsi cognitivamente al “tutto” tramite la lingua – potremmo dire – è uno dei modi fondamentali di interagire con la realtà stessa. Identificati, però, i meccanismi secondo cui questa procedura avviene più frequentemente in lingua italiana, vale a dire il ricorso al maschile non marcato, si profilano naturalmente i tratti del problema da considerare, il quale può essere definito nelle righe che seguono: servirsi della quantificazione universale in lingua italiana comporta obbligatoriamente l’utilizzo del maschile plurale non marcato, fatta eccezione per i casi in cui i soggetti sono entrambi al femminile, oppure quando l’attenzione è intenzionalmente indirizzata a insiemi di soli referenti femminili (“tutte le poetesse”, “tutte le ragazze”, “le professoresse”, ecc...). Questi esempi appena visti, dal punto di vista linguistico, sono marcati. Cognitivamente, la persona che li proferisce non intende insiemi di referenti di genere misto (maschile e femminile). Trattandosi infatti di un

femminile marcato, in italiano questi sintagmi rappresentano necessariamente insieme di referenti unicamente femminili; quindi non sono utilizzabili allorché occorre riferirsi alla realtà in senso generico, o potremmo anche dire neutrale.

Di primo acchito, sarebbe facile immaginarsi le conseguenze che la tal cosa comporti sul piano cognitivo e sociale, ma quello che interessa la presente trattazione è il problema del paradosso linguistico che viene ingenerato dalla problematica sollevata: che linea di comportamento deve adottare quindi chi scrive? Perché, dal punto di vista della norma, l'utilizzo del maschile plurale non marcato non è affatto percepito come violazione (del sistema grammatica); pertanto non è causa di squalifica sociale se impiegato nell'eloquio. Allo stesso tempo però esso non è sufficiente a rappresentare la totalità dei gruppi sociali che costituiscono una società, perché di essa, e conseguentemente del mondo, se ne dà una versione maschile attraverso il linguaggio. Al contempo, e qui si realizza il problema, scegliere di non adoperare il maschile plurale, nella fattispecie non marcata, adottando espedienti linguistici del tutto leciti e previsti dalla lingua italiana (es. la dittologia "ragazzi e ragazze", i nomi di professione al femminile "sindaca", "medica") per rendere inclusivo l'atto dell'eloquio, comporta di fatto prendere un'altra via rispetto all'abitudine inveterata del maschile non marcato. Non si parla quindi di errore, nel senso che non si riscontra alcuna violazione di regole strutturali o anche grammaticali della lingua, (Serianni, Antonelli, 2017: 226-228). *Sic stantibus rebus*, potremmo concludere che la cagione della problematica legata all'utilizzo del maschile non marcato che si realizza, scegliendo di divergere dalla norma consapevolmente, ha una natura intrinsecamente paradossale, poiché dal punto di vista di chi aderisce alla norma, tale maschile è grammaticalmente lecito, poiché fa parte di una norma riconosciuta dai e dalle scriventi italiani; allo stesso tempo dal punto di vista di chi è attento alle questioni di genere è perfettamente lecito, dato che grammaticalmente corretto in lingua italiana fare utilizzo delle dittologie già citate, dei nomi di professioni al femminile, per includere, nel più ampio senso di rappresentare le categorie di genere che non si riconoscono nella sclerotizzazione culturale derivante dal corredo antropologico del maschile (cioè l'immaginario comune legato al maschile). Infine, una volta tirate in ballo tutte le caratteristiche del problema in questione, e configurata la necessità di proce-

dere a implementare le raccomandazioni per un utilizzo non sessista della lingua, chi parla italiano è gettato in una *conditio* paradossale, dato che entrambi i punti di vista sono “corretti”. Ma osservando la faccenda da vicino si nota che l’uno, quello che guarda al maschile non marcato come norma, ha la ragione dalla sua unicamente secondo un punto di vista linguistico, l’altro, quello che tiene ad un uso non sessista di una lingua, ha la ragione dalla sua non solo dal punto di vista della lingua, ma anche di quello dell’etica.

3. Il lato semiotico e simbolico in rapporto ai costituenti linguistici

Quanto è stato possibile osservare nei paragrafi precedenti, ha permesso di impostare la questione del problema del maschile plurale non marcato nell’ambito di analisi linguistico e logico. Tuttavia, la prospettiva di tale studio risulterebbe non sufficiente per la presente indagine, se non venisse considerato un ulteriore ambito di ricerca: quello semiotico e simbolico. Ma in che modo può essere analizzata in una prospettiva semiotica la problematica del maschile non marcato? Ebbene, si tratta di partire dalla considerazione di un fatto assiomatico nella disciplina del segno:

Ogni persona, ogni oggetto, ogni elemento naturale o artificiale del nostro paesaggio, ogni forza o organizzazione «comunicano» continuamente. Comunicare in questo caso vuol dire semplicemente diffondere informazioni su di sé, presentarsi al mondo, avere un aspetto che viene *interpretato*, magari tacitamente, da chiunque sia presente (Volli, 2003: 3).

In altre parole, possiamo concludere che ogni aspetto e oggetto della e nella nostra vita non sono in grado di non comunicare; pertanto è lecito includere in questo insieme anche i dispositivi linguistici di una certa lingua (l’italiano nel nostro caso), cosa che fa capo ancora una volta alle parole di Ugo Volli:

Ma dappertutto non vi sono se non persone o cose che significano (se non altro, al minimo, che significano per me la loro mancanza di utilità e di interesse). Più le guardo e le studio, naturalmente, più il loro senso si arricchisce, in un processo che appare inesauribile (*Idem*).

A partire appunto da questi assiomi presi in considerazione, occorre però delineare le modalità in cui poter applicare un'indagine di tipo semiotico a costituenti sintattici di una lingua, quali ad esempio i sintagmi soggetto o complemento con maschile plurale non marcato, per essere studiati come oggetti dotati di significazione, cioè aventi ricchezza di senso.

3.1 Nelson Goodman: sistemi costruzionali e teorie simboliche

Avendo quindi illustrato i termini secondo cui è possibile dire che ogni oggetto del mondo è dotato di significazione, a fornire maggiore profondità a tale fatto possono offrire un notevole supporto le opere del filosofo statunitense Nelson Goodman (1906-1998). Esse investono ambiti plurimi, estendendosi dai domini della logica fino all'estetica, campi questi che fanno eco alla sua attività di docente presso l'università di Harvard. Nelle sue ricerche figura la teorizzazione della nozione di sistema costruzionale, la quale, secondo l'architettura teorica che ne dà il filosofo di Somerville, risulta, insieme al valore del simbolo³, di particolare interesse e utilità per il presente studio. Di queste due nozioni accennate, Goodman parla all'interno di *Ways of world making* e *Languages of art*, opere in cui delinea il modo in cui è possibile per l'essere umano "costruire" conoscenza, processo che prenderà il nome di "costruzione di mondo". Il filosofo parla di un vero e proprio creare mondi, intesi come strutture di conoscenza atte a rappresentare il mondo reale (o anche fisico); e la caratteristica innovativa (oltre che sorprendente) di questa teoria è che, valendo la proporzionalità "tanti i mondi quanti i sistemi simbolici esistenti", consegue che il soggetto vive in una pluralità simultanea di tali mondi. Ebbene, come forse intuibile, tra i sistemi costruzionali figura il linguaggio umano. Esso è un sistema in grado di produrre conoscenza, pur accettando che non si tratta dell'unico sistema possibile: l'arte e la

³ Simbolo, secondo Goodman, vale: ciò che assume significato nel momento in cui riferisce (o meglio denota) qualcosa. Tutto è simbolo a seconda del sistema di riferimento: in musica potrebbe essere simbolo una nota scritta sul pentagramma; in un museo una scultura; l'immagine in un quadro; nel linguaggio umano una parola, o anche una frase, come allo stesso tempo anche un significato può essere simbolo; ecc...

scienza ad esempio ne costituiscono altri due modelli. Entrambi infatti, dice Goodman, sono in grado di costruire cognitivamente concetti, poiché i sistemi richiedono l'impiego di rappresentazione simbolica, grazie alla quale è possibile garantire, o meglio produrre, una "versione" del mondo, distinta dalla realtà del mondo fisico, in altre parole come il mondo è per me che osservo tramite il sistema di simboli che seleziono.

In conclusione, dato che dire "conoscenza" per Goodman è dire "mondo", e che quest'ultimo si costruisce attraverso sistemi di diverso tipo che operano con degli strumenti altrettanto diversi, ma che sono universalmente "simboli", consegue che anche un sintagma di una frase in qualsiasi lingua esistente è considerabile simbolo, come risulta dalla seguente argomentazione:

1) (Come visto nella prima parte di questo terzo paragrafo), tutto è dotato di significazione, tutto è dotato di senso, cioè non può non comunicare.

2) I sistemi costruzionali si costituiscono di simboli attraverso i quali è possibile produrre conoscenza.

.'. Se tutto ha significazione (nell'accezione semiotica), allora anche i simboli hanno significazione.

Come ulteriore conseguenza, costruire conoscenza è garantito attraverso la combinazione di simboli, i quali a loro volta sono portatori di significazione.

3.2 Sistema-costruzionale-lingua

Nell'impossibilità di dedicare ulteriore spazio all'approfondimento e all'esposizione dei temi appena menzionati, la bontà delle due nozioni, quella di sistema costruzionale e di simbolo, occorre a sostegno della presente trattazione, poiché il sistema lingua, nel senso generalizzato di linguaggio umano, è facilmente sovrapponibile a quella di sistema costruzionale. Infatti, volendo dare uno sguardo su come opera il sistema lingua (considerato come sistema costruzionale), si può dire che esso è in grado di creare conoscenza o anche

una realtà coerente attraverso simboli a partire dalla traduzione di relazioni fra concetti, giudizi e proposizioni in veste linguistica *stricto sensu*. A questo punto, uno schema in grado di fornire un quadro di corrispondenze di strutture elementari fra sistema costruzionale e sistema linguistico generale, potrebbe senz'altro essere d'aiuto:

Constructional system:

1. Oggetti primitivi;
2. Termini;
3. Regole di costruzione.

Ai quali possono corrispondere senza difficoltà le componenti di un sistema linguistico generale:

1. Insieme lessico = {fonologia $\wedge$ morfologia};
2. Insieme semantica = {riferimento $\wedge$ rappresentazione};
3. Insieme della sintassi.

Giunti a questo punto della trattazione, si possono operare delle ulteriori considerazioni specifiche sul sistema-costruzionale-lingua: attraverso i simboli di cui esso è composto (parole, significati, regole) è in grado di operare riferimenti. E che rilievo ha tale affermazione nel presente studio? Ci permette di riprendere e integrare quanto è stato oggetto di analisi nei primi due paragrafi: la norma che adotta l'italiano nel dare veste linguistica alla quantificazione universale, in altre parole come tradurre in italiano il segno $\forall$. Essa è stata esaminata da un punto di vista linguistico e logico-semantico. Manca appunto il terzo: analizzare questa regola (o meglio norma) secondo una prospettiva che la identifica come simbolo, nel senso generale di oggetto dotato di significazione, come informa la semiotica, e allo stesso tempo la prospettiva del simbolo per Goodman, in cui esso è visto come parte costitutiva di un sistema costruzionale, in grado di caratterizzare una "versione" del mondo.

3.3 L'aspetto dei mondi

Come è stato visto, la conoscenza è il risultato di un'attività di ordinamento, organizzazione e manipolazione di un certo insieme di

simboli. Questi, a seconda del sistema entro cui si opera, hanno fattezze e funzionalità diversi, ma tutti concorrono alla rappresentazione del mondo con il quale il soggetto si interfaccia. Abbiamo anche detto che in italiano le modalità, previste per dare veste linguistica al simbolo logico, obbediscono a una norma che prevede il ricorso al maschile non marcato; quindi, doversi riferire alla totalità di un insieme implica in italiano la prevalenza del maschile.

Ora, nel corso della trattazione sono stati illustrati tutti i possibili punti di vista adottabili, per indagare la problematica di una norma linguistica, che è considerata corretta da una certa percentuale di parlanti, ma allo stesso tempo ingiusta da un'altra percentuale di parlanti italiano. Quello linguistico ci informa che la norma del maschile non marcato è accettata dalla comunità linguistica, ma contemporaneamente ingiusta, perché il maschile risulta veicolo di un ricettacolo di stereotipi di genere, nei quali grandi quantità di persone all'interno di una comunità non si riconoscono; quello logico-semantico informa della natura delle circostanze in cui viene utilizzata maggiormente la norma in questione: vale a dire dare forma linguistica al simbolo, stante a indicare tutti gli elementi di un insieme; il terzo punto di vista (quello semiotico-costruzionale) invece informa delle conseguenze, perché combinando le conclusioni a cui arrivano i primi due metodi di analisi, permette di elaborare un quadro definitivo sulla faccenda.

La norma del maschile non marcato, in quanto regola all'interno di un sistema costruzionale, è un simbolo, che come ogni oggetto nel mondo, ha significazione (perciò ricchezza di senso). Parlando, però, di sistema costruzionale, e specificamente di una regola sintattica per la formazione di frasi all'interno del sistema-costruzionale-lingua-italiana, sto dicendo che attraverso tale simbolo sono in grado di costruire conoscenza, che sarà da considerare una "versione del mondo", non come il mondo effettivamente è. Come è stato detto, a sistemi diversi corrispondono rappresentazioni, mondi (cioè conoscenze, secondo Goodman) diversi; quindi il simbolo che si sceglie è responsabile dell'aspetto del mondo in cui si "vive", cioè al modo in cui si concepisce la realtà in cui si vive. E quale può essere la conseguenza di un sistema costruzionale che adotta come simbolo produttivo al suo interno una norma, che è ricettacolo di stereotipi di genere orientati in senso maschilista-patriarcale? La risposta è semplice: data la stretta dipendenza di "versione del mondo" dal

sistema simbolico di partenza, una lingua⁴, cioè un sistema in grado di costruire mondi, in cui persistono dei tratti fortemente non inclusivi, non potrà fare altro che costruire un mondo sessista, all'interno del quale non è previsto alcuno spazio per l'inclusione.

Considerazioni finali

Un ultimo dato che può essere aggiunto alla presente trattazione è cercare di inserire nell'inquadramento delle teorizzazioni proposte un piccolo accenno alla relazione che vige tra linguaggio e realtà e alle conseguenze che essa comporta. A fronte dell'analisi linguistica *stricto sensu* il maschile non marcato risulta la norma corrente per dare forma linguistica al segno, fatto per cui risulta abbastanza evidente la sua indispensabilità. Frasi quantificate universalmente consentono all'essere umano di predicare (quindi stabilire relazioni tra un soggetto e un oggetto) che abbiano una validità duratura nel tempo; o meglio ancora, dare vita a inferenze, le cui estensioni verofunzionali siano sempre valide nel tempo. In altre parole si sta dicendo che tali tipologie di inferenze danno vita alla conoscenza necessaria alla vita della persona, come risuona nelle parole che seguono:

Regularities, expressed by general statements such as 'all Fs are Gs', are very valuable things to know when it comes to planning your day, driving a car, or just trying to survive for a while. Without knowledge of regularities, everything would always come as a complete surprise: be unexpected. It is easy to imagine how likely it would be that we all unintentionally killed ourselves within a couple of days without the capacity to have beliefs of regularities (possibly all by the same method, being unable to learn from the failure of others), given that our nosiness can overrule the guidance of our natural instincts (Cohnitz, Rossberg, 2006: 28).

Ma poiché è stato detto che la “versione del mondo” dipende dal sistema costruzionale scelto per produrre conoscenza; e ancora, poiché è stato detto che il linguaggio umano è un sistema costruzionale a sua volta, la “versione del mondo” che si dà in un dato linguaggio dipende dalle fattezze del linguaggio stesso. Se si applica

⁴ Si intenda una qualsiasi lingua, ma nel nostro caso l'italiano.

l'inferenza della precedente argomentazione per guardare da vicino le dinamiche atte a creare inferenze in lingua italiana, allora possiamo dire che in un ambito specifico per la creazione di conoscenza nella vita, vale a dire la fabbricazione di proposizioni quantificate universalmente, l'italiano prevede quasi esclusivamente il ricorso al maschile non marcato (es. "tutti i candidati", "tutti i vincitori", "tutti i ragazzi"). E la conseguenza di tutto ciò è la creazione di una conoscenza tramite simboli intimamente connessi a forti componenti sessiste; pertanto secondo quanto visto ricorrere al tutto in italiano secondo la norma del maschile plurale non marcato vuol dire "maschilizzare" la realtà.

Bibliografia

AA. VV. (2012). *Grammatica italiana*. Milano: Garzanti Libri.

Bloomfield Leonard (1996). *Il Linguaggio*. Milano: EST.

Cohnitz Daniel, Rossberg Marcus (2006). *Nelson Goodman*. Chesham: Acumen Publishing Limited.

Coşeriu Eugen (1971). *Teoria del linguaggio e linguistica generale: sette studi*. Bari: Laterza.

Dardano Maurizio, Trifone Pietro (1995). *Grammatica italiana. Con nozioni di linguistica*. Milano: Zanichelli.

Goodman Nelson (1978). *Ways of worldmaking*. Indianapolis: Hackett Pub Co Inc.

Goodman Nelson (2017). *I linguaggi dell'arte*. Milano: Il Saggiatore.

Lombardi Vallauri Edoardo (2024). *Le guerre per la lingua*. Torino: Giulio Einaudi editore.

Sabatini Alma (1993 [1987]). *Il sessismo nella lingua italiana*. Roma: Presidenza del Consiglio dei Ministri.

Santi Flavio, Viale Matteo (2012). *La grammatica italiana*. Roma: Istituto della enciclopedia italiana fondata da Giovanni Treccani.

Schwarze Christoph (2009). *Grammatica della lingua italiana*. Roma: Carocci.

Serianni Luca, Antonelli Giuseppe (2017). *Manuale di linguistica italiana. Storia, attualità, grammatica*. Milano: Pearson.

Serianni Luca, Castelvechi Alberto, Patota Giuseppe (a cura di) (2003). *Italiano*. Milano: Garzanti Libri.

Volli Ugo (2003). *Manuale di semiotica*. Roma: Laterza.

Per un linguaggio inclusivo e accessibile nella Pubblica Amministrazione

Virginia Laconi, virginia.laconi2@unibo.it
Università di Bologna

Introduzione¹

Preservare le diversità e riconoscerle come valore aggiunto, piuttosto che identificarle come fattore discriminante o di disparità sociale, è un tema centrale nella società attuale, che può essere applicato a diverse sfere lavorative e sociali, tra cui quella del linguaggio.

La lingua ha un forte impatto sul modo in cui noi rappresentiamo la realtà: contribuisce, infatti, per prima a costruire il mondo che ci circonda. L'uso che facciamo della lingua diventa pertanto fondamentale nel raggiungimento della piena parità sociale (Politecnico di Torino, 2024) e per questo, è importante imparare ad utilizzare un linguaggio accessibile e inclusivo, che non escluda nessuno.

I principali organi istituzionali che si occupano di comunicazione con la più ampia cittadinanza sono le Pubbliche Amministrazioni, che producono annualmente innumerevoli documenti di varia natura, in materia di servizi sociali, culturali ed economici. Tuttavia, ancora oggi, si riscontrano diverse problematiche nel linguaggio utilizzato dagli enti pubblici, a causa di un uso improprio della lingua e di alcuni suoi termini, sia dal punto di vista dell'inclusione che dell'accessibilità linguistica, la quale va a discapito della comprensione dei documenti e crea ulteriori limiti e discriminazioni.

Accessibilità e inclusione sono, quindi, due concetti chiave tanto a livello nazionale quanto europeo.

1. Il concetto di “accessibilità”

Il termine “accessibilità” ha diverse sfumature, la più diffusa è sicuramente quella di “accessibilità dei dati” in ambito informatico. Con la legge numero 4 del 2004², infatti, l'accessibilità viene definita come segue:

¹ La data di ultima consultazione dei siti citati in nota è il 6 febbraio 2025.

² <http://qualitapa.gov.it/sitoarcheologico/relazioni-con-i-cittadini/>

La capacità dei sistemi informatici, nelle forme e nei limiti consentiti dalle conoscenze tecnologiche, di erogare servizi e fornire informazioni fruibili, senza discriminazioni, anche da parte di coloro che a causa di disabilità necessitano di tecnologie assistive o configurazioni particolari. (Normattiva, 2004: articolo 2, comma 1, lettera a)

Allo stesso modo, l'enciclopedia Treccani definisce l'accessibilità come:

Proprietà che devono possedere le applicazioni per essere utilizzate con facilità dagli utenti, in particolare da coloro che si trovano in condizioni di disabilità o di svantaggio [...] per garantire che le funzionalità applicative dei servizi pubblici informatizzati siano accessibili al più ampio numero possibile di cittadini (Treccani, 2023).

Se sostituiamo le applicazioni e i servizi pubblici informatizzati con la lingua e il linguaggio, arriviamo a una possibile definizione di accessibilità in ambito linguistico, cioè l'uso di una lingua che possa essere fruita e capita con facilità dalla più ampia cittadinanza.

Esiste, infatti, una distinzione tra "accessibilità tecnica", ovvero l'adozione di soluzioni che garantiscono all'utenza *web* l'assenza di barriere ostacolanti la lettura e il reperimento di informazioni (Sacchi, 2024), e l'accessibilità cosiddetta "semantica", che deve invece garantire una comprensione il più completa e chiara possibile dell'informazione data (*idem*).

Affinché l'informazione sia accessibile semanticamente, due tipologie diverse di linguaggio entrano in gioco in linguistica: il linguaggio semplice e il linguaggio chiaro.

Si parla di linguaggio semplice o facile quando si adotta una strategia di scrittura vicina all'uso della lingua parlata: una sintassi lineare, parole facili, una grafia pulita e spaziosa e immagini che servono a spiegare meglio il contenuto del testo (FLO – Facile da Leggere Online, nd). Questo è il tipo di linguaggio da scegliere se ci si rivolge anche a persone con difficoltà cognitive.

Si parla invece di linguaggio chiaro per indicare la scelta di termini comuni e trasparenti, in ottica di divulgazione, al posto di termini tecnici più opachi e non comprensibili per un pubblico non esperto. È in questa seconda casistica che si inserisce il fenomeno della "semplificazione del linguaggio amministrativo".

Il linguaggio chiaro è già diffuso, con il nome di *plain language*, nei paesi anglofoni (Eagleson, 2023), dove esistono guide dettagliate sull'uso preferenziale di certi termini, così come di un determinato ordine degli elementi della frase per renderla il più chiara e fluida possibile. Ne è un esempio la guida "*Federal Plain Language Guidelines*" stilata dall'associazione americana *Plain Language Action and Information Network* (PLAIN) che si occupa di promuovere l'uso del *plain language* in tutte le comunicazioni amministrative del paese (Plainlanguage.gov, 2023).

Le pubbliche amministrazioni italiane stanno cercando di implementare il "linguaggio chiaro" allineandosi alle direttive europee nel settore, come testimonia la direttiva sulla semplificazione del linguaggio amministrativo dell'8 maggio 2002³.

Il Ministro della Funzione Pubblica desidera, con questa direttiva, contribuire alla semplificazione del linguaggio usato dalle amministrazioni pubbliche per la redazione dei loro testi scritti. Le amministrazioni pubbliche utilizzano infatti un linguaggio molto tecnico e specialistico, lontano dalla lingua parlata dai cittadini che pure ne sono i destinatari. [...] I numerosi atti prodotti dalle pubbliche amministrazioni, sia interni (circolari, ordini di servizio, bilanci) sia esterni, devono prevedere l'utilizzo di un linguaggio comprensibile, evitando espressioni burocratiche e termini tecnici (Ministro per la Pubblica Amministrazione, 2002).

2. Il concetto d'"inclusione"

Parallelamente all'uso di un linguaggio più chiaro e accessibile, un importante obiettivo della Pubblica Amministrazione è quello di rendere i propri testi più inclusivi, utilizzando un linguaggio non discriminatorio.

Alcuni paesi si sono già mobilitati per promuovere politiche di inclusione e uguaglianza tra genere maschile e femminile, introducendo convenzioni linguistiche da attuare nei testi amministrativi e statali. Tra questi paesi ci sono Italia e Francia, dove iniziative relative a questo argomento sono presenti dagli anni 1970-1980 (Vecchiato, 2004). Anche a livello europeo, nel 2018 sono stati pubblica-

³ <https://www.funzionepubblica.gov.it/articolo/dipartimento/08-05-2002>

ti due importanti documenti: *Una comunicazione inclusiva al Segretariato Generale del Consiglio dell'Unione Europea*, pubblicato dal Consiglio dell'Unione Europea⁴ con lo scopo di promuovere l'inclusione e la diversità, mediante l'uso della lingua e delle immagini; *La neutralità di genere nel linguaggio* pubblicato dal Parlamento Europeo⁵, il cui scopo «non è di limitare gli estensori dei testi del Parlamento europeo, obbligandoli a seguire una serie di norme vincolanti, ma piuttosto di incoraggiare i servizi amministrativi a tenere in debita considerazione la questione della sensibilità di genere nel linguaggio ogni qualvolta un testo viene redatto o tradotto, come pure in sede di interpretazione» (Parlamento Europeo, *online*, 2018).

Altre istituzioni internazionali che hanno scelto e che promuovono l'uso di un linguaggio inclusivo e privo di stereotipi sono, per esempio, l'ONU con il suo *Orientations pour un langage inclusif en français*⁶ e il Consiglio d'Europa nelle sue *Recommandation CM/Rec (2019)1 du Comité des Ministres aux États membres sur la prévention et la lutte contre le sexisme*⁷ (Haddad, 2023: 53).

È però importante sottolineare che il linguaggio inclusivo non riguarda esclusivamente la questione di genere, bensì deve tenere conto di tutti i fattori discriminanti. Il termine “inclusività” indica infatti l'eliminazione di stereotipi, discriminazioni, forme di esclusione e di violenza — che spesso derivano da un uso improprio della lingua stessa — nei confronti di minorità e categorie sociali. I fattori di discriminazione possono essere infiniti, ma tra quelli riconosciuti dalla normativa civile italiana in materia antidiscriminatoria, oltre al genere, troviamo l'origine etnica, il credo, l'orientamento sessuale, l'età e le disabilità (Regione Emilia-Romagna: 3). Tutti questi aspetti devono essere tenuti a mente durante la redazione di un documento amministrativo.

3. Il progetto E-MIMIC

In questo quadro generale, diffondere e promuovere una comunicazione inclusiva partendo dal linguaggio amministrativo, primo punto di contatto tra la cittadinanza e le istituzioni, diventa una

⁴ <http://europa.eu/!cG68fH>

⁵ https://www.europarl.europa.eu/cmsdata/288144/GNL_Guidelines_IT-original.pdf

⁶ <https://www.un.org/fr/gender-inclusive-language/guidelines.shtml>

⁷ <https://rm.coe.int/CoERMPublicCommonSearchServices/>

sfida importante, tanto interessante quanto urgente (Attanasio, 2021); ed è in questo contesto che si inserisce il progetto europeo E-MIMIC (*Empowering Multilingual Inclusive ComMunICation*), che vede la collaborazione fra un'*équipe* linguistica e un'altra informatica provenienti da tre università italiane, l'Università di Bologna, il Politecnico di Torino e l'Università di Roma Tor Vergata.

L'obiettivo principale del progetto è quello di eliminare gli stereotipi e le discriminazioni di genere ancora oggi insite nella lingua e soprattutto nei discorsi (Raus, 2016). Per fare questo si avvale dell'intelligenza artificiale, grazie alla quale è stato possibile sviluppare un dispositivo innovativo *Inclusively*, basato su tecnologie di reti neurali e *deep learning*. Si tratta di un applicativo che sarà in grado di riconoscere automaticamente gli elementi non inclusivi e non accessibili di un testo amministrativo, proponendo poi all'utenza delle riformulazioni inclusive e accessibili, eliminando quindi le forme inappropriate sulla base delle informazioni linguistiche inserite dal personale linguista (Politecnico di Torino, 2024).

Nei paragrafi seguenti viene presentata l'attività di ricerca svolta presso una Pubblica Amministrazione italiana nell'ambito del progetto E-MIMIC, con particolare attenzione alla terminologia amministrativa e all'accessibilità semantica di quest'ultima.

4. L'attività di ricerca presso Città Metropolitana di Torino⁸

L'ente Città metropolitana di Torino si è da sempre mostrato molto sensibile alle tematiche dell'inclusione e del contrasto alle discriminazioni. Nel suo statuto⁹ troviamo, infatti, i seguenti principi a cui l'ente ispira le sue attività:

d) promuovere il superamento di ogni discriminazione o disuguaglianza e consentire uguali opportunità per tutti, senza distinzione di genere, orientamento sessuale, credenza religiosa, convinzione filosofica, razza o etnia, opinioni politiche condizioni economiche e sociali, e in presenza di disabilità, tenendo al pieno sviluppo delle persone e della famiglia anche se svantaggiate e garantendo pari dignità alle minoran-

⁸ Un sentito ringraziamento al personale di Città metropolitana di Torino che mi ha accolta per svolgere uno *stage* di ricerca in sede da aprile 2024 a fine settembre 2024.

⁹ http://www.cittametropolitana.torino.it/speciali/2015/statuto_citta_metro

ze linguistiche del territorio, nell'ambito delle funzioni esercitate. (Città metropolitana di Torino, 2015: 1, punto 5, lett. d)

e) perseguire la realizzazione della parità di genere, adottando azioni positive idonee ad assicurare pari opportunità per tutti, [...] promuovendo azioni e politiche specifiche anche attraverso la collaborazione con altri enti, istituzioni e con l'associazionismo per agire sulle cause culturali e sociali del fenomeno (*ibidem*: art. 1, punto 5, lett. e).

f) favorire la realizzazione della parità di genere adottando in tutti gli atti dell'amministrazione, compresi i regolamenti, l'uso del linguaggio nel rispetto del genere (*ibidem*: art. 1, punto 5, lett. f).

Inoltre, la regione Piemonte ha stipulato una legge regionale per l'attuazione del divieto di ogni forma di discriminazione e della parità di trattamento nelle materie di competenza regionale (legge regionale n.5 del 2016¹⁰), grazie alla quale sono stati fissati alcuni principi generali relativi alla parità di trattamento tra le persone e al divieto di ogni forma di discriminazione in molteplici ambiti. L'articolo 12 della suddetta Legge istituisce anche la "Rete regionale contro le discriminazioni", con compiti di prevenzione e contrasto delle discriminazioni e assistenza alle vittime, di cui l'Ufficio Pari Opportunità e Contrasto alle Discriminazioni della Città Metropolitana di Torino è parte integrante. Quest'ultimo, per tutte le ragioni menzionate, si è sempre mostrato molto interessato alle tematiche del progetto E-MIMIC e di conseguenza ai risultati del progetto stesso.

Per questo motivo è stato possibile svolgere uno *stage* di ricerca di sei mesi presso l'ente torinese¹¹ e, in particolare, all'interno dell'Ufficio Pari Opportunità e Contrasto alle Discriminazioni.

Lavorare quotidianamente all'interno di una Pubblica Amministrazione si è rivelato interessante per comprendere il flusso lavorativo che si cela dietro la produzione di documenti amministrativi di vario genere, ma anche per individuare il ruolo di una persona esperta in linguistica, e di un applicativo come *Inclusively*, all'interno di un ente pubblico.

L'obiettivo principale è stato raccogliere materiale autentico prodotto dall'ente per:

¹⁰ <http://arianna.consiglioregionale.piemonte.it/iterlegcoordweb/dettaglioLegge>

¹¹ Lo stage di ricerca si è svolto precisamente dal 2 aprile al 30 settembre 2024.

- a) Analizzare il linguaggio amministrativo, quindi individuarne i tecnicismi di difficile comprensione e proporre delle riformulazioni più accessibili semanticamente, creando infine un glossario con questi termini e le rispettive riformulazioni.
- b) Utilizzare i documenti raccolti per addestrare l'applicativo *Inclusively*.

Nei seguenti paragrafi vengono descritte le fasi di lavoro svolte all'interno dell'ente, dalla ricerca documentaristica e la raccolta dati, fino alla creazione di un glossario finalizzato a rendere accessibili i tecnicismi, passando per la creazione di corpora e per l'analisi terminologica.

4.1 Ricerca documentaristica

La prima fase di lavoro è consistita nella ricerca documentaristica, ovvero nella ricerca di documenti pertinenti all'analisi linguistica da svolgere.

Questa prima fase si è sviluppata in due momenti: un primo momento di ricerca autonoma sul sito *web* dell'Ente, e un secondo momento di presa di contatto con colleghi e colleghe di Direzioni diverse dalla Direzione Istruzione e Sviluppo Sociale, di cui fa parte l'Ufficio Pari opportunità e Contrasto alle discriminazioni, che potessero fornire documenti interessanti e di vario tipo.

Inizialmente, sono stati analizzati i contenuti del sito *web* dell'Ente, che risulta diviso per canali tematici in base alle materie di competenza di ciascuna Direzione. Visitando ogni pagina dedicata a un canale tematico differente (tra cui Ambiente, Azioni integrate con Enti Locali, Istruzione e Diritto allo studio, Politiche Sociali, Viabilità, ecc.), sono stati scaricati e suddivisi in cartelle locali i documenti che risultavano più idonei e pertinenti agli obiettivi di ricerca. A tal proposito, i documenti privilegiati sono stati: bandi di concorso, comunicati, avvisi pubblici, schede progetto, pagine *web*, vademecum, rapporti e relazioni pubbliche, brochure e volantini — ovvero documenti principalmente rivolti alla cittadinanza.

È stata visitata, inoltre, la pagina dell'albo pretorio di Città Metropolitana di Torino¹², dove vengono pubblicati gli atti ammini-

¹² <https://stilo.cittametropolitana.torino.it/albopretorio/>

strativi dell'ente e da cui sono state scaricate alcune determine e decreti dirigenziali.

Parallelamente alla ricerca *online*, parte del materiale raccolto è stato scaricato dalle cartelle interne, accessibili esclusivamente ai membri dell'Ufficio Pari Opportunità e Contrasto alle discriminazioni.

In un secondo momento, come già menzionato, è stato chiesto a una persona referente per ogni Direzione dell'ente di fornire una decina di documenti tipici della propria direzione, che potessero presentare dei problemi di inclusione e/o di accessibilità. In questo modo è stata aumentata la varietà dei documenti e il contenuto degli stessi.

4.2 Raccolta dati

Al termine di questa prima fase di ricerca documentaristica autonoma e con la collaborazione del personale dipendente, sono stati raccolti circa 205 documenti di varia natura e tematiche differenti.

Nella Tab. 1 è possibile vedere la quantità di documenti ricavati da ogni Direzione.

Direzione	N. doc.	Direzione	N. doc.
Ambiente e Vigilanza Ambientale	6	Territorio, Pianificazione e Urbanistica	1
Azioni Integrate con Enti Locali	9	Trasporti e Viabilità	17
Edilizia	8	Ufficio Stampa	10
Integrazione processi finanziari e contabili	24	Centrale Unica Appalti	3
Rifiuti e Bonifiche	2	Finanza e Patrimonio	5
Risorse Umane	10	Flora, Fauna e Aree Protette	4
Sviluppo Economico	5	Istruzione e Sviluppo Sociale	99
Sviluppo Montano	1	Strategie Miglioramento e Organizzazione	1

Tab. 1. Documenti raccolti presso la Città metropolitana di Torino divisi per Direzione

Come si può notare, la direzione da cui sono stati raccolti più documenti è la Direzione Istruzione e Sviluppo Sociale (99 documenti), sia per questioni di accessibilità tecnica (accesso alle cartelle interne) e d'interazione con le colleghe presenti, sia per questioni di contenuto, poiché i temi trattati da questa direzione, come l'organizzazione scolastica, le pari opportunità, il contrasto alle discriminazioni, il servizio civile e la pubblica tutela, sono più suscettibili

alla produzione di documenti rivolti alla cittadinanza rispetto ad altre Direzioni più tecniche.

Per quanto riguarda le Direzioni restanti, come si può notare, alcune di queste hanno fornito più documenti rispetto ad altre; la quantità non è stato, tuttavia, un requisito fondamentale, poiché spesso sono risultati più idonei i pochi documenti ricevuti da alcune Direzioni piuttosto che i tanti ricevuti da altre. Infatti, tenendo conto degli obiettivi di ricerca, è stata fatta una selezione del materiale raccolto.

4.3 Creazione di corpora specialistici

Selezionando i documenti più pertinenti agli obiettivi di ricerca, sono stati selezionati 145 dei 205 documenti raccolti: alcuni risultavano troppo tecnici e poco discorsivi, o con molte tabelle che avrebbero creato problemi per la creazione di corpora e la suddivisione in segmenti per l'addestramento di *Inclusively*.

Con tali testi si è proceduto alla creazione di *corpora* tramite Sketch Engine¹³, un software disponibile *online* per la gestione dell'insieme dei testi raccolti (*corpora*) e per l'analisi linguistica di quest'ultimi.

Poiché i documenti selezionati sono di varia natura, sia dal punto di vista dei contenuti che della tipologia testuale, è stato necessario individuare un criterio di suddivisione affinché i risultati dell'analisi linguistica fossero attribuibili a una tipologia testuale piuttosto che a un'altra o a una determinata tematica.

Inizialmente si era creato un *corpus* unico di 145 documenti, con più sotto-*corpora* corrispondenti alle varie Direzioni di appartenenza dei testi. Tuttavia, questa suddivisione non ha dato risultati interessanti dal punto di vista linguistico.

Abbiamo quindi optato per un secondo criterio, che si è rivelato poi quello definitivamente adottato, che ha previsto la suddivisione del campione di documenti in due *corpora* distinti sulla base del pubblico destinatario dei singoli documenti: nonostante l'attenzione fosse stata inizialmente posta su materiale rivolto alla cittadinanza, con l'aggiunta di documenti di altre Direzioni, una parte dei documenti ricavati risultava, infatti, rivolta anche al personale dipendente dell'Ente.

¹³ <https://www.sketchengine.eu/>. L'accesso è stato possibile grazie alla licenza offerta dall'Università di Bologna.

Conseguentemente, abbiamo ritenuto opportuno creare due *corpora*, denominati “CMTO_esterno” e “CMTO_interno”, il primo concernente i documenti rivolti all'esterno dell'Ente, e quindi alla cittadinanza, il secondo i testi indirizzati al suo interno. Nella Tab. 2, riportiamo i dati dei due sotto-*corpora* (Tabella 2).

	CMTO_esterno	CMTO_interno
Contenuto	documenti di diverse Direzioni, rivolti a un pubblico esterno all'Ente	documenti di diverse Direzioni, rivolti a un pubblico interno all'Ente
N. di documenti	86	59
N. di parole	244 768	147 191
Tipo di documenti	brochure, vademecum, descrizione di progetti, rapporti, mail, bandi di concorso, comunicati stampa, pagine web	rapporti, decreti, comunicazioni interne
Pubblico target	la cittadinanza	il personale

Tab. 2. Dati dei sotto-corpora raccolti presso la Città metropolitana di Torino

Come si può notare dalla tabella, il *corpus* contenente documenti rivolti alla cittadinanza esterna è più voluminoso e rispecchia quindi i nostri obiettivi di ricerca, ovvero individuare elementi non inclusivi e non accessibili nel linguaggio amministrativo rivolto al più ampio pubblico. Nonostante ciò, è stato analizzato anche il linguaggio dei documenti interni all'ente per fare un eventuale comparazione dei dati raccolti.

4.4 Analisi terminologica

Una volta creati i due sotto-*corpora*, è stato possibile estrarre automaticamente, grazie alla funzione *Keywords* di Sketch Engine, i termini semplici e complessi che appaiono più frequentemente nei due campioni di testi.

Di seguito riportiamo un estratto del documento Excel utilizzato per l'analisi terminologica (Fig. 1). L'estrazione è stata fatta sotto forma di documento Excel per avere dei risultati ordinati e già predisposti per la creazione di un glossario.

Nel documento Excel della Fig. 1, le prime tre colonne sono state create automaticamente da Sketch Engine. Nella prima, troviamo i termini estratti automaticamente dal *corpus*, in ordine di frequenza; nella seconda, troviamo la frequenza di ciascuno di essi all'interno

name: extract_keywords
 corpus: user/virginia.laconi2/
 cmto_esterno

Termini	Frequency (focus corpus)	Frequency (reference corpus)	Categoria	Appartenenza (CMTO / italiano)	Accessibile / Non Accessibile	Proposta di riformulazione	Proposta di riformulazione 2
giudice tutelante		9194	Attore singolo	italiano standard	A.		
pubblica tutela	101	355	Attività	italiano standard	A.		
operatore volontario	100	2248	Attore singolo		A.		
decreto di nomina	111	4702	Documenti	italiano standard	A.		
amministratore di sostegno	136	9316	Attore singolo	italiano standard	N.A.	- assistente amministrativo a sostegno di persona non autosufficiente (figura che si prende cura di una persona non autosufficiente)	
città metropolitana	604	98589	Ente				
amministrazione di sostegno	87	6039	Attività	italiano standard	N.A.	assistenza per persone non autosufficienti	
tribunale di torino	86	7759	Ente	italiano standard	A.		
ufficio di prossimità	50	154	Ente	italiano standard	N.A.	- Sportello di orientamento - Sportello informativo locale - Ufficio informativo - Ufficio locale	- Sportello locale per l'orientamento giudiziario - Sportello giudiziario locale
tribunale di ivrea	44	997	Ente				
sezione decentrata	41	272	Attore collettivo	italiano standard	A.		
prova preselettiva	53	5730	Prassi	italiano standard	A.		
cancelleria tutele	36	39	Attore collettivo	italiano standard	N.A.	- Ufficio tutele del tribunale	
soggetto fragile	38	2073	Attore singolo	italiano standard	A.		
consigliera di parità	45	5985	Attore singolo	italiano standard	N.A.		
svolgimento delle prove	39	4201	Prassi	italiano standard	A.		

Fig.1. Estratto del documento Excel di analisi terminologica

del *corpus* specifico; nella terza, è indicata, invece, la frequenza di ogni termine in un *corpus* alternativo di lingua generale italiana, proposto da questo *software*, sulla base del quale è possibile capire quanto il termine in questione sia settoriale. A quest'ultimo riguardo, infatti, più il termine è presente nel corpus di lingua italiana generale proposto, meno sarà specifico rispetto al nostro corpus e quindi risulterà meno specializzato, e viceversa.

Le colonne successive della Fig. 1 sono state, invece, aggiunte da noi per categorizzare i termini: troviamo, infatti, indicazioni sulla categoria del termine (se si riferisce a un attore singolo o collettivo, a un documento, a un'attività, ecc.), sull'appartenenza del termine a un gergo proprio all'ente Città Metropolitana di Torino oppure al linguaggio standard, sulla distinzione tra termini "accessi-

bili” e “non accessibili” e, infine, sono riportate le colonne dedicate alle proposte di riformulazione dei termini individuati come “non accessibili”.

La distinzione tra termini accessibili e non accessibili è stata fatta sulla base di più fattori, basandoci sui criteri indicati nella norma ISO 704/2022¹⁴ per la formazione di nuovi termini (trasparenza, coerenza, appropriatezza, concisione, derivabilità e componibilità, correttezza linguistica).

Una volta individuati quali fossero i termini “non accessibili”, sono state fatte delle proposte di riformulazione. Al fine di proporre delle riformulazioni più trasparenti e chiare per un pubblico non esperto, si è fatto riferimento alle definizioni dei termini stessi o a contesti definitori trovati nel *corpus* da cui sono stati estratti.

Per esempio, il termine giuridico “beneficio di inventario”, parte della locuzione “accettare eredità con beneficio di inventario”, può risultare poco trasparente e, quindi, non accessibile. Come riformulazione, è stata proposta una perifrasi generale, ma più trasparente, che si rifà alla definizione del termine stesso: «È uno dei modi secondo i quali può essere compiuta l'accettazione dell'eredità. Produce l'effetto di tenere distinto il patrimonio dell'erede e quello del defunto» (Dizionario giuridico Broccardi, nd) La riformulazione proposta è stata quindi la seguente: «separazione del patrimonio della persona defunta da quello dell'erede».

Se proviamo a sostituire quest'ultima in una frase contenente il termine in questione, otteniamo il risultato riportato nella Tab. 3.

Testo originale	Riformulazione terminologica proposta
Se invece l'amministratore ha dei dubbi sulla passività dell'eredità, è bene che chieda l'autorizzazione ad accettare l'eredità con beneficio di inventario	Se invece l'amministratore ha dei dubbi sulla passività dell'eredità, è bene che chieda l'autorizzazione ad accettare l'eredità con separazione del patrimonio della persona defunta da quello dell'erede (beneficio di inventario)

Tab. 3. Esempio di riformulazione di un termine tratto dal corpus CMT0_esterno

4.5 Creazione del glossario

I termini individuati come “non accessibili” e le rispettive riformulazioni sono stati raccolti in un glossario.

¹⁴ <https://www.iso.org/obp/ui/fr/#iso:std:iso:704:ed-4:v1:en>

L'idea di crearne uno è nata dalla necessità di prendere nota di tutti quei termini che creano problemi di accessibilità e comprensione nei documenti amministrativi, e che attualmente non è possibile segnalare nella piattaforma d'inserimento dei dati utili all'addestramento dell'applicativo *Inclusively*. Trattando, però, sia l'aspetto inclusivo sia quello accessibile della comunicazione amministrativa è giusto tenerli in considerazione e addestrare il dispositivo a riconoscerli e riformularli, così come farà con gli altri elementi non inclusivi, come ad esempio l'accordo morfologico del femminile. Questo glossario verrà integrato all'applicativo in modo da rilevare termini non trasparenti e ottenere proposte di riformulazione affidabili e in linea con il materiale dell'ente.

Presentiamo un estratto del glossario nella Fig. 2, dove la prima e la seconda colonna contengono i termini individuati come "non accessibili", seguiti dalle possibili riformulazioni trovate e da un esempio tratto dal *corpus* e riformulato.

Termine	Collocazione	Proposta di riformulazione	Esempi	Esempi riformulati
amministratore di sostegno		- responsabile amministrativo (a sostegno) di persone non autosufficienti - rappresentante di persona fragile	"Di seguito le informazioni, complete di schede di sintesi, riguardo le principali istanze che, in qualità di amministratore di sostegno, vi potreste trovare nella necessità di presentare al Giudice."	"Di seguito le informazioni, complete di schede di sintesi, riguardo le principali istanze che, in qualità di assistenti amministrativi di persone non autosufficienti (amministratori di sostegno), vi potreste trovare nella necessità di presentare al Giudice."
amministrazione di sostegno		- assistenza per persone non autosufficienti - rappresentanza di persone fragili	-	
ufficio di prossimità		- Sportello locale per l'orientamento giudiziario - Sportello giudiziario locale	"collaborazione presso la sede dell'Ufficio di prossimità attraverso la presenza periodica, nell'accoglienza del pubblico per consulenze" "È continuata la collaborazione della Città metropolitana per la nascita degli Uffici di prossimità sul territorio del Tribunale di Ivrea [...]"	"collaborazione presso gli sportelli locali per l'orientamento giudiziario (Uffici di prossimità) attraverso la presenza periodica, nell'accoglienza del pubblico per consulenze" "È continuata la collaborazione della Città metropolitana per la nascita degli sportelli giudiziari locali (Uffici di prossimità) sul territorio del Tribunale di Ivrea [...]"
	cancelleria tutela	Ufficio Tutela del tribunale	"Gli utenti intervistati hanno saputo dell'esistenza dell'Ufficio Tutela di Città metropolitana più spesso tramite la cancelleria tutela del tribunale (38%)"	"Gli utenti intervistati hanno saputo dell'esistenza dell'Ufficio Tutela di Città metropolitana più spesso tramite l'Ufficio Tutela del tribunale (38%)"

Fig. 2.: Estratto del glossario terminologico che abbiamo proposto

Nelle prime due colonne si è effettuata un'ulteriore distinzione tra termini veri e propri e collocazioni. In linguistica, infatti, si tratta di una distinzione molto importante che viene effettuata sulla base di alcuni criteri sintattici e semantici.

Si parla di "collocazione" per indicare la co-occorrenza di due o più parole che tendono a presentarsi frequentemente insieme, contigue

o a distanza; si parla, invece, di “termine”, e più precisamente di “termine composto”, quando quest’ultimo non corrisponde a una parola singola (termine semplice), bensì a un gruppo di parole. Per distinguere un termine complesso da una collocazione esistono diversi criteri (Kübler, 2023), a cui abbiamo fatto riferimento:

- il criterio referenziale: se il gruppo nominale ha un solo referente nella realtà, allora si tratta di un termine;
- il criterio semantico: il senso globale del termine composto non deve essere uguale all’addizione dei sensi delle singole parole che lo compongono;
- il criterio sintattico: non è possibile trasformare né aggiungere altri elementi tra quelli che compongono un termine composto; se questo risulta possibile, si tratta allora di una collocazione.
- il criterio di traduzione: se il gruppo nominale si traduce con un termine semplice in un’altra lingua, si tratta allora di un termine e non di una collocazione.

4.6 Alcune osservazioni

In seguito alla creazione del glossario, è stato possibile individuare alcune caratteristiche comuni tra i termini non accessibili. Questi ultimi, infatti, possono essere distinti in tre categorie principali:

- a. Termini non accessibili, perché contengono sigle;
- b. Termini non accessibili, ma facilmente sostituibili con sinonimi più comuni e trasparenti;
- c. Termini estremamente tecnici e specifici a un ambito particolare, per i quali è stato più difficile trovare un termine sostitutivo semplice.

Le soluzioni adottate per la loro riformulazione variano, quindi, sulla base della categoria individuata.

La soluzione più semplice e immediata riguarda sicuramente i termini afferenti alla prima casistica, ovvero quelli che contengono sigle, poiché è sufficiente scioglierle per rendere il concetto più trasparente. Ne è un esempio la collocazione “Programma MIP” che è stata riformulata con il nome esteso, “Programma Mettersi In Proprio”.

In alcuni casi può essere utile, oltre a sciogliere la sigla, aggiungere delle informazioni per rendere il termine ancora più chiaro; è il caso di “sportello SAI”, ovvero lo “Sportello Ascolto Informazione” dell’ANFASS (Associazione Nazionale Famiglie di persone con disabilità Intellettive e/o Relazionali), per il quale, oltre a sciogliere la sigla, è stata aggiunta un’informazione sull’utenza a cui si rivolge questo sportello. La riformulazione proposta, quindi, è “Sportello Ascolto Informazione (SAI) per persone con disabilità”.

Anche nel secondo caso la soluzione adottata è semplice, ovvero sostituire i termini poco chiari con un sinonimo di comune utilizzo. Per esempio, il termine “Paesi Terzi”, nel contesto europeo indica tutti i paesi che non sono membri dell’Unione Europea, di conseguenza per rendere l’espressione più trasparente è stato riformulato con “Paesi extra Unione Europea” o l’equivalente “Paesi extra-UE”, in quanto l’acronimo UE è ormai ampiamente diffuso.

Le difficoltà principali sono sorte, invece, per i termini estremamente tecnici e specifici a un determinato ambito, in particolare all’ambito giuridico. Per questa tipologia, infatti, è stato difficile trovare un termine sostitutivo semplice o appartenente alla lingua comune. La soluzione che abbiamo adottato è stata quella di riformularli con perifrasi più ampie, ma più trasparenti. In questo modo, abbiamo sostituito il termine inizialmente utilizzato con una perifrasi seguita dal tecnicismo in modo che il pubblico possa recuperare successivamente il senso di esso quando sarà riutilizzato nel documento.

Ne è un esempio il termine “Ufficio di prossimità”. Per definizione, «dal punto di vista del sistema giudiziario sono degli sportelli che permettono ai cittadini di avere un riferimento vicino al luogo dove vivono e di usufruire di un servizio completo di orientamento e di consulenza» (Regione Piemonte, nd); in base a ciò, il termine è stato riformulato in due modi: “Sportello locale per l’orientamento giudiziario” e “Sportello giudiziario locale”.

Tuttavia, sostituire ogni occorrenza di “ufficio di prossimità” con “sportello locale per l’orientamento giudiziario” può risultare poco leggibile in un documento amministrativo. Di conseguenza, la soluzione ritenuta più opportuna è stata quella riportata nelle Tab. 4 e 5.

In questo modo, le occorrenze successive del termine “Ufficio di prossimità” all’interno dello stesso documento possono rimanere invariate, ma chi legge sa già di cosa si tratta.

Testo originale	Riformulazione terminologica proposta
collaborazione presso la sede dell'Ufficio di prossimità attraverso la presenza periodica, nell'accoglienza del pubblico per consulenze	collaborazione presso gli sportelli locali per l'orientamento giudiziario (Uffici di prossimità) attraverso la presenza periodica, nell'accoglienza del pubblico per consulenze

Tab. 4. Esempio di riformulazione terminologica tratto dal corpus CMT0_esterno

Testo originale	Riformulazione terminologica proposta
È continuata la collaborazione della Città metropolitana per la nascita degli Uffici di prossimità sul territorio del Tribunale di Ivrea	È continuata la collaborazione della Città metropolitana per la nascita di sportelli giudiziari locali (Uffici di prossimità) sul territorio del Tribunale di Ivrea

Tab. 5. Esempio di riformulazione terminologica tratto dal corpus CMT0_esterno

Il contatto diretto con il personale dipendente di Città Metropolitana di Torino si è rivelato di particolare importanza in questa fase dei lavori perché, in caso di dubbi, è stato possibile chiedere il parere di persone esperte e avere conferme sulla chiarezza e l'esattezza delle riformulazioni proposte.

È il caso del termine “tutore volontario”, che era stato inizialmente riformulato con “rappresentante legale di minore con cittadinanza straniera” sulla base della definizione stessa del termine: «I tutori volontari sono privati cittadini disponibili a esercitare la rappresentanza legale di un minorenne straniero arrivato in Italia senza adulti di riferimento» (Autorità garante per l'Infanzia e l'Adolescenza, nd). Dopo aver chiesto il parere dell'ufficio di Pubblica Tutela dell'Ente, è stato specificato che quel termine, nei loro documenti, viene utilizzato anche per riferirsi alla tutela legale di persone adulte e non solo minori. Di conseguenza, è stata aggiunta una seconda riformulazione più generica nel glossario, ossia “rappresentante legale di persona fragile”. In questo modo, i risultati generati dall'applicativo saranno più precisi e attendibili.

5. Annotazione e riformulazione in chiave inclusiva dei testi amministrativi

In parallelo all'attività di analisi linguistica della terminologia amministrativa in ottica di accessibilità, i documenti raccolti presso la Città metropolitana di Torino sono stati etichettati anche per esse-

re utili all'addestramento dell'applicativo *Inclusively*, dato che quest'ultimo è stato effettuato sia su dei corpora generici in fase di *pre-training* sia sui corpora di questo ente per il *fine tuning*, come indicato anche nella presentazione di Matteo Berta e di Tania Cerquittelli in questo volume.

In linea con gli obiettivi di ricerca, ovvero riformulare il linguaggio amministrativo rivolto al più ampio pubblico, si è cominciato a lavorare sui documenti del *corpus* "CMTO_esterno", composto esclusivamente da documenti rivolti alla cittadinanza. Questi sono stati analizzati ed etichettati, appunto, per insegnare al dispositivo quali fossero i termini e le forme frastiche non inclusive da un lato e quali invece fossero, dall'altro, le rispettive riformulazioni inclusive.

Per garantire una maggiore correttezza e affidabilità dei dati finali, la fase di annotazione è stata, infine, seguita da una fase di validazione dei *dataset* già annotati da parte di personale esperto dell'Università di Bologna e dell'Università di Roma Tor Vergata.

Conclusione e sviluppi futuri della ricerca

Per concludere, come menzionato inizialmente, uno degli obiettivi dello *stage* di ricerca svolto presso una Pubblica Amministrazione, la Città metropolitana di Torino, è stato comprendere quale potrebbe essere il ruolo di un applicativo come *Inclusively* nella quotidianità di un ente pubblico.

A tal proposito, dall'esperienza fatta emerge che, nonostante una crescente consapevolezza dell'importanza dell'uso di un linguaggio inclusivo e non discriminatorio, e il tentativo di utilizzarlo ove possibile, si fa ancora fatica ad abituarsi a un uso quotidiano di questo tipo di linguaggio. Le motivazioni sono principalmente due: da un lato, molti documenti vengono redatti sulla base di modelli in uso da tanti anni, che non prevedono l'uso di un linguaggio inclusivo; dall'altro, la mancanza di tempo a disposizione per potersi concentrare anche sull'aspetto puramente linguistico, sia nella redazione di nuovi documenti, sia nell'eventuale aggiornamento dei vecchi modelli. La presenza di un applicativo che riformuli automaticamente i documenti in chiave più inclusiva e accessibile permetterebbe una riscrittura veloce dei documenti stessi, così come la presenza di una figura che si occupa dell'aspetto puramente linguistico dei

documenti aiuterebbe a colmare le necessità laddove il personale esperto in altri ambiti si occupa del contenuto.

I lavori presentati in questo articolo sono quindi solo l'inizio di un'analisi approfondita del linguaggio amministrativo italiano, che dopo i documenti di Città Metropolitana di Torino, continuerà su un altro campione di testi ricavati da un'altra Pubblica Amministrazione regionale italiana, l'Assemblea Legislativa della Regione Emilia-Romagna. Si prevede, infine, di condurre un'analisi simile anche su *corpora* di documenti amministrativi francesi, sui quali al momento stanno lavorando altre colleghe, in ottica di operare una comparazione tra Italia e Francia, due paesi simili ma che, allo stesso tempo, risultano molto diversi dal punto di vista delle politiche linguistiche e delle tematiche trattate in questa ricerca (diversa tipologia e definizione di "testo amministrativo", diversa tradizione nel trattare la questione dell'inclusione, ecc.).

Bibliografia

Attanasio Giuseppe et alii (2021). *E-MIMIC: Empowering Multilingual Inclusive Communication*. 2021 IEEE International Conference on Big Data (Big Data), IEEE, pp. 4227–34. <https://doi.org/10.1109/BigData52589.2021.9671868>

Autorità garante per l'Infanzia e l'Adolescenza (nd). *Come diventare tutore volontario*. <https://www.garanteinfanzia.org/come-diventare-tutore-volontario-0>

Città Metropolitana di Torino. Albo pretorio. <https://stilo.cittametropolitana.torino.it/albopretorio/>

Consiglio dell'Unione Europea (2018). *Una comunicazione inclusiva al Segretariato generale del Consiglio dell'Unione Europea*. <http://europa.eu/!cG68fH>

Consiglio d'Europa (2019). *Recommandation CM/Rec (2019)1 du Comité des Ministres aux États membres sur la prévention et la lutte contre le sexisme*. <https://rm.coe.int/CoERMPublicCommonSearchServices/DisplayDCTMContent?documentId=090000168093b269>

Dizionario giuridico Broccardi (nd). *Beneficio di inventario*. <https://www.broccardi.it/dizionario/858.html>

Eagleson Robert (2023). Plainlanguage.Gov. *Short Definition of Plain Language*. <https://www.plainlanguage.gov/about/definitions/short-definition/>

FLO - Facile da Leggere Online (2023). *Che cos'è il linguaggio facile?* <https://www.faciledaleggere.online/index.html>

Haddad Raphaël et al. (2023). *L'écriture inclusive, et si on s'y mettait ?* Parigi: Le Robert.

Kübler Nathalie (2023). *Linguistica di corpus applicata alla terminologia*. Cours M2 ILTS, Université Paris Cité.

Ministro per la Pubblica Amministrazione (2002). *Direttiva sulla semplificazione del linguaggio amministrativo*. <https://www.funzionepubblica.gov.it/articolo/dipartimento/08-05-2002/direttiva-semplificazione-linguaggio>

Organizzazione delle Nazioni Unite (2018). *Orientations pour un langage inclusif en français*. <https://www.un.org/fr/gender-inclusive-language/guidelines.shtml>

Parlamento Europeo (2018). *La neutralità di genere nel linguaggio usato al Parlamento Europeo*. https://www.europarl.europa.eu/cmsdata/288144/GNL_Guidelines_IT-original.pdf

Plainlanguage.gov, 2023. *Federal plain language guidelines*. <https://www.plainlanguage.gov/guidelines/>

Politecnico di Torino (2024). *Un linguaggio più inclusivo grazie all'intelligenza artificiale*. <https://www.polito.it/ateneo/comunicazione-e-ufficio-stampa/poliflash/un-linguaggio-piu-inclusivo-grazie-all-intelligenza>

Pubblica Amministrazione di Qualità (2013). *Accessibilità. Come garantire un accesso universale a internet*. <http://qualitapa.gov.it/sitoarcheologico/relazioni-con-i-cittadini/open-government/comunicazione-istituzionale-online/portale-pubblico/internet/accessibilita/index.html>

Raus Rachele (2016). «Per una cittadinanza della lingua: promuovere la parità di genere nel linguaggio amministrativo». Piemonte Autonomie. <https://www.piemonteaautonomie.it/per-una-cittadinanza-della-lingua-promuovere-la-parita-di-genere-nel-linguaggio-amministrativo/>

Regione Emilia-Romagna sociale (nd). *Piccola guida contro le discriminazioni*. <https://sociale.regione.emilia-romagna.it/contro-le-discriminazioni/publicazioni/publicazioni-centro-regionale/piccola-guida-contro-le-discriminazioni>

Regione Piemonte (2016). *Legge regionale n.5 del 23 marzo 2016, "Norme di attuazione del divieto di ogni forma di discriminazione e della parità di trattamento nelle materie di competenza regionale"*. <http://arianna.consiglioregionale.piemonte.it/iterlegcoordweb/dettaglioLegge.do?urnLegge=urn:nir:regione.piemonte:legge:2016;5@2019-3-1>

Regione Piemonte, (nd). *Normativa. Uffici di prossimità*. <https://www.regione.piemonte.it/web/temi/fondi-progetti-europei/pon-governance/uffici-prossimita>

Sacchi Fabio (2024). *Purché siano accessibili: una desk review dei dati statistici e amministrativi italiani delle persone con disabilità*, in *L'integrazione scolastica e sociale* 90–105. <https://rivistedigitali.erickson.it/integrazione-scolastica-sociale/archivio/vol-23-n-2/purche-siano-accessibili-una-desk-review-dei-dati-statistici-e-amministrativi-italiani-delle-persone-con-disabilita/>

Treccani, Enciclopedia online (2023). *Accessibilità*. <https://www.treccani.it/enciclopedia/accessibilita>.

Vecchiato Sara (2004). *Le sexisme dans le langage. Notes sur l'italien et le français*. In Raus R. (a cura di), *Linguaggi e discriminazioni*, Torino, CIRSD, in <http://www.cirsde.unito.it>

Comunicación inclusiva en España y *Deep Learning*: adaptación metodológica del proyecto E-MIMIC al lenguaje de la administración española

Natalia Peñín Fernández, natalia.penin2@unibo.it
Università di Bologna

Introducción

El principal riesgo de utilizar soluciones basadas en IA (Inteligencia Artificial) radica en su uso irresponsable, el cual incluye la manipulación indebida de datos personales, la perpetuación de algoritmos sesgados y su capacidad para intensificar desigualdades existentes. Estos problemas no son meramente técnicos, sino también éticos y sociales. Por ejemplo, los datos sesgados reflejan patrones históricos de discriminación que los modelos de IA pueden amplificar si no se corrigen adecuadamente. Para minimizar estos sesgos, es fundamental que las aplicaciones de IA cuenten con supervisión humana en dos etapas clave: el entrenamiento del modelo y la evaluación de su desempeño.

En este contexto, la comunidad de lingüística computacional subraya la importancia de las anotaciones realizadas por expertos humanos como parte integral de los enfoques de aprendizaje experiencial (Seretan, Bouillon, Gerlach, 2014; Miyata Rei, Atsushi Fujita, 2021; Raus *et al.*, 2021). Estas anotaciones permiten identificar matices textuales complejos que los sistemas automatizados podrían pasar por alto. Los modelos de comprensión y generación de lenguaje natural se benefician significativamente de estas aportaciones al integrar conocimientos lingüísticos humanos mediante un enfoque autosupervisado, lo que permite reducir de manera notable los sesgos. Esta colaboración entre humanos y máquinas es esencial para desarrollar sistemas de IA más equitativos y responsables.

El proyecto E-MIMIC (*Empowering Multilingual Inclusive Communication*)¹, una iniciativa conjunta del Politécnico de Turín, de la Universidad de Bolonia y de la Universidad de Roma Tor Vergata, tiene como propósito principal explorar cómo pueden los avances más recientes en *Deep Learning* y PLN (Procesamiento del Lenguaje Natural) contribuir a combatir las formas discriminatorias del len-

¹ Para profundizar en los detalles del proyecto, se aconseja consultar la página *web* <http://dbdmg.polito.it/e-mimic/>

guaje mediante el desarrollo de una aplicación de apoyo a la escritura inclusiva para las administraciones públicas.

Este proyecto, fruto de la colaboración activa entre lingüistas encargados de anotar los segmentos textuales y científicos de datos que investigan el PLN, se centra especialmente en las comunicaciones formales, donde el lenguaje discriminatorio puede estar más normalizado o ser menos evidente. Para abordar estos desafíos, E-MIMIC adopta un enfoque de *Deep Learning* para analizar textos en busca de fragmentos que contengan elementos discriminatorios. Una vez identificados, propone reformulaciones intralingüísticas basadas en estándares lingüísticos inclusivos previamente aprendidos por el sistema. Este proceso no solo fomenta una comunicación más respetuosa y equitativa, sino que también sensibiliza sobre el impacto del lenguaje en la perpetuación de desigualdades.

En este artículo se describe la adaptación metodológica² de las premisas del proyecto E-MIMIC a la lengua española³.

1. La comunicación inclusiva en España: breve historia de las propuestas

En España, el análisis del lenguaje desde una perspectiva que aborda su carácter discriminatorio, especialmente en relación con el género, comienza a adquirir relevancia en la década de 1970. Este interés se ve impulsado por la publicación de *Lenguaje y discriminación social* en 1977, una obra de Álvaro García Meseguer que muchos consideran fundacional en el debate sobre la inclusividad del idioma español y la necesidad de establecer políticas lingüísticas que eliminen los usos discriminatorios (Guerrero Salazar, 2013 y 2021).

A partir de esta obra, la investigación lingüística en España durante los años 80 del siglo XX se alinea con la aparición de las primeras guías en español enfocadas en la promoción de un lenguaje no sexista, dirigidas inicialmente al ámbito administrativo. En este contexto, el Instituto de la Mujer, fundado en 1983, se convierte en un actor clave en la promoción de políticas que fomentan la igualdad lingüística.

² Desarrollo e incorporación de un conjunto de criterios lingüísticos (procesamiento sintético de datos) y discursivos que contribuyan al entrenamiento de las redes neuronales subyacentes a la traducción automática.

³ La aplicación se está ultimando para la lengua italiana y se está desarrollando para el francés.

En el ámbito político, un hito significativo en la historia del lenguaje inclusivo en España será la aprobación de la Ley Orgánica 3/2007, de 22 de marzo, para la igualdad efectiva de mujeres y hombres. Esta normativa marcó avances cruciales en los terrenos social y legislativo. En su apartado «Políticas públicas para la igualdad», se establece como principio general de actuación la adopción de un lenguaje no sexista en la administración pública y su promoción en todas las esferas sociales, culturales y artísticas. En 1990, el Instituto de la Mujer presentó la guía para el *Uso no sexista del lenguaje administrativo*, que se consolidó como un recurso fundamental tanto para entidades públicas como privadas. Este documento proporciona recomendaciones prácticas específicas para la redacción de textos administrativos y las comunicaciones institucionales. Asimismo, destacan otras guías elaboradas por el Instituto de la Mujer, que a menudo colaboran con otras instituciones en la publicación de estas recomendaciones, además de aquellas desarrolladas por ayuntamientos y diputaciones provinciales. Además, asociaciones de mujeres, sindicatos, parlamentos, empresas, universidades, etc. han elaborado guías para fomentar el uso de un lenguaje igualitario en sus respectivas prácticas laborales (Guerreiro Salazar, 2022).

También en 2011, el Instituto Cervantes publicó su *Guía de comunicación no sexista*, que destaca por su exhaustividad y por abordar de manera integral el uso inclusivo del lenguaje. Su reedición en 2021 reafirmó su vigencia y relevancia.

En las últimas décadas, movimientos como el LGTBQ+ y colectivos de personas con discapacidad han comenzado a cuestionar las prácticas lingüísticas tradicionales, proponiendo nuevas formas de expresión más inclusivas que combatan todo tipo de discriminación, ampliando el enfoque más allá del género. Este debate sigue siendo objeto de controversia, con posturas polarizadas que encuentran eco en la prensa y en las redes sociales, donde el enfrentamiento es constante.

En este contexto, la Real Academia Española (RAE), como parte de la Asociación de Academias de la Lengua Española (ASALE), defiende el uso del masculino genérico por economía lingüística a pesar de que recomienda un uso moderado en los casos en que no sea estrictamente necesario y rechaza estrategias ajenas al sistema lingüístico español (Academia, 2020).

La postura de la RAE se ha manifestado de manera significativa en tres ocasiones. La primera fue en 2012 con el informe del académico Ignacio Bosque Muñoz, que analizó nueve guías de lenguaje no sexista publicadas en España principalmente por universidades. En este informe expresó las reservas de la institución sobre ciertas propuestas estableciendo que contenían «recomendaciones que contravienen no solo normas de la Real Academia Española y la Asociación de Academias, sino también de varias gramáticas normativas, así como de numerosas guías de estilo elaboradas en los últimos años por muy diversos medios de comunicación» (Bosque Muñoz, 2012: 1). La segunda ocasión tuvo lugar en 2020, cuando, por solicitud de la vicepresidenta del Gobierno español, la RAE elaboró un informe valorando la adecuación de la Constitución española a un lenguaje que reflejara «de manera correcta y fiel la realidad de una democracia compuesta por hombres y mujeres» (Salazar, 2022: 5) en cuya conclusión afirmaba que

el lenguaje con el cual ha sido redactada la Carta Magna es “claro e inteligible”, a lo cual se agrega que la antigüedad del texto no presenta problemas importantes para su interpretación y acomodación al contexto social actual, aunque se reconoce la posibilidad de adecuar algunos sustantivos referidos a cargos y oficios unipersonales (Salazar, 2022: 4).

Finalmente, en diciembre de 2022, la RAE se distanció de las *Recomendaciones para un uso no sexista del lenguaje en la Administración parlamentaria*, acordadas en la Reunión de la Mesa de las Cortes Generales.

Desde posturas menos conservadoras, han surgido propuestas exitosas en contextos comunicativos como las redes sociales y la prensa. Entre estas estrategias destacan el desdoblamiento de sustantivos (los beneficiarios y las beneficiarias), las diversas alternativas gráficas (l@s beneficiari@s, lxs beneficiarixs, l*s beneficiari*s) el uso del neomorfema *-e* (les empleades) y el femenino sobreextendido.

Tradicionalmente el debate sobre el lenguaje inclusivo se ha centrado en el género, lo que limita su alcance a una visión binaria (hombre/mujer). De hecho, el término más comúnmente utilizado en las primeras guías ha sido ‘lenguaje no sexista’. Son, en cambio, escasos los estudios que examinan la discriminación lingüística en un sentido más amplio, considerando el lenguaje como reflejo de es-

tereotipos y prejuicios relacionados con la edad, la etnia, el origen cultural, la situación socioeconómica o necesidades específicas. A pesar de ello, existen guías como la *Guía para una comunicación sin barreras* de la Fundación IPAR HEGOA (2009) o el *Manual de lenguaje inclusivo* de COCEMFE (2022), que abordan la comunicación inclusiva hacia personas con discapacidad desde una óptica social y pedagógica. No obstante, estas aproximaciones distan mucho de integrar un enfoque propiamente lingüístico, y mucho menos una perspectiva translingüística.

En la adaptación de la metodología del proyecto E-MIMIC y de los criterios preestablecidos a la lengua de la administración española, adoptamos una definición de lenguaje inclusivo en su sentido más amplio, entendiéndolo como aquel «que promueve la no discriminación y la igualdad para todos los grupos de población que se sienten excluidos o discriminados en base a determinados usos lingüísticos» (López Fraguas, 2019). Esta perspectiva busca garantizar la inclusión y equidad lingüística para todos los colectivos afectados por prácticas comunicativas excluyentes.

Para ello, en el proyecto, seguimos igualmente las directrices establecidas en la *Comunicación inclusiva de la Secretaría General del Consejo de la Unión Europea*⁴ y la Oficina de Naciones Unidas en Ginebra sobre el uso de lenguaje inclusivo en el ámbito de la discapacidad⁵. Asimismo, se decidió emplear estrategias de neutralización adecuadas para textos administrativos y que, en general, estuvieran más alineadas a las políticas lingüísticas peninsulares y a las indicaciones de la Real Academia Española.

⁴ <https://www.consilium.europa.eu/es/documents-publications/publications/inclusive-comm-gsc/>

⁵ La Oficina de las Naciones Unidas en Ginebra ha elaborado estas directrices en cumplimiento de la Estrategia de Inclusión de la Discapacidad de las Naciones Unidas, adoptada en 2019. Esta estrategia constituye un marco fundamental de políticas y acciones diseñado para integrar de manera transversal la inclusión de las personas con discapacidad en el sistema de las Naciones Unidas. Su objetivo principal es eliminar barreras y fomentar la participación plena de las personas con discapacidad en todos los ámbitos de la vida y el trabajo, promoviendo así un avance sostenido y transformador hacia la inclusión. En particular, el indicador 15 de la Estrategia, relacionado con la comunicación, establece que tanto las comunicaciones internas como externas deben ser respetuosas y accesibles para las personas con discapacidad (ONU, 2021).

2. Metodología

La metodología propuesta por el proyecto E-MIMIC, adaptada a los textos administrativos en lengua española se divide en diferentes fases:

1. Recogida de datos. Creación de un corpus homogéneo compuesto por documentos de diferentes géneros textuales provenientes de la administración española.
2. Preparación del etiquetado de los textos españoles.
 - Identificación de las expresiones que deben evitarse por su potencial ofensivo o discriminatorio.
 - Identificación y redacción de propuestas de reformulación desde una perspectiva inclusiva.
 - Redacción de directrices de preedición.
3. Etiquetado de textos. Anotación manual de segmentos del corpus, tanto a nivel de frases como de unidades léxicas, destacando los elementos que no cumplen los criterios de lenguaje inclusivo.
4. Modelización de datos. Preentrenamiento y puesta a punto del dispositivo de *Deep Learning* para la detección de lenguaje no inclusivo, la reformulación y la generación de textos para la lengua española.

En este artículo se ofrece una descripción profunda de cada una de las tres primeras fases llevadas a cabo en el año 2024.

3. Recogida de datos

El corpus está compuesto por más de 150 textos que abarcan una amplia variedad de géneros textuales administrativos procedentes

⁶ Nuestra clasificación de los textos administrativos utilizados en el corpus se basa en el trabajo elaborado por Sara Pistola, Susana Viñuales-Ferreiro “Una clasificación actualizada de los géneros textuales de la Administración pública española” por contener la lista más reciente (2021) de géneros textuales del ámbito de la Administración española, clasificada en función de su emisor y su receptor.

de diversas instituciones públicas españolas⁶. Entre los géneros representados se encuentran:

- Resoluciones. Documentos que reflejan decisiones formales de las instituciones.
- Convocatorias. Anuncios dirigidos al público sobre eventos, cursos o procesos administrativos.
- Normativas. Reglas, leyes o disposiciones oficiales publicadas por las entidades.
- Módulos. Materiales formativos o informativos destinados a actividades específicas.
- Notas de prensa. Comunicados formales que informan al público sobre eventos, decisiones o actividades.

Las instituciones que aportaron los textos incluyen entidades académicas, administrativas y sanitarias, representan diferentes niveles de la administración pública.

- Universidades: Universidad Complutense de Madrid, Universidad de Salamanca, Universidad Rey Juan Carlos, Universidad Autónoma de Madrid.
- Instituciones sanitarias: Hospital 12 de Octubre.
- Ayuntamientos: Ayuntamiento de Madrid, Ayuntamiento de Jaén.

Los textos se clasifican dentro de la categoría "up-down", es decir, están dirigidos desde las administraciones hacia los ciudadanos. Este enfoque refleja la función de las instituciones como emisoras de información oficial.

El corpus incluye un total de 6.000 frases. Este enfoque en la segmentación y el análisis de las unidades lingüísticas — como sesgos, inclusividad lingüística o estructuras recurrentes — hace que el corpus sea relevante para investigaciones en lingüística aplicada, procesamiento del lenguaje natural (PLN) y estudios de comunicación institucional.

4. Preparación del etiquetado de los textos españoles

El análisis de los mecanismos lingüísticos presentes en los documentos administrativos en lengua española revela fenómenos discriminatorios ampliamente documentados en la literatura sobre el estudio del

lenguaje inclusivo desde una perspectiva lingüística (Bengoechea Bartolomé, 2009, 2019; Díaz Hormigo, 2009; Guijuelmo, 2019; Guerrero Salazar, 2013, 2020, 2021; Pano Alamán, 2022). Estos fenómenos se manifiestan principalmente en dos niveles: gramatical y semántico.

1. A nivel gramatical:

- Invisibilización de la mujer: El uso reiterado de formas masculinas genéricas como sustantivos, pronombres y determinantes masculinos con valor genérico puede generar confusión e interpretaciones erróneas. Aunque estas formas incluyan a las mujeres en sus referencias, en la literatura se denuncia que su uso sistemático oculta su presencia de manera simbólica.
- Asimetrías en el tratamiento lingüístico de hombres y mujeres: Se evidencian diferencias en el contenido semántico de las palabras (sexismo léxico) y en la estructura sintáctica de los enunciados (sexismo sintáctico).
- Ausencia de formas femeninas para profesiones prestigiosas: Tradicionalmente muchas denominaciones profesionales han carecido de variantes femeninas, especialmente en los cargos más destacados. Asimismo, se recurre a términos que no reflejan el género de la persona, perpetuando la ambigüedad y la exclusión de las mujeres.
- Jerarquización en el orden de mención: Los hombres suelen mencionarse en primer lugar, reflejando y reforzando un orden jerárquico.
- Reducción de las mujeres a roles relacionales: Se las menciona principalmente en función de sus vínculos con otros (como madres, esposas, etc.) o a través de tratamientos de cortesía (por ejemplo, “señora” o “señorita”).
- Asociaciones verbales discriminatorias: Se vincula a la mujer con conceptos como debilidad, pasividad, labores domésticas, histeria o infantilismo, lo que contribuye a desvalorizar su figura.

2. A nivel semántico:

- Expresiones y fórmulas fijas discriminatorias: Uso de frases hechas que perpetúan estereotipos de género.
- Representación estereotipada de colectivos sociales: Los discursos reflejan prejuicios asociados no solo al género, sino también a la edad, diversidad familiar, estado de salud, entre otros.

5. Reformulación de los fragmentos desde una perspectiva inclusiva

En una segunda fase, se procedió a reformular los fragmentos identificados como no inclusivos aplicando estrategias lingüísticas tanto gramaticales como semánticas que promueven la inclusión.

A continuación, se muestran los casos más frecuentes en nuestro corpus junto con las estrategias de neutralización empleadas con mayor recurrencia en las reformulaciones. La columna de la izquierda contiene los casos originales con los fragmentos considerados no inclusivos, mientras que la derecha presenta las reformulaciones de neutralización inclusivas propuestas en el proyecto.

5.1 En el plano léxico y sintáctico

a) Sustitución de las formas de masculino genérico con sustantivos epicenos o de género común sin el artículo:

<u>El participante</u> tiene que presentarse a la entrevista en horario.	<u>La persona participante</u> tiene que presentarse a la entrevista en horario. <u>El sujeto participante</u> tiene que presentarse a la entrevista en horario.
Es necesario ofrecer un servicio de apoyo a <u>los agredidos</u>	Es necesario ofrecer un servicio de apoyo a <u>las víctimas de agresión</u>
Podrán optar al concurso <u>los profesionales</u> con experiencia	Podrán optar al concurso <u>profesionales</u> con experiencia

b) Sustitución de las formas de masculino genérico con sustantivos colectivos o de género neutro:

Los resultados presentados en este nuevo trabajo están permitiendo a <u>los científicos</u> catalogar estos organismos	Los resultados presentados en este nuevo trabajo están permitiendo a <u>la comunidad científica</u> catalogar estos organismos
--	--

c) Sustitución de las formas de masculino genérico con sustantivos abstractos (sustituciones metonímicas):

Sus acuerdos se adoptarán por mayoría simple, decidiendo, en caso de empate, el voto de calidad <u>del presidente</u>	Sus acuerdos se adoptarán por mayoría simple, decidiendo, en caso de empate, el voto de calidad <u>de la presidencia</u>
---	--

d) Sustitución de las formas de masculino genérico con pronombres relativos o indefinidos:

La calificación de este ejercicio será de 0 a 10 puntos, <u>quedando eliminados aquellos aspirantes</u> que no alcancen una nota mínima de 5 puntos	La calificación de este ejercicio será de 0 a 10 puntos, <u>se eliminará a quienes</u> no alcancen una nota mínima de 5 puntos
---	--

e) Sustitución de las formas de masculino genérico con determinativos sin género marcado:

El formulario de solicitud se encuentra a disposición de <u>los interesados</u>	El formulario de solicitud se encuentra a disposición de <u>cada solicitante</u>
---	--

f) Sustitución de las formas de masculino genérico con fórmulas impersonales o pasivas con 'se':

<u>El participante</u> debe prestar atención	<u>Es necesario</u> prestar atención, <u>se</u> <u>prestará</u> atención
Para <u>ser admitidos</u> a las presentes pruebas selectivas, <u>los concursantes</u> deberán reunir los siguientes requisitos generales: a) <u>Ser español</u> o nacional de un Estado miembro de la Unión Europea o nacional de aquellos Estados, a los que, en virtud de Tratados Internacionales celebrados por la Unión Europea y ratificados por España, sea de aplicación la libre circulación de <u>trabajadores</u> en los términos en que ésta se halla definida en el Tratado Constitutivo de la Comunidad Europea.	Para <u>que se les admita</u> a las presentes pruebas selectivas, <u>las personas concursantes</u> deberán reunir los siguientes requisitos generales: a) <u>Tener nacionalidad española</u> o de un Estado miembro de la Unión Europea o de aquellos Estados, a los que, en virtud de Tratados Internacionales celebrados por la Unión Europea y ratificados por España, sea de aplicación la libre circulación de <u>personas trabajadoras</u> en los términos en que ésta se halla definida en el Tratado Constitutivo de la Comunidad Europea

g) Sustitución de las formas de masculino genérico con formas adjetivales:

[...] experiencia con estudiantes <u>extranjeros</u>	[...] experiencia con estudiantes <u>de nacionalidad extranjera</u>
--	---

h) Eliminación del desdoblamiento⁷:

A partir del día siguiente a la resolución, el candidato podrá solicitar su carta de admisión <u>al coordinador/ra</u> del programa a través de la propia aplicación	A partir del día siguiente a la resolución, la persona candidata podrá solicitar su carta de admisión <u>a la coordinación</u> del programa a través de la propia aplicación
adaptación del puesto de trabajo <u>a la/s persona/s afectada/s</u>	adaptación del puesto de trabajo <u>a la persona o personas afectadas</u>

i) Sustitución del masculino genérico con el femenino marcado⁸:

<u>Directores</u> : [Nombre Apellido F] (Instituto Universitario de Investigación Ortega-Marañón) y [Nombre Apellido M] (ACCEM)	<u>La directora</u> [Nombre Apellido F] (Instituto Universitario de Investigación Ortega-Marañón) y <u>el director</u> [Nombre Apellido M] (ACCEM)
---	--

⁷ Se adoptó esta estrategia considerando que la duplicación contraída puede dificultar la lectura para personas con trastornos de aprendizaje o quienes dependen de herramientas de accesibilidad, como lectores de texto, los cuales podrían no interpretar correctamente estas construcciones.

⁸ Con respecto a los cargos y profesiones Mariottini y Palmierini confirman que "Con respecto a la dimensión del uso y de la norma compartida, destacamos, por parte de los

5.2 En el plano semántico

Los siguientes casos observan fenómenos lingüísticos menos estudiados hasta ahora desde una perspectiva lingüística. Se trata de expresiones discriminantes y estereotipadas que perpetúan mecanismos de discriminación en cuanto al género, la edad, la composición familiar, el tratamiento de las discapacidades o de las enfermedades.

a) Roles de género estereotipados: Se reformularon expresiones que representaran normas, expectativas o ideas preconcebidas sobre cómo deben comportarse, actuar, sentir o contribuir las personas en función de su género, es decir, si son hombres o mujeres.

De esta manera, mientras el nivel de riesgo sea bajo, todas las personas que accedan a este tipo de instalaciones podrán salir de la habitación para pasear por las áreas establecidas, si los médicos o enfermeras así determinan

De esta manera, mientras el nivel de riesgo sea bajo, todas las personas que accedan a este tipo de instalaciones podrán salir de la habitación para pasear por las áreas establecidas, si el personal médico y de enfermería así determinan

b) Discriminación por edadismo: Se reformularon expresiones que reflejaban prejuicio o estereotipo hacia las personas por su edad cronológica, especialmente cuando ésta es elevada, o a través de términos despectivos que connotan pasividad y/o dependencia.

En todo caso, y si el desplazamiento es estrictamente necesario y está dentro de los supuestos previstos en el Real Decreto por el que se declara el estado de alarma, aconseja a los ancianos y a los enfermos que eviten, por motivos de salud, utilizar el Metro y el transporte público en general

En todo caso, y si el desplazamiento es estrictamente necesario y está dentro de los supuestos previstos en el Real Decreto por el que se declara el estado de alarma, aconseja a las personas mayores y a quienes padecen alguna enfermedad que eviten, por motivos de salud, utilizar el Metro y el transporte público en general

c) Discriminación por diversidad familiar: La reformulación tiene en cuenta la diversidad de estructuras familiares y las posibles asimetrías de corresponsabilidad familiar incluyendo la posibilidad de la inexistencia de una familia en la vida de una persona.

Como novedad, ahora podrán ser beneficiarias también todas las asociaciones que trabajen para familias con otras dificultades, como las de más de tres hijos, con un solo padre o las de padres adoptivos

Como novedad, ahora podrán ser beneficiarias también todas las asociaciones que trabajen para familias con otras dificultades, como las numerosas, monoparentales o las adoptantes

A ellos se añaden estrategias dirigidas a mejorar la información y comunicación con los jóvenes y sus padres

A ellos se añaden estrategias dirigidas a mejorar la información y comunicación con jóvenes y sus figuras parentales

Proper, que cerrará el sábado estas divertidas jornadas para toda la familia.

Proper, que cerrará el sábado estas divertidas jornadas para todos los públicos.

hablantes españoles, atención, interés y actitud positiva hacia el cambio lingüístico, que es reflejo de un cambio de sensibilidad social y cultural” (2022: 335).

La organización remitirá a los Centros que obtengan plaza un modelo de Autorización <u>paterna/materna</u> y una ficha médica individual que se deberá entregar junto con la demás documentación antes de las 12:00 horas del viernes 12 de enero de 2024.	La organización remitirá a los Centros que obtengan plaza un modelo de Autorización <u>de la figura parental</u> y una ficha médica individual que se deberá entregar junto con la demás documentación antes de las 12:00 horas del viernes 12 de enero de 2024
d) Discriminación en ámbito sanitario: Se eliminaron estereotipos negativos así como expresiones estigmatizadoras o pasivizantes.	
3.000 litros de leche donada pasteurizada que podrán beneficiar cada año a más de 600 <u>bebés prematuros o recién nacidos enfermos</u> de la Comunidad de Madrid [...]	3.000 litros de leche donada pasteurizada que podrán beneficiar cada año a más de 600 <u>bebés con una enfermedad o que han nacido de forma prematura</u> de la Comunidad de Madrid [...]
e) Discriminación en el ámbito de las discapacidades: El criterio predominante es sustituir los estereotipos negativos o el lenguaje estigmatizador sustituyendo expresiones con fuerte connotación negativa por soluciones inclusivas.	
Dado el carácter excepcional y único de esta convocatoria, que no podrá reiterarse en el futuro, las plazas reservadas al turno de <u>discapacitados</u> que no llegaran a cubrirse se acumularían al turno general.	Dado el carácter excepcional y único de esta convocatoria, que no podrá reiterarse en el futuro, las plazas reservadas al turno de <u>personas con discapacidad</u> que no llegaran a cubrirse se acumularían al turno general.

6. Los retos de la reformulación inclusiva neutra

Las dificultades que han surgido durante este proceso merecen una reflexión. En primer lugar, hay que señalar los desafíos que plantea la falta de equivalencia semántica entre sustantivos colectivos y abstractos y el masculino genérico en algunos casos concretos. Términos como ‘infancia’ o ‘niñez’ solo en casos muy restringidos equivalen al plural ‘los niños’ y podrán sustituirlo en un discurso oral o escrito. Esta falta de equivalencia evidencia algunas de las limitaciones del lenguaje para expresar inclusividad y precisión en ciertos contextos.

Asimismo, el uso de estructuras pasivas impersonales, muy comunes en español, presenta un valor comunicativo diferente al de las estructuras activas con sujeto explícito. Por ejemplo, frases como “Los aspirantes accederán a los lugares de realización de los ejercicios...” contrastan en claridad y estilo con su alternativa impersonal: “Se accederá a los lugares...”. Este matiz exige una cuidadosa evaluación para preservar la precisión y claridad del mensaje.

Igualmente, no se puede ignorar la necesidad de cumplir con los principios generales de redacción de textos administrativos: claridad, precisión, uniformidad, sencillez y economía. Aunque son objetivos fundamentales, su implementación efectiva puede ser compleja, especialmente cuando se busca equilibrar estos principios con las demandas de inclusividad y corrección formal.

Finalmente, la repetición abusiva de ciertas estructuras en el texto administrativo puede dificultar la lectura y comprensión. Este aspecto debe ser atendido mediante estrategias de reformulación y variación que mantengan la coherencia sin comprometer la fluidez del documento.

Estas reflexiones no agotan el tema, sino que representan un punto de partida para continuar explorando soluciones lingüísticas que mejoren tanto la calidad como la eficacia de los textos administrativos en español.

Conclusiones

El uso de la Inteligencia Artificial en el ámbito de la lingüística computacional y el procesamiento del lenguaje natural presenta oportunidades relevantes para mejorar la inclusión y la equidad en los textos administrativos. Sin embargo, el principal desafío radica en la manipulación irresponsable de los datos, lo que puede perpetuar sesgos históricos y generar discriminación. Para mitigar estos riesgos, es fundamental contar con supervisión humana en la creación y evaluación de los modelos de IA, como se evidencia en el enfoque del proyecto E-MIMIC, que pone énfasis en la colaboración entre lingüistas y científicos de datos para desarrollar un lenguaje inclusivo. El corpus utilizado en este proyecto, compuesto por textos administrativos en español, muestra la necesidad de reformular los textos de manera inclusiva para evitar fenómenos de discriminación de género, edad y otras formas de exclusión social. Las estrategias de reformulación, tanto a nivel léxico como semántico, demuestran que es posible transformar los textos administrativos en documentos más inclusivos y respetuosos, utilizando un enfoque que contemple la diversidad de las personas y sus realidades sociales.

A pesar de los avances, existen retos significativos, como la falta de equivalencia semántica entre algunos términos y la dificultad para

mantener la claridad y precisión en el lenguaje administrativo al introducir cambios inclusivos. Estos desafíos requieren un análisis continuo y la implementación de estrategias que logren un equilibrio entre la inclusividad y los principios fundamentales de la redacción administrativa.

El proyecto E-MIMIC, al aplicar IA en el análisis y reformulación de textos administrativos, representa un paso importante hacia una comunicación institucional más equitativa y responsable. No obstante, es necesario seguir investigando y perfeccionando las metodologías de inclusión lingüística, con el fin de crear sistemas de IA más justos y conscientes de la diversidad, contribuyendo así a una sociedad más inclusiva y respetuosa.

Bibliografía

- Bengoechea Bartolomé Mercedes (2005). "Sexismo y androcentrismo en los textos administrativo-normativos". Alcalá: Universidad de Alcalá. Recuperado de <https://drive.google.com/file/d/0B1OhCa48TcPeNGNOdVFyZkUtVVU/edit>
- Bengoechea Bartolomé Mercedes (2019). "El informe de la RAE: sexismo lingüístico y visibilidad de las mujeres, un texto político". *Revista con la A*, 61, 1-4.
- Bengoechea Bartolomé Mercedes (coord.) (2009). *Efectos de las políticas lingüísticas, antisexistas y feminización del lenguaje en los medios (2006-2009)*. Madrid: Instituto de la Mujer. Recuperado de <http://www.inmujer.gob.es/areasTematicas/estudios/estudioslinea2010/docs/efectosPolíticasLinguistas.pdf>
- Bosque Muñoz Ignacio (2012): "Sexismo lingüístico y visibilidad de la mujer". Boletín de Información Lingüística (BILRAE), RAE. Recuperado de http://www.rae.es/sites/default/files/Bosque_sexismo_linguistico.pdf
- COCEMFE (2022). *Manual de lenguaje inclusivo*.
- Díaz Hormigo María Tadea (2009). "Androcentrismo social, discriminación lingüística y propuestas para un uso igualitario de la lengua". En Fuentes Catalina y Alcaide Esperanza (eds.): *Manifestaciones textuales de la descortesía y agresividad y verbal en diversos ámbitos comunicativos*, Sevilla: Universidad Internacional de Andalucía, 98-117.
- Fundación Ipar Hegoa (2009). *Guía para una comunicación sin barreras*.
- García Meseguer Álvaro (1977). *Lenguaje y discriminación sexual*. Barcelona: Montesinos.
- Grijelmo Álex (2019). *Propuesta de acuerdo sobre el lenguaje inclusivo*. Madrid: Taurus.
- Guerrero Salazar Susana (2013). "Las guías de uso no sexista del lenguaje editadas en castellano por las universidades españolas (2008-2012)". En R. Palomares Perraut (ed.), *Historia(s) de mujeres en homenaje a M.ª Teresa López Beltrán*. Málaga: Perséfone. vol. I, 118-132. Recuperado de http://www.aehm.uma.es/persefone/Homenaje_Maite1_ISBN2.pdf
- Guerrero Salazar Susana (2020). "El debate social en torno al lenguaje no sexista en la lengua española". *IgualdadES*, 2, 191-211. Doi: <https://doi.org/10.18042/cepc/IgdES.2.07>
- Guerrero Salazar Susana (2021). "El discurso metalingüístico sobre 'mujer y lenguaje' en la prensa española: Análisis del debate lingüístico y su repercusión social". En S. A. Flores Borjabad & R. Pérez Cabaña (coords.), *Nuevos retos y perspectivas de la investigación en Literatura, Lingüística y Traducción*, Madrid: Dykinson, 1630-1646.

Guerrero Salazar Susana (2022). "Repercusión mediática del informe de la RAE sobre el lenguaje inclusivo en la Constitución española". *CLAC*, 89, 1-17.

López Fraguas Isabel (2019). "Lenguaje inclusivo, comunicación no sexista, género". *Puntoycoma*, 163, 7-18.

Mariottini Laura, Palmerini Monica (2022). "El léxico femenino de las profesiones en español e italiano. Un modelo de análisis entre sistema, norma y habla aplicado al caso de ministra". *Annali — sezione romanza*, LXIV, 1, Università di Napoli L'Orientale.

Miyata Rei, Atsushi Fujita (2021). "Understanding pre-editing for black-box neural machine translation". En *Proceedings of the 16th conference of the European chapter of the association for computational linguistics*, 1539-1550. <https://aclanthology.org/2021.eacl-main.132.pdf>

Pano Alamán Ana (2022). "La imagen de las mujeres en las campañas institucionales del Instituto Andaluz de la Mujer en redes sociales". *PRAGMÁTICA SOCIOCULTURAL*, 10, 5-25.

Parlamento Europeo (2018). *Gender-neutral language in the European Parliament*. https://www.europarl.europa.eu/cmsdata/151780/GNL_Guidelines_EN.pdf

Pistola Sara, Viñuales-Ferreiro Susana (2021). "Una clasificación actualizada de los géneros textuales de la Administración pública española". *Revista de Lengua i Dret, Journal of Language and Law*, 75, 181-203.

Raus et al. (2021). "E-MIMIC: "Empowering Multilingual Inclusive Communication"". En IEEE International Conference on Big Data (Big Data). <https://ieeexplore.ieee.org/document/9671868>

Real Academia Española (2020). *Informe de la Real Academia Española sobre el lenguaje inclusivo y cuestiones conexas*. Recuperado de https://www.rae.es/sites/default/files/Informe_lenguaje_inclusivo.pdf

Seretan Violeta, Bouillon Pierrette, Gerlach Johanna (2014). "A large-scale evaluation of pre-editing strategies for improving user-generated content translation". En *Proceedings of the 9th edition of the language resources and evaluation conference (LREC)*, 1793-1799. http://www.lrec-conf.org/proceedings/lrec2014/pdf/676_Paper.pdf

Communication inclusive et apprentissage profond : adaptation méthodologique du projet E-MIMIC à la langue de l'administration française*

Michela Tonti, michela.tonti@unibg.it
Università di Bergamo

Martina Ailén García, martinaailen.garcia2@unibo.it
Università di Bologna

Introduction

Au sein du projet *Empowering a Multilingual Inclusive Communication* (E-MIMIC), le personnel de l'École Polytechnique de Turin coordonne des personnes expertes en science des données ou en informatique, ainsi que des équipes linguistiques (en italien, français et espagnol), afin de créer le dispositif *Inclusively* qui propose des reformulations d'éventuels segments non inclusifs dans les textes de l'administration. L'inclusion est entendue ici dans le sens large de la non-discrimination concernant le genre, le handicap, l'âge, l'origine, etc.

Dans cette étude, nous nous focaliserons sur la version française d'*Inclusively*, en considérant le seul français de France, et sur l'importance de prédisposer de corpus *ad hoc* et de données linguistiques qualitativement élevées pour l'entraînement d'un modèle supervisé d'IA qui puisse respecter avant tout la diatopie linguistique (Raus, Tonti, 2025).

Le présent article s'organise comme suit : d'abord, nous illustrerons la conceptualisation sous-jacente à la constitution d'un corpus de textes administratifs de la fonction publique française, ainsi que les difficultés survenues lors du repérage et de la récolte des documents pour la création d'un corpus de travail. Ensuite, nous présenterons les critères linguistiques à la base des solutions morphosyntaxiques et sémantiques retenues afin de prédisposer des annotations qui permettent au dispositif *Inclusively* de proposer des corrections des segments non inclusifs. Dans la majorité des cas, la non-inclusion de ces derniers est due à la présence du masculin générique, ce qui caractérise le langage de l'administration française.

* L'introduction et les sections 1 et 2 ont été rédigées conjointement par Michela Tonti et Martina Ailén García. La section 3 et les conclusions ont été rédigées par Michela Tonti, qui a aussi révisé le français de l'article.

1. Notre parcours heuristique pour créer un corpus ad hoc

Afin de réunir nos données pour la construction du corpus, nous sommes parties de la définition suivante de « document administratif » en France, que l'on peut définir de la manière suivante :

Sont considérés comme administratifs tous les documents produits ou reçus par une administration publique (administrations d'État, collectivités territoriales, établissements publics). Il en va de même pour les documents détenus par les organismes privés chargés d'une mission de service public pour autant qu'ils sont liés, par leur nature, leur objet, ou leur utilisation à la gestion de cette mission. L'article L. 311-2 du code retient la notion de rattachement à une mission de service public pour qualifier la nature administrative des documents. (<http://www.cada.fr/particulier/le-document-est-il-administratif>)¹.

Une fois défini le type de document recherché, nous avons mis en place des critères de sélection des documents pertinents pour la création du corpus français en nous appuyant sur les pages institutionnelles du site officiel français d'information et de démarches administratives *service-public.fr* et sur celles de la Commission d'accès aux documents administratifs (CADA)².

La fiche pratique classée par événement de vie : « Papiers – Citoyenneté – Élections », dénommée « Accès aux documents administratifs » et publiée par le site *service-public.fr*, nous informe qu'un document administratif peut prendre une forme écrite, d'enregistrement sonore ou visuel ou apparaître sous forme numérique ou informatique³. De surcroît, l'article 300-2, issu de la loi 2016-1321 du 7 octobre 2016 rappelle comme suit :

Constituent de tels documents notamment les dossiers, rapports, études, comptes rendus, procès-verbaux, statistiques, instructions, circulaires, notes et réponses ministérielles, correspondances, avis, prévisions, codes sources et décisions (https://www.legifrance.gouv.fr/codes/texte_lc/LEGITEXT000031366350).

¹ Tous les sites consultés ont été vérifiés le 31 janvier 2025.

² Cette commission est une autorité administrative indépendante chargée de veiller à la liberté d'accès aux documents administratifs et aux archives publiques ainsi qu'à la réutilisation des informations publiques. Son site est le suivant : <http://www.cada.fr>

³ <http://www.service-public.fr>

Le site CADA précise quels sont les documents qui ne sont pas considérés comme administratifs : les ordonnances, les jugements, les juridictions administratives et financières, les documents de nature judiciaire, d'état civil ou produits depuis plus de 75 ans, ainsi que les documents privés (actes immobiliers ou actes notariés) ou les actes des assemblées parlementaires.

Comme la fonction publique française brasse des typologies textuelles ciblées, nous avons songé à cerner la variété de typologies textuelles par rapport aux critères cités ainsi qu'en fonction de leur facilité d'accès et repérage. Par conséquent, les textes publiés dans le *Bulletin Officiel* n'ont pas été retenus, car, étant des textes législatifs publiés et approuvés, ils ne pouvaient pas être reformulés de manière inclusive⁴. En outre, nous avons veillé à ce que chaque corpus soit géographiquement situé en France métropolitaine, car le projet E-MIMIC se focalise sur le français de l'Hexagone. Par conséquent, les textes relatifs aux Départements et Régions d'Outre-Mer et Collectivités d'Outre-Mer (DROM-COM) ont été omis.

Si l'on s'en tient donc à l'administration française métropolitaine, trois macro-catégories d'organismes publics énonciateurs ont été prises en compte : l'État, y compris les universités, les collectivités territoriales et les hôpitaux. Il s'agit de trois organismes administratifs reconnus qui fournissent à la citoyenneté des documents administratifs accessibles directement à partir de leurs sites web respectifs.

Procès-verbaux, comptes rendus et règlements ont été retenus à condition de dater au plus tard de 2020 afin de rassembler un corpus récent, au net des temps prévus pour la publication des actes et des documents administratifs, qui s'élèvent normalement à deux mois. Pour proposer un large éventail d'institutions, nous avons puisé dans l'annuaire du site officiel de la République française qui précise les critères administratifs et juridiques pour repérer les établissements publics⁵.

⁴ Voir la circulaire du ministre Édouard Philippe du 21 novembre 2017 relative aux règles de féminisation et de rédaction publiée au Journal officiel de la République française, au lien suivant <https://www.legifrance.gouv.fr/jorf/id/JORFTEXT000036068906>

⁵ www.annuaire-entreprises.data.gouv.fr

2. Problèmes de construction d'un corpus français pour l'apprentissage de l'IA

Dans cette section, nous énumérons les principaux problèmes survenus lors de la constitution du corpus français qui sont les suivants :

1. l'« imprévisibilité » des textes ;
2. le format des documents ;
3. les limites des plateformes d'où l'on télécharge les textes ;
4. l'incohérence des titres avec le contenu réel des documents.

2.1 L'« imprévisibilité » des textes

L'une des principales difficultés relevées lors de la phase de création du corpus français a été l'« imprévisibilité » des textes (Paveau, 2017), à savoir l'accessibilité apparente et la difficulté de consulter des documents pertinents pour les objectifs de la recherche, contrairement à ce qui devrait être normalement garanti à la communauté citoyenne par une administration transparente. C'est le cas de certaines municipalités qui obligent les usagers à créer un compte pour consulter certains documents, comme, par exemple, les comptes rendus des réunions du conseil municipal. Les documents qui devraient être en libre accès pour le grand public se révèlent être en accès restreint. En outre, il faut souligner que, selon le site *service-public.fr*, pour les communes dont la densité de population est inférieure à 3500 habitants, le conseil municipal est libre de choisir et de modifier à tout moment le mode de publication des décisions syndicales et municipales.

2.2 Le format des documents

En droit français, la multimodalité des formats des documents administratifs est autorisée et reconnue comme telle. Preuve en est que les sites des municipalités ne permettent pas le téléchargement de leurs documents dans un format unique. Tantôt les documents sont scannés en PDF, tantôt ils sont disponibles sous forme de texte

en HTML, ou sous forme de vidéo sans qu'une pièce jointe soit prévue. Cela rend l'information difficilement accessible aux personnes malvoyantes ou malentendantes, étant donné l'absence de sous-titres ou de fichiers pouvant être lus par des lecteurs d'écran.

Il s'avère également que la présence de documents de grande taille ou de fichiers volumineux pose des problèmes de consultation à cause du nombre de pages : le seul moyen d'y parvenir est de surfer en ligne, tout téléchargement de fichier étant exclu. Des images au format JPEG sont parfois utilisées pour fournir l'instantané d'un rapport manuscrit. Par ailleurs, la présence de nombreux tableaux a également été un motif de rejet de textes, car le dispositif de segmentation des documents utilisé dans le projet E-MIMIC⁶ ne parvient pas à les traiter de manière automatique.

2.3 Les limites des plateformes concernées

Chaque institution met en ligne ses propres documents administratifs dans son site institutionnel au lieu de les afficher dans une plateforme partagée, comme cela pourrait l'être la plateforme existante *data.gouv.fr*. La recherche et l'analyse au cas par cas de chaque site institutionnel a ralenti de beaucoup le travail de recherche. Dans le cas des organismes municipaux, par exemple, la typologie de documents que nous avons retenus est classée dans la catégorie « assemblées municipales », et, pour les musées et les hôpitaux, dans la section « règlements intérieurs ».

Même la présence d'une coopération intercommunale, telle l'Intercommunalité, qui pourrait simplifier la recherche et mieux répondre aux besoins des usagers ne réussit pas, pour l'instant à faciliter la tâche de consultation des documents.

2.4 L'incohérence des titres avec le contenu réel des documents

L'incohérence de l'étiquetage des fichiers a représenté un ralentissement supplémentaire dans la phase de collecte des textes exploi-

⁶ Ce dispositif a été réalisé exprès en Python par l'École Polytechnique de Turin et est mis à disposition des équipes du projet par la plate-forme *Label Studio* (<https://labelstud.io>).

tables : en effet, le titre du fichier ne correspondait souvent pas au type de texte repéré dans le document. En guise d'exemple : le fichier comportant dans le titre la mention « compte-rendu », s'est avéré être, après téléchargement, un procès-verbal⁷. A l'égard de ce type de document, nous signalons qu'à partir du 1er juillet 2022, la réforme faite par la Direction Générale des Collectivités locales (DGCL), l'agence gouvernementale chargée de respecter les règles juridiques appliquées aux Collectivités locales, a stipulé que les comptes rendus des réunions municipales sont remplacés par la Liste de Délibérations, une liste à puces qui contient l'ordre du jour de chaque session afin de fournir au public des informations plus rapides et plus directes sur l'organe délibérant. Par conséquent, la quantité de comptes rendus dans le corpus semble être inférieure à celle des règlements et des procès-verbaux.

2.5 Traitement préalable du corpus récolté

Bien que la phase de collecte des textes dans des formats adaptés à la plateforme E-MIMIC ait pris plus de temps qu'elle n'aurait dû en raison des problèmes mentionnés, nous avons quand même réussi à créer un corpus répondant aux exigences du projet, à savoir un corpus de 411780 *tokens* et 145 textes. Une fois rassemblé ce corpus, nous avons procédé aux étapes suivantes :

a) Polissage

Pour que la machine accepte les textes, ceux-ci doivent être réaménagés. Ainsi, les textes comportant des segments mal coupés ont été nettoyés. La police et la taille du texte, l'espacement et le positionnement du contenu dans la page ont été ajustés, les listes à puces ou numérotées ont été éliminées, car la machine doit être entraînée par segments, plutôt que par paragraphes.

b) Anonymisation

Les textes doivent être anonymisés afin de respecter la vie privée et l'individu. C'est pourquoi, lors de l'annotation, tous les noms et

⁷ Voir les documents du site de la commune de Devecey : <https://www.devecey.fr/page/les-comptes-rendus-du-conseil>

titres ont été remplacés et les patronymiques ont été anonymisés, comme dans l'exemple suivant : M. Martin Dupont [Prénom_M] [Nom]. L'intégration de l'élément _M et _F, à savoir l'initial de masculin/féminin, a pour but de préserver l'inclusion de genre.

Pour des raisons pratiques et de temps, ce remplacement a été effectué à l'aide des raccourcis et des fonctions de Word qui permettent de remplacer de manière automatique plusieurs mots en même temps dans le même texte. Parfois cela n'a pas été possible puisque certains textes contenaient plusieurs termes d'adresse pour la même personne (Madame, Mme, Monsieur, M.). Lorsque des abréviations étaient présentes et qu'il était possible de connaître s'il s'agissait d'un homme ou d'une femme dans la suite de la phrase, l'initiale du nom a été remplacée par l'expression [Prénom_M], s'il s'agissait d'un homme, et par [Prénom_F], s'il s'agissait d'une femme. Dans d'autres cas, en revanche, lorsque le genre était indéchiffrable, l'abréviation a été omise et remplacée par [Nom].

3. Les critères de reformulation de la version française d'*Inclusively*

Après avoir cerné et récolté le corpus, nous avons défini les critères linguistiques de reformulation inclusive dans la lignée de ceux que les équipes italienne et espagnole ont mise en place pour ces deux langues traitées également par le dispositif *Inclusively*. Nous allons donc présenter ces critères tout en précisant d'emblée que l'écriture inclusive est un sujet fort débattu en France et que déjà cette appellation, choisie parmi d'autres (par exemple, « langage inclusif ») puise ses racines dans la langue-culture de ce pays. En l'occurrence, l'écriture inclusive a été définie en France, comme l'« ensemble des attentions graphiques et syntaxiques qui permettent d'assurer une égalité des représentations des deux sexes » (Haddad, 2017 : 4). Cependant, par-delà tous les débats que cette définition, et la question de l'inclusion plus généralement ont suscité en France (cfr. Raus *et alii*, 2022 ; Raus, Tonti, 2025), il nous a fallu essayer d'accorder les démarches européennes de communication inclusive qui ont inspiré le projet E-MIMIC dans toutes les versions linguistiques considérées avec les politiques françaises (notamment la Circulaire du 21 novembre 2017 relative aux règles de féminisation et

de rédaction des textes publiés au *Journal officiel*), ce qui a donné au final les critères que nous énumérons de suite.

3.1. Flexion du genre des noms de personnes

En français, l'identité de genre des personnes et le genre grammatical, féminin ou masculin, sont étroitement associés. Si le masculin générique est imposé par la circulaire du 21 novembre 2017 relative aux règles de rédaction des textes publiés au *Journal officiel*, néanmoins des tentatives d'écriture inclusive se manifestent sous de multiples formes qui posent parfois problème à nos yeux car il s'agit de solutions plus excluantes qu'inclusives.

3.1.1 Les noms collectifs pour corriger la double flexion partielle ou doublets abrégés

Nous classons la double flexion partielle ou doublets abrégés parmi les stratégies peu inclusives. Il s'agit d'un processus langagier consistant à ajouter une marque morphologique de genre féminin à un mot déjà fléchi au masculin et utilisé en emploi générique. Cette marque est signalée par l'insertion de signes typographiques comme le point, le point médian, les parenthèses et/ou les barres obliques.

Ex. : Une séance sera consacrée plus particulièrement aux doctorant.e.s comparatistes, avec la participation d'un ou d'une directeur.rice de recherche en littérature comparée, afin d'aborder plus spécifiquement les questions de méthodologie propres à cette spécialité. (Procès-verbal du conseil d'ED du 20 février 2023, Université Paris Sorbonne).

Nous avons exclu cette solution ayant recours au point bas car inefficace du point de vue de la lisibilité (aussi, mais pas seulement, pour les personnes malvoyantes) et de la visibilité du féminin ; une solution envisagée pour remédier à ce manque de visibilité est celle que nous proposons dans le Tab. 1 :

Forme non inclusive circulant dans notre corpus	Forme envisagée en tant qu'inclusive
Une séance sera consacrée plus particulièrement aux doctorant.e.s comparatistes, avec la participation d'un ou d'une directeur.rice de recherche en littérature comparée.	Une séance sera consacrée plus particulièrement à la communauté doctorante des comparatistes, avec la participation du personnel de direction de recherche en littérature comparée.

Tab. 1. Reformulation du point médian

La collocation « communauté doctorante », composée de la base « communauté » et du collocatif « doctorante » reçoit une spécialisation du sens grâce au collocatif qui sémantise de manière univoque l'état, le caractère de ce qui est commun à plusieurs personnes étant inscrites à un cursus doctoral. Il s'agit d'une solution inclusive car le nom collectif « communauté » permet de contourner le souci de l'accord. Le nom de personne « doctorant » fait l'objet d'une conversion en adjectif.

3.1.2 Les noms collectifs pour corriger la double flexion totale ou doublets complets

Le doublet complet est constitué de la forme masculine d'un mot ainsi que de la forme féminine correspondante, coordonnée généralement par la préposition « et », comme dans l'exemple suivant :

Ex 1. : Le téléphone portable peut être utilisé à des fins pédagogiques à la demande des formateurs et des formatrices. (Institut de formation IFSI IFAS Règlement intérieur, juillet 2021).

Ou encore :

Ex 2. : Comme la demande avait été faite par certains de nos concitoyens et concitoyennes [...]. (Procès-verbal Ville de Roanne, séance du jeudi 4 mai 2023).

Nous estimons que les désignations masculines et féminines coordonnées posent problèmes dans l'ensemble de la phrase et de la préservation de l'inclusion au vu de la tendance attestée à garantir l'accord uniquement au masculin comme dans notre exemple.

Le pronom indéfini au pluriel « certains » est notamment accordé au masculin. Afin de pallier cette visibilité lacunaire accordée au féminin et pour privilégier des désignations neutres, nous envisageons les solutions proposées dans le Tab. 2 :

Forme non inclusive circulant dans notre corpus	Forme envisagée en tant qu'inclusive
Le téléphone portable peut être utilisé à des fins pédagogiques à la demande des formateurs et des formatrices.	Le téléphone portable peut être utilisé à des fins pédagogiques à la demande du personnel formateur.
Comme la demande avait été faite par certains de nos concitoyens et concitoyennes [...].	Comme la demande avait été faite par certaines personnes de notre corps citoyen [...].

Tab. 2. Reformulation de la double flexion

L'introduction du nom collectif « personnel » suivi du collocatif spécificateur, à savoir « formateur » et du nom collectif « corps » sémantisé par le spécificateur « citoyen » s'avère efficace aux fins de l'inclusion.

3.1.3 Féminisation du genre pour les postes à « impact social »

L'invisibilisation des femmes lorsqu'elles sont affectées à un poste de responsabilité ou de prestige est à bannir. Si en général nous prôtons l'épicénisation en tant que processus linguistique visant à rendre le langage plus inclusif, l'exemple suivant pose un problème de visibilité des rôles professionnels au féminin :

Ex : La coordination de la formation et son fonctionnement sont assurés par Mme Sylvie DUVAL responsable la coordination de l'institut. (Centre hospitalier Lézignan-Corbières, Règlement intérieur du 28 juillet 2023).

La perception du terme épïcène « responsable » désignant une personne qui a la charge d'une fonction de responsabilité est celle d'invisibiliser une position professionnelle valorisante. Nous proposons la solution suivante dans le Tab. 3 :

Forme non inclusive circulant dans notre corpus	Forme envisagée en tant qu'inclusive
La coordination de la formation et son fonctionnement sont assurés par Mme [Prénom_F] [Nom] responsable de la coordination de l'institut.	La coordination de la formation et son fonctionnement sont assurés par Mme [Prénom_F] [Nom], cheffe de la coordination de l'institut.

Tab. 3. Reformulation par la féminisation quand nécessaire

Le stock lexiculturel de la langue française dispose du nom de personne « cheffe », désignant une personne qui a sous sa direction la responsabilité d'un service. Donc, nous préconisons le critère de la féminisation du genre des noms de personne ou de métier lorsque ceux-ci sont socialement impactants.

3.2 La « réactivation » lexicale

La « réactivation » d'unités lexicales, de régularités ou de morphèmes déjà existants dans la langue mais tombés en désuétude (Viennot, 2014) est à solliciter afin de proposer des reformulations plus inclusives, tout en étant respectueuses de la morphologie de la langue. Nous proposons, dans le Tab. 4, un cas de figure puisé dans notre corpus avec une reformulation attentive à l'inclusion.

Forme non inclusive circulant dans notre corpus	Forme envisagée en tant qu'inclusive
Chaque document doit porter la désignation précise de son auteur, sans confusion possible avec l'institution. (Centre hospitalier de Lézignan-Corbières)	Tous les documents doivent porter la désignation de la personne autrice, sans confusion possible avec l'institution.

Tab. 4. Réactivation d'unités lexicales pré-existantes à de fins d'inclusion

Une solution possible porterait à introduire le mot générique « personne », à mettre en place une conversion grammaticale de « auteur » en adjectif et à puiser la forme fléchie au féminin « autrice » ayant disparue de l'usage alors qu'elle existe depuis longtemps.

3.3 Reformulation métonymique

Nous proposons d'adopter le critère défini par nos soins « reformulation métonymique » pour adapter l'énoncé qui suit à une expression plus inclusive :

Communication inclusive et apprentissage profond :
adaptation méthodologique du projet E-MIMIC
à la langue de l'administration française

Ex : Allocation d'activité contrat apprenti (www.service-public.fr/particuliers/vosdroits/F33375)

Afin d'éviter d'avoir recours au masculin générique au niveau du nom de métier (apprenti), nous envisageons la solution proposée dans le Tab. 5 :

Forme non inclusive circulant dans notre corpus	Forme envisagée en tant qu'inclusive
Allocation d'activité contrat apprenti	Allocation d'activité contrat d'apprentissage

Tab.5. Reformulation métonymique

Nous penchons pour la collocation « contrat d'apprentissage », passant de la désignation de l'actant : l'apprenti, à l'action d'apprendre un métier, en particulier une formation professionnelle organisée permettant d'acquérir une qualification.

3.4 Accord de proximité

Pourvu que le recours à l'accord de proximité soit résiduel en tant que solution compensatoire inclusive, elle constitue à notre sens une alternative valable.

Dans le cas rapporté dans le Tab. 6, nous devons faire face à une énumération de sujets.

Forme non inclusive circulant dans notre corpus	Forme envisagée en tant qu'inclusive
Tous les étudiants, étudiantes et élèves ont accès au self du site Sainte-Anne au tarif équivalent CROUS.	Toutes les étudiantes, étudiants et élèves ont accès au self du site Sainte-Anne au tarif équivalent CROUS.

Tab. 6. Reformulation et accord de proximité

Pour fournir une solution alternative à celle qu'on a préconisée dans les sections 3.1.1 et 3.1.2, nous proposons d'opérer une inversion dans l'ordre de présentation des sujets favorisant la forme au féminin en premier et prévoyant ainsi l'accord de proximité.

3.5 Hyperonymie de genre

La propriété de l'hyperonymie de genre est la propriété qui permet à l'unité de représenter tous les genres. Nous estimons donc

qu'en cas de formes fléchies, la stratégie à envisager est celle que nous présentons dans le Tab. 7 :

Forme non inclusive circulant dans notre corpus	Forme envisagée en tant qu'inclusive
Nous devons d'abord approuver le procès-verbal de la séance ... représentons ici un certain nombre d'électeurs montrougiens (Procès-verbal ville de Montrouge 15 déc. 2022)	Nous devons d'abord approuver le procès-verbal de la séance ... représentons ici une certaine portion de l'électorat montrougien.

Tab. 7. Reformulation par l'hyperonymisation

Nous nommons « hyperonymisation de genre » le processus langagier qui consiste à préférer un terme qui ne peut être fléchi qu'à un seul genre (« l'électorat ») à un terme qui peut être fléchi à plusieurs genres mais qui n'est utilisé qu'au masculin en emploi générique. Si besoin, l'hyperonymie de genre attribue cette propriété à ce nouveau mot (ou syntagme) si celui-ci ne la possède pas déjà.

3.6 Épicénisation

Le recours à des mots épicènes qui existent déjà et à la création de nouveaux épicènes pour éviter l'emploi générique du genre masculin relève d'un processus langagier nommé « épicénisation ». Il peut se définir comme processus relevant du français inclusif qui, pour éviter l'emploi générique du genre masculin, consiste à avoir recours à un mot épicène (ex. « élève » plutôt qu'« étudiant »).

Afin d'éviter d'avoir recours au masculin générique, le processus langagier de l'épicénisation puise dans le stock lexiculturel de la langue pour accéder à des solutions plus soutenables, comme le montre l'exemple dans le Tab. 8.

Forme non inclusive circulant dans notre corpus	Forme envisagée en tant qu'inclusive
C'est une autorité administrative indépendante qui ne dépend ni du pouvoir politique, ni des directeurs du projet. (Compte-rendu Préfecture de La Drôme, 30 nov. 2022)	C'est une autorité administrative indépendante qui ne dépend ni du pouvoir politique, ni des responsables du projet.

Tab. 8. Reformulation par l'épicénisation

Il faut tout de même remarquer que les emprunts aussi, quand ils ne sont pas lexicalisés, peuvent devenir des ressources précieuses

en tant que mots épïcènes, comme par exemple *leader*, *manager* (mais la forme lexicalisée « *manageur* », qui n'est plus épïcène, existe aussi), etc.

3.7 Recours à des mots génériques

Comme le précise le dictionnaire de l'Académie française en ligne, « *membre* » désigne, au sens figuré, « chacune des personnes qui composent un groupe, une communauté, une partie d'un ensemble organisé ». De manière similaire, le mot « *individu* » renvoie à tout être vivant dont se compose une société, un groupe, une collectivité.

Du point de vue sémantique, des mots génériques comme « *membre*, *individu*, *personne*, etc. », permettent alors de reformuler de manière inclusive qui a la charge d'une fonction ou la capacité de prendre des décisions dans une organisation, comme le montre l'exemple dans le Tab. 9.

Forme non inclusive circulant dans notre corpus	Forme envisagée en tant qu'inclusive
Les visites autonomes de groupes se font sous la conduite d'un(e) responsable du musée. (Règlement de visite du Palais de la Porte Dorée, 2 juin 2021)	Les visites autonomes de groupes se font sous la conduite d'un membre responsable du musée. Les visites autonomes de groupes se font sous la conduite d'un individu responsable du musée.

Tab. 9. Reformulation à l'aide des mots « *membre* » et « *individu* »

Nous signalons que ces mots sont souvent considérés comme étant des épïcènes bien que leurs fonctionnements syntaxiques et sémantiques restent différents : en effet, les mots épïcènes ne produisent pas de généralisation et demandent donc l'accord genré des déterminants (par exemple, « *un(e) gestionnaire courageux/euse* » *vs* une *personne courageuse* / un *individu courageux* / un *membre courageux*).

En guise de conclusion

Ce tour d'horizon, nous a permis de montrer notre acheminement heuristique vers la découverte et le remaniement de l'écriture

inclusive administrative de France. Reste que le français inclusif s'avère encore lacunaire pour atteindre une totale inclusivité. C'est qu'une langue française véritablement inclusive doit peut-être commander, non des solutions éparses et qui se concurrencent, mais une systématisation de ces nouveaux traitements du genre grammatical. Un travail à double volet — morphosyntaxique et théorique — pourrait peut-être atteindre un degré de français inclusif qui, pour l'instant, est bien loin du compte.

Bibliographie

Haddad Raphaël (éd.) (2017). *Manuel d'écriture inclusive : faites progresser l'égalité femmes · hommes par votre manière d'écrire*. Paris : Mots-clés.

Paveau Marie-Anne (2017). *L'analyse du discours numérique. Dictionnaire des formes et des pratiques*. Paris : Hermann.

Raus Rachele, Tonti Michela (2025). « L'intelligence artificielle pour préserver le français et l'italien : le projet E-MIMIC ». *Langages*, 237, 21-41.

Raus Rachele, Tonti Michela, Cerquitelli Tania, Cagliero Luca, Attanasio Giuseppe, La Quatra Moreno, Greco Salvatore (2022). « L'analyse du discours et l'intelligence artificielle pour réaliser une écriture inclusive : le projet E-MIMIC ». In : *8^e Congrès Mondial de Linguistique Française*, HS Web Conf, 138, 1-15. https://www.shs-conferences.org/articles/shsconf/pdf/2022/08/shsconf_cm1f2022_01007.pdf

Tonti Michela (2024). « L'annotation humaine au croisement de l'intelligence artificielle et de l'écriture inclusive : le dispositif *Inclusively* », *Synergies Italie*, 20, 147-178.

Viennot Eliane (2014). *Non, le masculin ne l'emporte pas sur le féminin !*. Donnemarie-Dontilly : Éditions iXe.

Sitographie

www.cada.fr

www.service-public.fr

www.annuaire-entreprises.data.gouv.fr

www.legifrance.gouv.fr

www.data.gouv.fr

www.collectivites-locales.gouv.fr

<https://www.legifrance.gouv.fr/jorf/id/JORFTEXT000036068906>

Riflessioni a *latere*

Il linguaggio discriminatorio tra etica ed etimologia

Danio Maldussi, danio.maldussi@unibg.it
Università degli Studi di Bergamo

Introduzione

Il percorso di riflessione intorno al linguaggio discriminatorio del presente saggio si articola in due parti: la prima è dedicata all'etica intesa come comportamento dell'uomo nei suoi risvolti morali, tra cui rientra il rispetto del prossimo e pertanto anche una scrittura che sia inclusiva; la seconda è incentrata sull'etimologia, cruciale e ineludibile per risalire alle radici delle parole e individuare il punto di rottura ossia dove queste ultime si snaturano, da "neutre" diventano "orientate" e quindi, nel nostro caso, oggetto di critiche di mascolinizzazione. Due poli apparentemente lontani ma che in realtà convergono: da un lato, infatti, si tratta di considerare gli interventi volti a favorire la scrittura inclusiva come insieme di azioni a valle di un movimento più ampio che vede, a monte, il comportamento umano nel suo complesso. Il lessico che usiamo incide sui comportamenti, o sono piuttosto i comportamenti che determinano il lessico che utilizziamo? Forse nessuna opzione prevale sull'altra, di certo si alimentano a vicenda. In ogni caso, riteniamo che un progetto di grande respiro quale *Empowering Multilingual Inclusive Communication*, debba necessariamente inserirsi a sua volta in uno spazio più vasto, di cui costituisce la punta dell'*iceberg*. Se è indubbia la necessità di intervenire *ex post* sulla lingua, crediamo che sia di primaria importanza agire *ex ante* sui comportamenti umani e relazionali. Dall'altro lato, si tratta di non procedere a una facile epurazione, ignorando così le radici etimologiche, come cercheremo di illustrare con l'analisi della locuzione latina "*bonus pater familias*", espunta dal diritto francese e sostituita, a seconda dei contesti giuridici, dall'aggettivo "*raisonnable*" e dall'avverbio "*raisonnablement*". Il tutto in nome della lotta agli stereotipi di genere e della neutralizzazione del linguaggio, quest'ultimo tacciato nello specifico di patriarcalismo e pertanto di palese discriminazione nei confronti delle donne (Rivais, 2014).

1. Rivalutare la parola “rispetto”

Il linguaggio discriminatorio vale a dire l'uso di «parole, immagini, simboli e atti comunicativi che rinforzano posizioni di subordinazione o diffondono pregiudizi e stereotipi» (Tumminello, 2021) deve necessariamente essere affrontato nel quadro di un disegno più ampio. Lo dimostrano le varie e recenti iniziative, anche a livello comunale e regionale, nonché le proposte di legge regionale, volte a sviluppare l'affettività tra i giovani e di conseguenza delle relazioni personali di qualità, improntate al rispetto reciproco, al controllo dei propri impulsi, prevenendo, perlomeno negli auspici, la violenza di genere. Proliferano le iniziative nei vari paesi. In Francia, ad esempio, il rapporto denominato *Programme d'Éducation à la vie affective, relationnelle et sexuelle* (EVARS), il cui testo doveva essere discusso a dicembre 2024 dal *Conseil supérieur de l'éducation*, poi rinviato, mira sostanzialmente a rinnovare i corsi di educazione affettiva e sessuale a scuola. Non senza tuttavia suscitare l'opposizione dei cattolici che temono un partito preso ideologico. Si tratta pertanto, nell'ambito dell'educazione civica, di ampliare e rendere più incisiva l'educazione al vivere insieme in senso lato.

Taluni, come Haidt (2024), puntano il dito contro i media sociali, accusati di avere depauperato anche gli ultimi scampoli di contatto umano a colpi di *likes* e di avere dato vita a una generazione ansiosa. Possibile, ma non tutti concordano: i media sociali non potrebbero avere semplicemente esacerbato comportamenti già *in nuce* all'interno della nostra società e quindi essersi limitati a rispecchiarli? Esistono casi in cui il mutamento nel sistema di comunicazione ha indotto profondi cambiamenti, come nel caso del concetto di 'democrazia' (Bentivegna, Campus, Valeriani, 2024: 64). In assenza di disintermediazione e a fronte del conseguente *empowerment* dei cittadini, l'arena mediale si fa necessariamente complessa facendo spazio, ad esempio, a un nuovo concetto di democrazia: al concetto di democrazia rappresentativa incarnata dai partiti politici si è sostituito quello di "democrazia dei pubblici" (*ibidem*: 55). Viviamo indubbiamente in una società complessa e in continua evoluzione.

Sono domande che ci interpellano e che richiedono una risposta globale, non soltanto dalla scuola come comunità educante, ma anche e soprattutto dalla famiglia che dovrebbe riappropriarsi del proprio ruolo nella formazione dell'individuo, ritornando a fare, molto

semplicemente, la propria parte. Accogliamo pertanto favorevolmente la notizia che “rispetto” è la parola dell’anno per la Treccani. Come spiegano su RaiNews Valeria Della Valle e Giuseppe Patota, condirettori del Vocabolario Treccani, questa parola

dovrebbe essere posta al centro di ogni **progetto pedagogico**, fin dalla prima infanzia, e poi diffondersi nelle relazioni tra le persone, in famiglia e nel lavoro, nel rapporto con le istituzioni civili e religiose, con la politica e con le **opinioni** altrui. Il termine “rispetto”, continuazione del latino “*respectus*”, va oggi rivalutato e usato in tutte le sue sfumature, proprio perché la mancanza di esso è alla base della violenza esercitata quotidianamente nei confronti delle donne, delle minoranze, delle istituzioni, della natura e del mondo animale. (Barbàra, 2024; il grassetto è nel testo).

Certo, se la sfera genitoriale è uno dei molteplici ambiti ad essere coinvolti, come bene sottolineano Della Valle e Patota (*ibidem*), riteniamo che la sua centralità andrebbe ribadita in funzione proprio del recupero del valore dell’*exemplum* nella sua forma più nobile. Anche perché, come ha più volte sottolineato Gino Cecchettin, padre di Giulia uccisa dall’ex fidanzato l’11 novembre 2023, la violenza di genere non si combatte con le pene ma con la prevenzione¹. E la prevenzione non può prescindere dal “*respectus*” e dall’“*exemplum*”.

2. La locuzione ***bonus pater familias***: tra buon senso e radici etimologiche

Passiamo ora al secondo punto della nostra riflessione ossia alla locuzione latina “*bonus pater familias*”. Quanto segue ci dà la misura del livello di esacerbazione. Nell’agosto del 2014, la locuzione latina “*bonus pater familias*” è stata sostituita nel Codice civile francese dall’aggettivo “*raisonnable*” e dall’avverbio “*raisonnablement*” (Rivais, 2014). “*Raisnable*” e “*raisonnablement*” sono da intendersi non tanto nell’accezione filosofica derivata da “*raison*” che sta a indicare la capacità di ragionamento, propria di chi è dotato di ragione, quanto in quella appartenente al lessico comune di “ragionevole” ossia permeato di buon senso (CNRLT). Il caso della sostituzione in-

¹ <https://www.avvenire.it/attualita/pagine/gino-cecchettin-parla-ai-giovani-sir>

veste il contesto giuridico, quindi specialistico, e questo costituisce un elemento centrale come vedremo qui di seguito.

Già nel 2004, Tosi e Visconti nel loro saggio sull'europeizzazione della lingua italiana ci avevano visto bene paventando la prossima sostituzione, nel contesto del diritto civile francese, della locuzione "*bonus pater familias*" con "*raisonnable*". Ora, ci dicono i due autori, la "*reasonableness*" è un concetto cardine della *common law*. I giuristi anglofoni fanno continuamente riferimento ad espressioni quali "*reasonable step*" oppure "*reasonable person*". Se è vero che le lingue continentali hanno una tradizione filosofica basata sul concetto di ragione «il problema è che gli uomini di legge del diritto civile continentale rifuggono dal concetto di ragionevolezza, totalmente estraneo al diritto romano, basato su diritti scritti e imprescindibili» (Tosi, Visconti, 2004: 159). La vaghezza di "*reasonable*" è evidente nell'atteggiamento del giurista di *common law*, che è persona ragionevole che valuta caso per caso se è stata usata più forza del dovuto. *Common law* e *Civil law* sono due sistemi giuridici che operano in modo radicalmente diverso. «Eppure» scrivono Tosi e Visconti, «il concetto di ragionevolezza si sta diffondendo in lingue che non sono l'inglese», il che dovrebbe dare fastidio ai francesi «così fieri della loro definizione di *négligence* come la '*déviation du standard du bon père de famille*'». «Il termine esiste», concludono gli autori, «ma è il concetto a richiedere di essere esaminato con più cura nella traduzione» (*ibidem*: 159-160). In italiano, ad esempio, la locuzione "*bonus pater familias*" è tuttora presente nei contratti assicurativi ed equivale a "con la diligenza del buon padre di famiglia". Essa sta a indicare «la **diligenza media**, intesa come impegno adeguato di energie e di mezzi per il soddisfacimento dell'interesse del creditore, che è legittimo attendersi da qualunque soggetto di media avvedutezza e accortezza» (Mazitelli, 2022; il grassetto è nel testo).

A questo punto della nostra riflessione, s'impone un breve *excursus* etimologico. La parola "*pater*" deriva da una figura mitologica. Nella dimensione prelatina, che ne ha costituito la matrice semantica, il *pater* è il Dio degli indoeuropei e non il padre (padrone) che, come scrive bene Mercadante, è «inteso nella propria natura biologica [ed è] indicato piuttosto da *pārens* e *gēnītōr*» (Mercadante, 2022: 78). Il *pater*, in origine, non ha una connotazione personale, ovvero non corrisponde a un affetto di tipo personale, fisico, e solo successivamente diventa *pater familias*, *dominus* ossia colui che possiede e prende la decisione.

La lingua è caratterizzata da slittamenti semantici, e noi, proprio in funzione del principio saussuriano, ereditiamo il linguaggio da chi ci ha preceduto. La diatriba, pertanto, che coinvolge la locuzione latina “*bonus pater familias*”, è prettamente etimologica. Il *pater* è una condizione divina che resta stabile fino a tutto il periodo prelatino per poi subire uno slittamento nel diritto romano, dove assume identità legale. L'espressione, come già evidenziato, non possedeva una componente affettiva, che è altresì riscontrabile nella parola “matrimonio”. Come scrive in proposito Mercadante:

Nell'indoeuropeo, di fatto, non esisteva, un termine specifico per designare il matrimonio nel senso in cui lo intendiamo oggi. L'*Enciclopedia delle scienze sociali* Treccani riporta, addirittura, che gl'indoeuropei, nel fare riferimento a questa istituzione, utilizzavano termini diversi, a seconda che essa fosse riferita all'uomo o alla donna: per l'uomo, si usavano verbi derivati dalla radice *wedh-*, ‘condurre una donna nella propria casa’, mentre, per la donna, non erano usati verbi, ma sostantivi, a indicare che ella non compiva l'atto di sposarsi, ma lo subiva. A Roma, invece, lo spotalizio vero e proprio era indicato col termine *nuptiae*, che, derivando da *nūbēre*, ‘velare, coprire’, esprimeva un forte valore figurativo: allo stesso modo in cui una nube copre il cielo, così la donna è coperta dal velo. Il riferimento è al velo giallo o arancione – il *flammēum* – che veniva posto sul capo della sposa e che scendeva a velare il volto (Mercadante, 2025b).

A nostro avviso, considerare “*bonus pater familias*” come una sorta di locuzione maschilista è un'esacerbazione frutto forse di un eccesso di militanza. Prova ne è la matrice originaria ed etimologica di *pater*.

A riprova della necessità di risalire alle origini etimologiche, citiamo l'esempio dello svuotamento semantico subito dalla parola *fēmīna* (femmina):

Nel XV secolo, di fatto, l'annientamento del carattere sessuale della *donna* e la sua relativa demonizzazione erano già fenomeni bell'e compiuti. Il nucleo semantico di *fēmīna* era stato svuotato per sostituzione: il ricorso al termine *donna* quale derivato del tardo latino *dōmna*, per i letterati dell'epoca, era stato strumentale, per così dire; essi, in altri termini, ossessionati dal timore della dannazione e dal peccato della carne, ricostruirono, attraverso un modello linguistico, una figura femminile pura e impalpabile. Di conseguenza, l'uso di *donna* sopravanzò quello di *femmina*, generando, col tempo, quel proposito spregiativo che tutti noi oggi attribuiamo a chi dica *femmina* in luogo di *donna*. Dalla poesia trobadorica in lingua d'oc e dal romanzo cortese-cavalleresco fino al

Dolce stil novo e a Dante, il desiderio sessuale divenne oggetto di costante smaterializzazione, cui gli autori giunsero mediante febbrili e, talora, quasi maniacali meccanismi di sublimazione (Mercadante, 2025a).

Oggi, come sappiamo la parola “femmina” si è talmente caricata di connotazioni negative e offensive nei confronti di una donna (derivato del tardo latino “*dōmna*”), da diventare una parola tabù. In realtà, la radice di “femmina” (**fē-*) sta a indicare

colei che viene succhiata, cioè colei grazie alla quale l'essere umano può farsi *perfetto* ed *eccellente*, essere appagato e divenire, a propria volta, fecondo, produttivo. Dalla sua *mammella* il piccolo trae il *nutrimento*: ciò va detto, nonostante le linee linguistiche patriarcali che, nei secoli, hanno occultato il ruolo precipuo e dominante della donna (*idem*).

“Femmina”, «che, nel registro quotidiano, ormai, a causa d'una moda bislacca e sconclusionata, è percepito come dispregiativo rispetto a ‘donna’, [...] per etimologia, invece, esprime sostanziale nobiltà di genere» (*idem*). Al contrario, la parola “donna” che possiede la radice proto-indo europea **dem-*, indica la casa (Gavetti, 2021). “Donna”, a propria volta, è «forma sincopata di *dominā*, ovverosia *signora, padrona, moglie, sposa*». Per concludere questa breve disamina, «il legame tra la *donna* e la *casa* sembra scontato, oltre che immediato» (Mercadante, 2025a).

Per concludere, risalire alle radici etimologiche potrebbe aiutare a non gettarsi lancia in resta contro le parole, a non cadere negli eccessi di quella che D'Amico definisce «la dittatura del linguaggio anti discriminatorio» (D'Amico, 2023: 199). Chi intende «vedere pure nella filologia una sorta di maschilismo» (Mercadante, 2022: 79), dovrà, in taluni casi, farsene una ragione.

Conclusioni

In questo nostro breve saggio abbiamo affrontato la questione del linguaggio discriminatorio da due prospettive a nostro avviso complementari. In primo luogo, abbiamo volutamente inserito la questione del linguaggio in un disegno più ampio, consapevoli che l'educazione all'affettività e il rispetto del prossimo costituiscano i

preamboli ineludibili per dare vita a un linguaggio non discriminatorio. In secondo luogo, abbiamo voluto ribadire come la lotta alla discriminazione debba, a nostro avviso, fare i conti con le radici etimologiche delle parole che spesso, nel tempo, subiscono slittamenti semantici che poco hanno a che vedere con la matrice originaria ed etimologica. Auspichiamo pertanto uno sguardo più sereno sulle nostre radici che ci consenta di non fare *tabula rasa* di parole che poco hanno a che vedere con la discriminazione nei confronti delle donne. Potrebbe in proposito essere utile rileggere e magari mandare a memoria la poesia *Le parole* di Gianni Rodari musicate da Sergio Endrigo (1974):

Abbiamo parole per vendere,
Parole per comprare,
Parole per fare parole.

Andiamo a cercare insieme
Le parole per pensare.
Andiamo a cercare insieme
Le parole per pensare.

Abbiamo parole per fingere,
Parole per ferire,
Parole per fare il solletico.

Andiamo a cercare insieme,
Le parole per amare.
Andiamo a cercare insieme
Le parole per amare.

Bibliografia

La data di ultima consultazione dei siti è il 31 gennaio 2025.

Barbàra Ugo (2024). "Per Treccani "Rispetto" è la parola dell'anno". *AGI* <https://www.agi.it/cronaca/news/2024-12-16/rispetto-parola-dell-anno-29181085/#:~:text=Il%20Dizionario%20dell'italiano%20Treccani,esprimere%20con%20azioni%20o%20parole%22>

Bentivegna Sara, Campus Donatella, Valeriani Augusto (2024). *La comunicazione politica contemporanea*. Bologna: Il Mulino.

CNRTL (*Centre national de ressources textuelles et lexicales*) <https://www.cnrtl.fr/definition/raisonnable>

D'Amico Marilisa (2023). "Linguaggio discriminatorio e garanzie costituzionali". *Rivista AIC (Associazione italiana dei Costituzionalisti)*, 1, 198-254.

Endrigo Sergio (1974). *Ci vuole un fiore*. Testi di G. Rodari; musiche di L. Bacalov, S. Endrigo; orchestra diretta da Luis Bacalov, Dischi Ricordi.

Gavetti Elisabetta (2021). "Donna, s.f.". *Massime dal passato*. <https://massimedalpassato.it/donna-s-f/>

Haidt Jonathan (2024). *La generazione ansiosa. Come i social hanno rovinato i nostri figli*. Milano: Rizzoli.

Mazzitelli MariaFrancesca (2022). "Il buon padre di famiglia nel diritto privato". *Blog Simone* <https://edizioni.simone.it/2022/09/22/buon-padre-di-famiglia-nel-diritto-privato/#:~:text=La%20diligenza%20del%20buon%20padre%20di%20famiglia%20%C3%A8%20la%20diligenza>

Mercadante Francesco (2022). *Le parole dell'economia. Viaggio etimologico nel lessico economico*. Milano: Il Sole 24 ORE.

Mercadante Francesco (2025a). "Donna, non più femmina: semantica della desessualizzazione". *Lingua italiana*. Roma: Treccani https://www.treccani.it/magazine/lingua_italiana/speciali/Trame/T_Mercadante4.html

Mercadante Francesco (2025b). "Sposo e sposa, coloro che si prendono a garanti". *Lingua italiana*. Roma: Treccani https://www.treccani.it/magazine/lingua_italiana/speciali/Trame/S_Mercadante5.html

Rivais Rafaële (2014). "Comment le "bon père de famille" vient de disparaître du code civil". *Le Monde*. https://www.lemonde.fr/vie-quotidienne/article/2014/08/21/comment-le-bon-pere-de-famille-vient-de-disparaitre-du-code-civil_6003843_5057666.html

Tosi Arturo, Visconti Jacqueline (2004). “L’‘europeizzazione’ della lingua italiana”. *Lingua italiana d’oggi*, I, 151-173.

Tumminello Fabio (2021). “Ai tempi dell’odio: linguaggio, hate speech e diritti umani”. *Ius in Itinere*. <https://www.iusinitinere.it/ai-tempi-dellodio-linguaggio-hate-speech-e-diritti-umani-40538>

Chi ha dato i dati all'intelligenza artificiale?

Una riflessione al tempo dei Large Language Models

Francesca Dragotto, francesca.dragotto@uniroma2.it
Università di Roma Tor Vergata

Introduzione

L'addestramento del modello di intelligenza artificiale, ovvero il processo che porta un codice ad attraversare le fasi di acquisizione dell'intelligenza, riproduce più di quanto si fosse propensi a ritenere i meccanismi fondamentali della costruzione della conoscenza umana, in particolare di quella dei primi semestri di vita che culmina nella capacità di ogni singolo “nuovo” cucciolo di *Sapiens* di diventare competente della/nella propria L1. Proprio come avviene per la predisposizione biologica all'acquisizione di una lingua da parte dell'individuo, il solo “codice”, benché condizione necessaria per lo sviluppo di capacità d'interpretazione e apprendimento da parte di una IA, necessita di “esposizione” ai dati perché il processo possa diventare operativo e performante. Questo parallelismo suggerisce un approfondimento di domanda sull'impatto che quantità e qualità dei dati esercita sulle “forme dell'imparare” proprie dell'intelligenza, di quella umana e conseguentemente di quella artificiale. Per sostanziare la riflessione, si assumerà un punto di vista vicino alle scienze del linguaggio.

Quella sul ruolo dei dati costituisce infatti probabilmente la domanda per antonomasia da porsi, anzi da ri-porsi, quando ci confronta con il tema dell'apprendimento da parte della macchina. Che sussistano delle differenze è cosa ovvia, persino banale. La stessa Chat GPT 4.0, interrogata (in più sessioni, l'ultima delle quali il 14 gennaio 2025) su possibili somiglianze tra come il sapere interpretativo si costruisce in *Sapiens* e nell'IA — «l'IA acquisisce il linguaggio come un bebè?» — risponde in modo *tranchant* che: «No, l'intelligenza artificiale non acquisisce il linguaggio come un bebè. Il modo in cui un'IA impara a usare il linguaggio è completamente diverso da come avviene l'acquisizione del linguaggio umano».

Il *chatbot* passa quindi ad elencare le differenze principali, in due blocchi.

Nel caso del primo, che fa precedere dall'indicazione «Acquisizione del linguaggio in un bebè», cita in successione:

1. Interazione sociale
2. Esposizione graduale
3. Cognizione e creatività
4. Neuroplasticità.

Nel caso del secondo, a cui antepone «'Acquisizione' del linguaggio in un'IA», cita invece:

1. Apprendimento dai dati:
2. Nessuna cognizione o contesto reale
3. Velocità e quantità
4. Mancanza di adattamento biologico.

Tra i due soggetti apprendenti sembrerebbe perciò sussistere una distanza difficilmente colmabile, sulla quale parrebbe invece opportuno un supplemento di riflessione, stando, per esempio, a nuovi elementi provenienti dalla sperimentazione di modelli di apprendimento automatico di un'IA alla quale sono servite solo poche decine di ore di filmato registrato a intermittenza — poi trascritto — da una fotocamera posta sulla testa di un bambino per iniziare a capire il significato di alcuni sostantivi¹.

I risultati di un simile esperimento, fondato su una quantità di dati di esposizione irrisoria se paragonata a quella dei LLM, rendono necessario, come si diceva, almeno un supplemento di riflessione sui dati e, a cascata, sulla complessità del meccanismo di acquisizione umana, che potrebbe rivelarsi meno astratto di quanto finora immaginato. A beneficio di questo supplemento, e sempre

¹ Vong et al. hanno indagato la questione in un modo longitudinale del tutto innovativo, utilizzando registrazioni video di esperienze svolte in prima persona da un singolo bambino, Vic, in un intervallo di tempo compreso tra i suoi 6 e 25 mesi di vita. Le immagini ottenute dalla videocamera montata sulla sua testa hanno ripreso l'interazione del bambino in contesti autentici (naturalistic settings). Applicando tecniche di apprendimento automatizzato, Vong et al. hanno introdotto il modello Child's View for Contrastive Learning (CVCL) e, sincronizzando i fotogrammi ricavati dal video con le parole corrispondenti, che venivano pronunciate simultaneamente, hanno incorporato immagini e parole in spazi rappresentativi condivisi, riproducendo in tal modo l'esperienza testuale multimodale a cui è esposto il cervello dell'infante in via di acquisizione.

all'interno di un framework multi- inter- e transdisciplinare, del quale difficilmente si potrebbe fare a meno vista la complessità di ciò che chiamiamo “intelligenza”, potrebbe essere compulsato in modo più approfondito il sapere relativo dell'acquisizione della lingua (L1 *in primis*) prodotto in epoca recente in ambito di scienze del linguaggio e di altre discipline che si occupano di descrivere e spiegare questo processo.

Riconsiderare la questione dei dati all'interno di un perimetro coerente con le osservazioni del processo acquisizionale potrebbe inoltre contribuire a un altro aspetto oggetto di dibattito nel corso di questo convegno, quello della decolonizzazione della conoscenza, intendendo con ciò il contrasto all'incistamento e pervasività nel modello di conoscenza, sia esso umano o no, dei costrutti socialmente e culturalmente marcati presenti nei (modelli mentali soggiacenti ai) testi di esposizione.

1. Le tappe dell'acquisizione

Riconsiderare analogie e differenze del processo di acquisizione della conoscenza, e quindi dell'affinamento del dispositivo interpretativo che con essa e da essa scaturirà, è senz'altro il punto da cui avviare la riflessione.

La differenza principale tra il processo umano e quello che tende a replicarlo² consiste nelle caratteristiche dei *setting* testuali e contestuali che caratterizzano l'esposizione³.

Entrambi, quello a cui è esposto il cervello e quello con il quale si addestra la macchina, sono infatti caratterizzati da ricchezza multimodale e dunque da sincretismo di linguaggi, ma quelli umani, almeno finora, comprendono una ricchezza espressiva maggiore (verbale, paraverbale e non-verbale insieme) di quella propria dei testi a cui viene “esposta” la macchina che impara a partire da *dataset* preaddestrati. Questa gestisce per lo più testi prodotti in una modalità

² L'impiego di artificiale è fuorviante perché, attivando il recupero di una idea infondata di non umanità, induce a immaginare che l'IA non sia soggetta a quelle scorciatoie in cui si esprime l'economia cognitiva, tendenza che nella specie umana è presente al massimo grado.

³ Si intende con *setting* l'ambiente popolato da testi — discorsi e produzioni multimodali di ogni tipo —, ecosistema ideale per l'acquisizione della lingua.

che è scritta per quanto riguarda il formato ma che è spesso più prossima all'oralità per varietà linguistica di riferimento, ma che dell'oralità schiacciano, opacizzandole, le caratteristiche prosodiche, che nell'esperienza umana guidano invece l'acquisizione della grammatica della lingua in tandem con l'*input* visivo.

Nel caso dell'acquisizione del linguaggio verbale da parte di un infante è prassi ripartire il processo in tappe, utili soprattutto in ambito clinico per circoscrivere e distinguere casi di *late-blooming* (tempi rallentati in un contesto comunque tipico) da casi di *late-talking* (indizio di sviluppo atipico e di potenziali *language disorders*).

Tali fasi sono di prassi identificate in: vocalizzazione, lallazione, protoparole e prime parole in senso stretto, fase olofrastica e successivamente frasi espanse, via via più ricche per lessico e morfosintassi. A fare da stimolo per l'attraversamento delle diverse basi, e da innesco del sistema acquisizionale tutto, c'è proprio il contesto situazionale con i suoi attori e i linguaggi di cui questi si servono per la produzione di testi multimodali. Oltre che dalle interazioni dirette, in particolare negli ultimi anni questi *setting* sono stati popolati da interazioni mediate da schermi, che hanno posto l'infante in una situazione in fondo assai simile a quella "esperita" dalla macchina, vista la natura per lo più monodirezionale e lo schiacciamento (linearizzazione) dei messaggi/testi che attraversano lo schermo.

Per dirlo con la sintesi conclusiva di ChatGPT, nel suo complesso «Un bebè acquisisce il linguaggio come parte della sua esperienza di vita: è un processo biologico, sociale e cognitivo, radicato nella necessità di comunicare e comprendere il mondo».

Nel caso della macchina, secondo una concezione data finora per scontata, questi *set* di esposizione sarebbero popolati da soli messaggi/testi scritti: «L'IA 'impara' il linguaggio come un processo di estrazione e generalizzazione di modelli statistici dai dati. Non comprende il linguaggio o il mondo nel senso umano, ma lo utilizza per simulare risposte coerenti» (ChatGPT). La ricorrenza di successioni simili di parole-*tokens* consentirebbe perciò all'IA di maturare delle attese che, una volta confermate, vanno a diventare lo schema interpretativo di riferimento (*pattern*). Un *pattern* che *in itinere* va a perfezionarsi grazie alle successive interpretazioni.

Questo meccanismo, a ben vedere, non si differenzia molto dai meccanismi alla base della significazione che consente all'infante di procedere nella progressiva attribuzione di significato a stimoli sen-

soriali esterni ricorrenti. E dunque nell'opera di interiorizzazione del mondo per effetto dell'esposizione a testi situazionati.

Benché non si tratti di un adattamento biologico, il funzionamento dell'IA sembrerebbe infatti distaccarsi, semplificandosi, dalle attese alimentate dalle *performance* realizzate nei training basati sui *Large Language Models* (LLM) allorquando il training stesso è stato modificato per avvicinamento proprio all'esperienza umana.

Nel caso dell'episodio citato in apertura, quello dell'addestramento a partire dal materiale filmico incamerato dalla videocamera posta in testa al bambino — filmato multimodale, perciò paragonabile per tipo di stimolo a quello esperito dal bambino stesso — i risultati della ricerca sembrano mostrare, benché ancora preliminarmente, che si può fare a meno di grandi repertori di dati quando si disponga di un materiale ricco di spunto di corrispondenza e, verrebbe da dire, autentico.

Benché lo scopo della ricerca fosse diverso, i risultati inaspettati possono infatti essere indicativi di una strada differente, caratterizzata da maggiore complessità e ripetizione nel tempo dello stimolo situazionato piuttosto che di un'estensione e copiosità di dati però non situazionati. Una strada che avvicina ancora di più le modalità acquisizionali di *Sapiens* e della macchina che tenta di riprodurne il funzionamento. Poiché sia nell'archetipo che nell'apografo la scintilla della costruzione di conoscenza sembra risiedere nella riproposizione nel tempo delle stesse associazioni stimolo (pluri)sensoriale-significato (estratto dal contesto), ovvero nell'esposizione al segno come parte di un testo più ampio e più complesso (impossibile non richiamare l'intreccio alla base proprio del significato del latino *textus*), inevitabile sarà anche la costruzione per assimilazione di pregiudizi e stereotipi.

2. Le conseguenze dell'acquisizione

Rinforzati da *bias* soprattutto di conferma, questi meccanismi di conoscenza facilitata, che nell'individuo possono subire occasionali e/o casuali inceppamenti dai quali assumere prospettive diverse (un imprevisto, un incidente occorso a sé o a una persona cara, una patologia...), nella macchina rischiano di perpetuarsi, trascinandosi dietro/dentro rappresentazioni delle società e di chi le popola anco-

ra più asfittiche di quanto lo sia già la cosiddetta realtà, con le sue asimmetrie e diseguaglianze.

L'universo rappresentazionale dell'IA autogeneratosi per assimilazione dei testi di addestramento, in assenza di cortocircuiti provocati, non può così che tendere a sovrarappresentare ciò che è già sovrarappresentato, a sottorappresentare o non rappresentare ciò che c'è già poco o che non c'è, a rappresentare ciò che c'è amplificando il modo con cui è rappresentato dalla massa di utenti.

L'orizzonte di conoscenza già infeltrito nell'essere umano, mediamente poco incline al confronto con ciò che si differenzia dai tratti che caratterizzano la sua "comfort zone", nella macchina equivale a un infeltrirsi proporzionato all'efficienza delle sue performance.

3. Un ragionamento conclusivo

Quello per i dati di addestramento delle IA è un entusiasmo che senz'altro dovrebbe essere calmierato, senza demonizzazione ma con criticità, alla luce delle implicazioni con cui hanno profilato e continuano a profilare la conoscenza a cui si fa riferimento per dare riscontro a un *prompt*. E ciò specialmente laddove, in nome della quantità e della reperibilità a basso costo, si è attinto da social e altre fonti di analoga fattura.

Le criticità o almeno problematicità del training basato su LLM sono peraltro multiple e riguardano in diversi modi le lingue e i testi, oltre che l'ecosistema per ragioni dovute al consumo di energia e alla produzione di CO₂ richieste dalle specificità di questo tipo di addestramento.

Il ricorso massiccio alla produzione testuale dei *social media*, per esempio, schiaccia l'architettura della lingua e generalizza una varietà ibrida, che di scritto spesso ha la modalità trasmissiva ma non le caratteristiche morfosintattiche, ortografiche e lessicali di riferimento.

Per quanto riguarda il contenuto, inoltre, il modo in cui i *social media* stimolano esibizionismo, narcisismo e inclinazione allo sfogatoio determina, complice l'anonimato o la possibilità di essere chi si vuole, un allargamento difficilmente contenibile degli spazi comunicativi e della semantica propria di un tipo di interazione sociale che, prima dell'*onlife* (Floridi 2017), coesisteva di necessità e per lo

più in regime di minoranza con diverse altre varietà linguistiche, generi testuali e stili, anche nella stessa persona.

Se si aggiunge poi, stavolta in relazione a fonti di apprendimento magari più bilanciate per forma e sostanza, il problema del ricorso alla traduzione automatizzata per l'addestramento delle IA nel caso di lingue con bacini di utenti ridotti, la complessità e problematicità del tema si fa ancora maggiore.

Con tutta l'importanza che merita, l'entusiasmo per la performatività nello svolgimento di un *task* da parte di un'IA dovrebbe fare maggiormente i conti con le conseguenze della costruzione di una conoscenza di riferimento che non include pezzi spesso consistenti di mondo e porzioni di biodiversità sociale e che presta il fianco ad una colonizzazione beata.

Ecco perciò che la politica dei dati, ebbene sì, la politica dei dati è il problema all'epoca dell'IA.

Bibliografia

Cavagnoli Stefania, Dragotto Francesca (2021). *Sessismo*. Milano: Mondadori.

Dragotto Francesca, Cavagnoli Stefania (2023). *Alla scoperta delle scienze del linguaggio*. Milano: Mondadori.

Floridi Luciano (2017). *La quarta rivoluzione. Come l'infosfera sta trasformando il mondo*. Milano: Raffaello Cortina Editore.

Vong Wai Keen, Wang Wentao, Orhan A. Emin, Lake Brenden M. "Grounded language acquisition through the eyes and ears of a single child", *Science*, vol. 383, issue 6882, 504-511.

“The limits of LLMs are the limits of the world”: considerations on the role of linguistics in enhancing AI-generated language quality

Valeria Zotti, valeria.zotti@unibo.it
Università di Bologna

Introduction

In this short article I will present some observations that arose from the activities I carried out in 2024 as part of the PRIN 2022 E-MIMIC project.

In the first part, I will briefly describe the increasing focus on inclusive communication and the risks associated with the use of artificial intelligence systems, such as the loss of language diversity and the spread of biases, which I will illustrate with a few examples from neural machine translation.

In the second part, I will discuss how the flattening of language and the loss of multilingualism caused by AI systems may coincide with the loss of the ability to conceptualize the world differently from dominant models. Observation of language acquisition and development in children reveals the social conditioning that permeates a language and the role of social experience in the comprehension process, a factor that distinguishes human competence from artificial language production.

Linguistics proves to be of central importance in understanding the complexity of human language, as well as the need for strong human involvement in data-driven methods, as is being done in the E-MIMIC project whose innovative deep-learning methodology aims to develop artificial intelligence systems that can serve to promote equality in communication.

1. Deep Learning & Gender Bias in Language and in AI

Awareness of the risks that the development of Artificial Intelligence, particularly Generative AI (GenAI), may cause for the preservation of linguistic and cultural diversity and the protection of gender equality and minority rights has grown significantly in the scientific community over the past few years (Raus *et al.*, 2022; Raus, 2024; Cennamo *et al.*, 2023; Mattioda, 2024).

The rapid and unprecedented progress of artificial intelligence in all areas of human endeavor, and particularly a branch known as Deep Learning, which enables neural network-based algorithms to learn “autonomously” on large amounts of data (Le Cun, 2019), has raised serious questions, both in academia and in industry and institutions, about the provenance and quality of data used to train these systems.

As noted by New York Times tech journalist Cade Metz, in the 2010s, “As Google and other tech giants adopted the technology, no one quite realized it was learning the biases of the researchers who built it. These researchers were mostly white men, and they didn’t see the breadth of the problem until a new wave of researchers — both women and people of color — put a finger on it” (Metz, 2021).

The lack of diversity and inclusion that characterizes the design of AI systems is thus a major concern today. In the field of text generation, AI systems reproduce what humans say and write. If a system is fed on human biases (conscious or unconscious), it will produce results full of these same biases. In the current Large Language Models on which AI systems are trained, the data are heterogeneous and not always traceable. These data risk reinforcing discrimination and bias by giving them the appearance of objectivity.

To illustrate how these systems risk amplifying stereotypes, I will give a first example related to the massive spread of bias in neural machine translation, which, as Yvon (2023: A19) explains, is mainly based on machine learning methods exploiting large corpora of parallel texts aligned at sentence level.

In a second example I will briefly show that another disadvantage of generative AI is that, because of the dominance of training data in English, the preservation of multilingualism is endangered.

1.1 Bias in Neural Machine Translation systems

The following example shows “gender bias” present in the data generated by AI systems taken from a small test performed in 2021 on the well-known neural machine translation software, DeepL.

While in Italian it is possible to omit the subject personal pronoun preceding the verb, resulting in neutral sentences that do not indicate whether the subject is masculine or feminine, in English and French it must be specified.

The output of this NMT system in both English and French (Tab. 1) suggests that numerous stereotypes were present in the data used to train it, made explicit by the use of the female personal pronoun for roles traditionally delegated to the female sex, such as childcare. Other experiments have shown that if the subject personal pronoun is not specified, DeepL generally has a preference for the masculine.

<i>Traduci da Italiano</i>	<i>Traduci in Inglese (US)</i>
Pratica sport. Si occupa dei bambini. Lavora. Guida l'auto.	He plays sport. She takes care of children. He works. He drives a car.
<i>Traduci da Italiano</i>	<i>Traduci in Francese</i>
Pratica sport. Si occupa dei bambini. Lavora. Guida l'auto.	Il fait du sport. Elle s'occupe des enfants. Il travaille. Il conduit une voiture.

Tab. 1. Test DeepL 2021

Three years later, when the same test was repeated in 2024 the same stereotypes were generated again. Activities such as playing sports, working and driving a car are still presented by the British and American English translations as masculine activities, while, in the sentence related to childcare, the software hesitates and omits the subject pronoun thus returning an incorrect sentence (Tab. 2). In the French translation, on the other hand, childcare is attributed to a masculine subject (Fr. pronoun “il” = En. “he”), suggesting that the French data may be less “biased” than the English data.

<i>Traduci da Italiano</i>	<i>Traduci in Inglese britannico</i>
Pratica sport. Si occupa dei bambini. Lavora. Guida l'auto.	He plays sport. Takes care of children. Works. He drives a car.
<i>Traduci da Italiano</i>	<i>Traduci in Inglese americano</i>
Pratica sport. Si occupa dei bambini. Lavora. Guida l'auto.	Plays sport. Takes care of children. He works. He drives a car.
<i>Traduci da Italiano</i>	<i>Traduci in Francese</i>
Pratica sport. Si occupa dei bambini. Lavora. Guida l'auto.	Il pratique le sport. Il s'occupe des enfants. Il travaille. Il conduit une voiture.

Tab. 2. Test DeepL 2024

This simple test raises further questions about the preponderance of training data in English as a potential source of bias.

1.2 English language dominance in Human and Machine Translation

Given that these technologies rely above all on data in English, a language which in many cases is not gender-specific, their ability to generate inclusive texts in other languages (e.g. in the Romance languages involved in the E-MIMIC project), is rather limited.

English is dominant not only for economic reasons, with leading research on AI taking place largely in English-speaking countries, but also for linguistic reasons. Morphologically complex languages (such as Italian, French, and Spanish) are much more difficult to handle than English, while English has the advantage of being morphologically simple (e.g. it has few verb inflections, no gender markings), so creating IT language models for other languages requires more computer engineering work and is more expensive (Vetere, 2023).

An example of the impoverishment of lesser-represented languages due to this sort of “colonization” of English is provided by a study by Henkel (2024) focusing on a comparison of human translation and machine translation, which shows that machine translations have reduced lexical variety. The example reported in this study is about the verb “think” in English which, in the case of human translation, is translated in French in 29% of cases by “*penser*” and in 28% by “*croire*” and for the remaining 40% by other lexemes or expressions, while DeepL translates in 74% of cases with “*penser*” and a few other variants. As noted by Henkel (2024), machine translation applies a limited number of “recipes”: it tends to favour one single translation strategy with a small number of alternative solutions.

The real risk of NMT-induced language impoverishment is thus astoundingly obvious in the statistical data, the dangers of which cannot be overstated. Indeed, linguistic diversity is an expression of diversity of thought, but also of human diversity with respect to gender and ethnic or age differences, disabilities, etc. As follows from Ludwig Wittgenstein's famous aphorism “the limits of language are the limits of my world”, if human language becomes impoverished, under the influence of machine-generated language, the risk in a not-too-dystopian future of an impoverishment of human thought is real, a recurring theme in science fic-

tion¹ literature and film. In other words, with this flattening of language generated by AI systems, the loss of multilingualism may coincide with the loss of the ability to conceptualize the world differently from dominant models. This ability is innate in human beings but the social conditioning we undergo from the moment we acquire a language affects our ability to conceptualize reality differently.

2. Social conditioning in Human and AI generated language

An example taken from my personal experience of observing how children acquire language reveals the extent to which language is a social product and is imbued with unconsciously transmitted stereotypes.

When my daughter was attending kindergarten at the age of 4 and could not yet read or write, I informed her one day that she would be attending a party the next day with ‘the children’ (in Italian, “*i bambini*”) from her class. She stared hard at me and replied, “Mom, why do you say ‘*I bambinI*’ (masculine in Italian) and not ‘*lE bambinE*’ (feminine in Italian)? In my class there are more girls than boys!”.

I explained to her in simple words that in Italian we use the generic masculine to refer to both sexes. After protesting indignantly that “it’s not fair,” she suggested “*bambin**” as a new form that we would use in our family to indicate that we were all equal. Having begun at that time to take an interest in inclusive language, I recognized in this attempt to eliminate sexism in the Italian language the motive for the well-known *Recommendations* (Sabatini, 1987) and the many guidelines that have been published in the last 50 years to promote gender-neutral and non-discriminatory use of language.

This brief exchange with my daughter resonated in me as it made me more aware of how much social conditioning we undergo through language acquisition. Children are born pure with an innate sensitivity that is gradually lost as they learn a language.

¹ As observed during the *ALMA AI Film Festival. Pre-Visions of Artificial Intelligence*, sponsored by the ALMA AI research center and organized by Milano Michela and Zotti Valeria in spring 2024 at the *Cineteca* di Bologna.

My daughter naturally perceived an injustice conveyed by language, namely the occultation of women, which is the reflection of centuries of women's inferior social status. As observed by the father of linguistics, Ferdinand de Saussure, “language is social”, it is a social product of the language faculty and a collective reality, a set of conventions that the community adopts (Saussure, 1916).

Human beings, as they learn a language also integrate stereotypes conveyed by language that are socially induced, the result of language habits and modes of interaction repeated over time within society. Over time, a child who becomes an adult loses the innate ability to recognize these stereotypes and unconsciously integrates social injustice.

My daughter's example shows another key factor, the fundamental role of “experience.” She was aware that three-quarters of her class were females, and this experience, which was also social in nature, enabled her to grasp on a cognitive level the contradiction that exists between reality and the fictitious representation of reality conveyed by language.

Unlike children, AI systems have access to unimaginably large quantities of data but do not experience it socially in the form of real interactions with other living beings within a community (probably one of the reasons why they have no sense of humour yet, as demonstrated by Veale, 2021). Although there is research on “Embodied AI”, which is generative IA inserted in a robotic body which should improve its performance, machines are a long way from social and pragmatic language competency (e.g. knowing what to say in the correct situation). This leads us to the concept of discursive diversity of which IT developers must become aware. Linguistics can be very helpful to raise this awareness.

3. Discursive diversity and the role of linguistics

Linguistics is the scientific study of human language in all its manifestations and structures. Linguists are fully aware of the essential notion of “discourse type”. Language is not monolithic and is characterized by the co-existence of many “discourse types” and “registers,” with different, and sometimes even contradictory or incompatible characteristics. The problem with LLMs is that they contain

raw data taken largely from the Internet (Naik *et al.*, 2024) in which some discourse types are under-represented, others are over-represented, and all are amalgamated as if they were the same type. These data cannot be considered as representative of human language.

The question of “representativeness” has been debated for many years in the field of corpus linguistics, well before the advent of AI. Chomsky (1957: 15) rejected corpus linguistics because he considered that corpus data would always be insufficient to describe language competency or the language system as a whole. Even a text corpus of several billion words is only a small sample compared to the amount of language that hundreds of millions of speakers produce in one day, on and off the internet. We can only comprehend language in its diversity and complexity and obtain higher-quality results from AI, in text generation or machine translation, by distinguishing different discourse types one at a time, using specialized datasets.

The experience of the E-MIMIC Project is innovative in this respect insofar as the material used to train the neural networks is authentic and the analysis and preparation of the data was carried out by linguists following linguistic-discursive criteria designed to ensure inclusiveness (Raus *et al.*, 2022).

The role of linguistics is thus fundamental in IT research in order to understand the fundamental characteristics of human language and the reasons for which we recognize AI-generated language as “artificial”. These are hotly debated issues among those involved in natural language generation and machine translation.

A fascinating book entitled *Intelligence artificielle, intelligence humaine: la double énigme* (The Double Enigma: Artificial Intelligence and Human Intelligence), by French philosopher and mathematician Daniel Andler (2023), offers a compelling enquiry into what “human” or “natural” intelligence is and the various difficulties that AI has encountered. Indeed, since its inception in the 1960s, AI has been claimed to have the potential to equal human intelligence, yet this initial ambition remains an illusion.

As far as language is concerned, generative AI has made astonishing progress in emulating human beings, but machines are still a long way from ‘autonomy’ with respect to humans because they lack true understanding of the complexity of human language, and they lack real-world human experience. Thus, linguistics is essential to

take into account the role that humans play as the era of AI unfolds. If AI lowers the quality of the language humans speak, if humans allow language to be impoverished and if people become accustomed to communicating and interacting in machine-like ways, the risk that we might lose the ability to distinguish between human and artificial intelligence and forget our humanity is real.

Conclusion

Political actions that have been conducted over the past few years have been aimed at re-educating speakers of a language so that they adopt ostensibly “new” but actually “innate” communicative practices that give recognition to all individuals. Some initiatives, such as that of the president of the University of Trent who, in March 2024, published University Regulations written using the ‘overextended feminine’ for offices and gender references (e.g., “*decana*” “dean-fem.” “*rettrice*” “rector-fem.” “*segretaria*” “secretary-fem.” even for men), is, as stated by the president himself, “A symbolic act to demonstrate equality starting with the language of our documents.” Pr. Flavio de Florian said that he felt excluded from this all-female document and understood the necessity of promoting neutral and non-sexist language.

It is paradoxical that in the same era in which we are becoming aware of the pervasiveness of sexist stereotypes in language, generative AI systems “amplify” these biases. Thus, it is essential to develop a linguistic data policy, to make computer scientists aware of this need, of the fact that methods used to collect data are crucial, that “*le manque de données textuelles de qualité entraîne une sous-représentation de certaines langues-cultures*”² (Mattioda, 2024: 5), and that relying on linguistic criteria is efficient. The E-MIMIC project and other commendable initiatives, such as the national AI research network FAIR, have shown how linguistics can help improve the quality of language output from AI and avoid the kind of linguistic impoverishment and biases that have so far plagued AI-generated language.

² “the lack of quality textual data leads to an under-representation of certain language-cultures” (Mattioda 2025: 5).

References

Andler Daniel (2023). *Intelligence artificielle, intelligence humaine: la double énigme*. Paris: Gallimard.

Cennamo Ilaria, Cinato Lucia, Mattioda Maria Margherita, Molino Alessandra (2023). “Promoting multilingualism and inclusiveness in educational settings in the age of AI.” In: Tasa Fuster Vicenta, Monzó-Nebot Esther and Rafael Castelló-Cogollos (eds). *Repurposing language rights. Guiding the uses of artificial intelligence*. València: Tirant lo Blanch, 277-315.

Chomsky Noam (1957). *Syntactic Structures*. La Haye: Mouton.

Henkel Daniel (2024). “Verbs of cognition in translation between English and French”. In: Dister Anne and Longrée Dominique (eds). *JADT 2024 Mots comptés, textes déchiffrés*. Louvain-La-Neuve: Presses Universitaires de Louvain, 399-408.

Le Cun Yann (2019). “Deep Learning Hardware: Past, Present, and Future”. In: *2019 IEEE International Solid-State Circuits Conference - (ISSCC)*, San Francisco, CA, USA, 12-19.

Mattioda Maria Margherita (2024). “Intelligence artificielle et diversité linguistique. Quelle gestion équitable pour la garantie des droits linguistiques?”. *Language Problems and Language Planning*, December 2024, 1-22.

Metz Cade (2021). *Genius Makers: The Mavericks Who Brought AI to Google, Facebook, and the World*. New York: Dutton.

Naik Dishita, Naik Ishita, Naik Nitin (2024). “Large Data Begets Large Data: Studying Large Language Models (LLMs) and Its History, Types, Working, Benefits and Limitations”. In: *Contributions Presented at The International Conference on Computing, Communication, Cybersecurity and AI, July 3–4, 2024, London, C3AI 2024*. Cham: Springer, 239-314. https://doi.org/10.1007/978-3-031-74443-3_18

Raus Rachele (2024). “Deep learning e traduzione intralinguistica: riformulare i testi della pubblica amministrazione in modo inclusivo”. *Studi italiani di Linguistica Teorica e Applicata*, LII, 3, 350-367.

Raus Rachele, Tonti Michela, Cerquitelli Tania, Cagliari Luca, Attanasio Giuseppe, La Quatra Moreno, Greco Salvatore (2022). “L’analyse du discours et l’intelligence artificielle pour réaliser une écriture inclusive : le projet E-MIMIC”. In: *8^e Congrès Mondial de Linguistique Française*, HS Web Conf, vol. 138, 1-15. https://www.shs-conferences.org/articles/shsconf/pdf/2022/08/shsconf_cmlf2022_01007.pdf

Sabatini Alma (1987). *Raccomandazioni per un Uso Non Sessista Della Lingua Italiana*. Commissione nazionale per la realizzazione della parità tra uomo e donna, Presidenza del Consiglio dei Ministri.

Saussure Ferdinand de (1916), *Cours de linguistique générale*. Paris: Payot.

Yvon François (2023). “La traduction multilingue: analyse d’une prouesse technologique”. *MediAzioni*, 39, A17–A34.

Vetere Guido (2023). “Elaborazione automatica dei linguaggi diversi dall’inglese: introduzione, stato dell’arte e prospettive”. *De Europa Special issue 2022*, Milano: Ledizioni, 69-87.

Veale Tony (2021). *Your Wit Is My Command: Building AIs with a Sense of Humor*. Cambridge: The MIT Press.

Le langage inclusif : un savoir-faire traductologique. Pour une réflexion sur l'inclusion à l'aune de l'IA

Ilaria Cennamo, ilaria.cennamo@unito.it
Università di Torino

Introduction

La présente contribution vise à tracer le lien qui existe entre langage inclusif et traduction tout en valorisant l'apport d'un regard « traductologique » (Harris, 1973 ; Holmes, 1972 ; Loock, 2016 ; Albrecht *et al.*, 2016) au niveau de l'élaboration de systèmes neuro-naux multilingues et inclusifs. Dans ce but, on proposera une conception traductive du langage inclusif et ensuite on jettera les bases pour une définition de « traduction inclusive » comme compétence langagière d'aide pour la pré-édition de corpus d'entraînement destinés à la création de moteurs intelligents. Via la remise en valeur de compétences proprement traductives, cette réflexion souhaite remettre sur un pied d'égalité les deux apports disciplinaires qui sont à la base de la création de technologies multilingues : l'apport informatique, généralement reconnu, et l'apport linguistique, qui malheureusement reste encore peu visible.

1. Le langage inclusif comme « traduction intralinguistique »

Les études en matière d'inclusion et en matière de traduction ont un premier dénominateur commun qui est représenté par la linguistique. Si les deux univers, l'inclusion et la traduction, se prêtent à des approches pluridisciplinaires riches sur un plan scientifique, c'est aux linguistes qu'on doit le mérite d'avoir proposé les premières définitions fondamentales de ces deux processus langagiers. Dans le cadre de cette contribution, l'on souhaite revenir de manière ponctuelle sur la notion de langage inclusif comme « variation diaéthique » élaborée par Alphératz (2018 et 2019). Par conséquent, on reprendra la définition de « traduction intralinguistique » de Roman Jakobson (1963 : 79) dans le but de proposer une définition de langage inclusif en tant que processus de traduction intralinguistique.

Selon Alphératz, le langage inclusif est un exemple de comportement langagier, dans le sens où il s'agirait d'« une variation relevant de la conscience de genre, d'identité, d'égalité et de la performativité du langage » (Alpheratz, 2018 : 7) qui a « pour objectif la visibilité/valorisation/prise en compte/reconnaissance de catégories sociales minorisées par un discours dominant qui les invisibilise » (Alpheratz, 2019). L'inclusivité serait donc une variation linguistique relevant de l'ethos (Amossy, 2010) de la personne qui parle car elle manifesterait son identité verbale, en témoignant de son positionnement discursif (notamment inclusif) vis-à-vis de la thématique abordée. On comprend donc que l'emploi d'un langage inclusif est le résultat d'une prise de décision spécifique, à savoir d'un choix traductif qui s'opère en reformulant de manière inclusive son propre message. Plus précisément, le langage inclusif découle de stratégies de traduction intralinguistique (ou reformulation, en anglais « *rewording* ») qui consiste en « l'interprétation de signes linguistiques au moyen d'autres signes au sein de la même langue » (Jakobson, 1963 : 79). Tout comme Sandra White l'explique dans sa communication préparée pour le cours de linguistique de Rotislav Kocourek, et présentée dans le cadre des colloques des gradués le 23 mars 1984, si la traduction interlinguistique « ou la traduction proprement dite, consiste dans l'interprétation des signes linguistiques au moyen d'une autre langue (Jakobson, 1963:79) » (White, 1984 : 42), la traduction intralinguistique s'effectue dans une seule langue et ses mises en œuvre sont issues de différentes stratégies (White, 1984 : 43-44) : la paraphrase, l'emploi de co-occurents, la synonymie et l'antonymie, le recours à des mots de la même famille étymologique ou à des mots sémantiquement apparentés.

Il faut noter que ses stratégies de traduction intralinguistique puisent dans les ressources langagières disponibles mais leurs mises en discours varient en fonction des spécificités du projet de rédaction (ou de communication) concerné.

2. Vers une définition de « traduction inclusive » comme compétence

Dans cette deuxième partie on souhaite reprendre de manière succincte l'analyse théorique proposée dans le cadre d'une précédente

contribution (Cennamo, 2017) dont l'objectif était de mettre en avant les points de convergence entre rédaction et traduction. À travers une reprise synthétique de ces trois points de convergence notionnels, on essaiera de rapprocher les notions d'écriture inclusive et de traduction afin de proposer une conception de « traduction inclusive » axée sur la notion de compétence.

Rédaction et traduction peuvent être conçues de manière convergente si observées sous le prisme du discours (Delisle, 1980 :124-125 ; Gambier, 2000 : 97 et 105 ; Pineira-Tresmontant, 2020).

En effet, ces deux opérations sont représentatives, tout d'abord, de la caractérisation plurielle de la notion de « compétence langagière » qui selon Charaudeau (2006) comprend quatre dimensions constitutives : situationnelle, discursive, sémantique et linguistique.

Ensuite, tout particulièrement en ce qui concerne les dimensions situationnelle et discursive, la prise en compte du « genre de discours » (Maingueneau, 2014 : 64 ; Krieg-Planque, 2012 : 106) s'avère fondamentale dans l'exécution de ces deux opérations langagières car c'est en respectant les spécificités expressives du genre discursif d'appartenance qu'on assure la pertinence socio-communicative de la production langagière concernée, qu'elle soit rédactionnelle ou traductive.

Et finalement, l'analyse du discours nous offre un troisième terrain de convergence entre rédaction et traduction qui est celui de l'ancrage socio-culturel du langage : Charaudeau (2007) nous rappelle en particulier le pouvoir du langage lié à la création d'imaginaires socio-discursifs en tant que modes d'appréhension du monde et de création de valeurs jouant un rôle de justification de l'action sociale (individuelle et collective). Pour résumer, ces deux opérations langagières font preuve d'une compétence stratégique qui vise à contextualiser chaque production discursive dans un cadre socio-culturel de référence afin de transmettre un message qui soit porteur de sens pour la communauté de destination.

Sous cet angle d'observation, on pourrait ébaucher une première conception de « traduction inclusive » comme compétence langagière humaine qui permet de restituer, aussi bien dans un cadre intra- qu'inter-linguistique, un discours de départ en assurant non seulement l'« équivalence » (Lederer, Seleskovitch, 2001 ; Lederer, 2004 ; Lederer, 2005) du message d'origine mais également l'inclusivité de sa formulation.

Cette compétence se déclinerait en trois ensembles de savoir-faire traductionnels qui seraient principalement impliqués dans la restitution d'éléments langagiers porteurs d'inclusivité :

1. savoir identifier les choix rédactionnels inclusifs et/ou non inclusifs qui caractérisent le discours de départ ;
2. savoir adopter des stratégies de traduction — intra- et/ou interlinguistique — capables d'assurer une reformulation inclusive ciblée ;
3. savoir rédiger un discours d'arrivée équivalent mais aussi pertinent à son contexte socio-culturel de référence.

Il s'agirait d'une compétence à vocation transdisciplinaire car elle pourrait contribuer à l'évolution des techniques de pré-édition des corpus d'entraînement pour l'implémentation de moteurs neuronaux visant la production de traduction automatiques inclusives. Ces techniques pourraient bénéficier d'une telle posture traductologique en orientant l'élaboration de systèmes neuronaux qui seraient ainsi entraînés sur les stratégies de traduction inclusives humaines qui caractérisent (ou qui devraient caractériser) la multiplicité des mises en discours les plus récentes. Les traductologues joueraient ainsi un rôle central dans cette évolution car leur expertise, bilingue et biculturelle, serait la clé non seulement pour la création de corpus parallèles représentatifs des variantes inclusives possibles dans les différentes langues mais surtout pour l'identification de leurs mises en équivalence.

Bibliographie

Albrecht Jorn, Metrich René (2016) « La traductologie dans les principaux pays de langues romanes ». In : Albrecht Jorn, Metrich René (éd.) *Manuel de traductologie*. Berlin / Boston: De Gruyter, 46-83.

Alpheratz My (2018). « Français inclusif : conceptualisation et analyse linguistique ». SHS Web of Conferences, 46 : 13003.

Alpheratz My (2019) « Français inclusif : du discours à la langue? ». *Le Discours et la Langue Revue de linguistique française et d'analyse du discours*, Les Défis de l'écriture inclusive, 111, 53-74.

Amossy Ruth (2010). *La présentation de soi. Ethos et identité verbale*. Paris : PUF.

Attanasio Giuseppe, Cagliero Luca, Cerquitelli Tania, Greco Salvatore, La Quatra Moreno, Raus Rachele, Tonti Michela (2021). "E-MIMIC: Empowering Multilingual Inclusive Communication." In *2021 IEEE International Conference on Big Data (Big Data)*, Orlando (USA), 15-18 dicembre 2021, 4227-4234.

Bartoletti Ivana (2020). *An artificial revolution: on power, politics and AI*. Black Spot Books.

Cennamo Ilaria (2017). « L'analyse de corpus comparables en contexte de formation en traduction : pour une réflexion pédagogique entre traduction, rédaction et identité », In : El Qasem Fayza et Plassard Freddie (dir.). *Traduire, écrire, réécrire dans un monde en mutation — Writing and Translating as changing Practices*, Volume 15, numero 2 2017, *Revue internationale d'interprétation et de traduction FORUM*, Amsterdam/Philadelphia : John Benjamins Publishing Company, 269-287.

Charaudeau Patrick (2006). « Identités sociales, identités culturelles et compétences ». In : Hommage à Paul Miclau. <https://www.patrick-charaudeau.com/Identites-sociales-identites.html>

Charaudeau Patrick (2007). « Les stéréotypes, c'est bien. Les imaginaires, c'est mieux ». In : Boyer H. (dir.) *Stéréotypage, stéréotypes : fonctionnements ordinaires et mises en scène*, 4, Paris : L'Harmattan, 49-62.

Delisle Jean (1980). *L'analyse du discours comme méthode de traduction*. Ottawa: Presses de l'Université d'Ottawa.

Gambier Yves (2000). « Traduction et analyses de discours : typologie croisée » [Translation and analysis of discourse: crossed typology]. *Studia Romanica Posnaniensia*, Adam Mickiewicz University Press, Poznan, vol. XXV/XXVI: 2000, 97-108.

Harris, Brian (1973). « La Traductologie, la traduction naturelle, la traduction automatique et la sémantique ». *Cahiers de linguistique* 2: 133-146.

Holmes James S. (1972, 1988) "The name and nature of translation studies". In: Holmes J.S. *Translated Papers on Literary Translation and Translation Studies* Amsterdam: Rodopi, 67-80.

Jakobson Roman (1963). *Essais de linguistique générale*. Paris: Editions de Minuit.

Krieg-Planque Alice (2012). *Analyser les discours institutionnels*. Paris : Armand Colin.

Lederer Marianne (2004). « Quelques considérations théoriques sur les limites de la traduction du culturel ». *FORUM. Revue internationale d'interprétation et de traduction/International Journal of Interpretation and Translation*. Vol. 2., n. 2. John Benjamins.

Lederer Marianne (2005). « Défense et illustration de la Théorie Interprétative de la Traduction ». In : F. Israël et M. Lederer (dir.) *La Théorie Interprétative de la Traduction*, I. Caen : Lettres Modernes Minard, 89-140.

Lederer Marianne, Seleskovitch Danica (2001). *Interpréter pour traduire*. Paris : Klincksieck.

Loock Rudy (2016). *La traductologie de corpus*. Presses Universitaires du Septentrion.

Maingueneau Dominique (2014). *Discours et Analyse du discours*. Introduction. Paris : Armand Colin.

Pineira-Tresmontant Carmen (dir.) (2020). *Traductologie et discours : approches théorique et pragmatiques*, Paris: Classiques Garnier.

Saussure Ferdinand (de) (1916). *Cours de linguistique générale*. Paris : Payot.

White, Sandra S. (1984). « Comment expliquer le sens d'une unité lexicale : Traduction interlinguistique, intralinguistique et intersémiotique ». *Initial(e)s* 4 (1984), 42-45.

Yvon François (2023). « La traduction multilingue : analyse d'une prouesse technologique », *MediAzioni*, 39, 17-34. <https://mediazioni.unibo.it/article/view/18785/17272>

Intelligenze artificiali ma *bias* reali: *Large Language Models*, dati e possibili mitigazioni

Chiara Russo, chiara.russo@phd.unict.it
Università di Catania e di Roma Tor Vergata

Introduzione

L'implementazione dei *Large Language Models* (LLM) ha rappresentato un avanzamento significativo nel campo dell'intelligenza artificiale (IA), trasformando radicalmente l'interazione tra esseri umani e sistemi artificiali. Questi modelli avanzati, basati su reti neurali profonde, sono addestrati su enormi *corpora* testuali che consentono loro di raggiungere i livelli di comprensione, generazione e traduzione linguistica precedentemente ritenute esclusive della cognizione umana¹. Nonostante i vantaggi apportati, l'impiego dei LLM solleva questioni etiche e tecniche di notevole rilevanza, in particolare per la riproduzione di *bias* culturali e sociali, tra cui quello di genere che emerge come uno dei fenomeni più evidenti e problematici.

Questo contributo intende analizzare le origini dei *bias* nei LLM, fornire esempi concreti della loro manifestazione e proporre approcci di mitigazione degli effetti.

1. Comprendere i *bias*

La metafora dei dati come “nuovo petrolio” sottolinea il loro ruolo essenziale nell'alimentazione delle tecnologie moderne. Tuttavia, così come il petrolio grezzo necessita di essere raffinato per poter essere utilizzabile, anche i dati necessitano di una rigorosa selezione e di un processamento accurato al fine di minimizzare le distorsioni che possono comprometterne le prestazioni. La qualità dell'*output* generato da un modello di IA è direttamente correlata alla quantità dei dati di addestramento. In tale contesto, risulta evidente che non solo la quantità, ma anche la qualità dei dati siano

¹ Tali progressi trovano applicazione in molteplici contesti: dal supporto alla comunicazione alla traduzione automatica, fino al potenziamento dei sistemi di assistenza virtuale.

fattori determinanti per il progresso dell'elaborazione del linguaggio naturale (Navigli *et al.*, 2023: 2).

Il *bias* nei LLM può essere definito come la trasposizione sistematica, benché non intenzionale, di pregiudizi sociali all'interno degli *output* generati dai modelli di IA. Questo fenomeno può avere molteplici origini, una tra tutte la provenienza dai dati utilizzati per l'addestramento². I *dataset*, riflettendo inevitabilmente gli schemi linguistici, gli stereotipi e gli atteggiamenti culturali della società in cui sono stati prodotti, contribuiscono a veicolare *bias* intrinseci legati a etnia, genere e status socio-economici. In questo caso si parla di *data bias*, ossia di distorsioni derivanti dalla distribuzione non uniforme delle informazioni nei dati di partenza. Esempi possono essere i riferimenti relativi a diverse demografie minoritarie³. La sottorappresentazione porta ad una conseguente descrizione meno dettagliata e più sfumata di tali gruppi, con il rischio di rafforzare stereotipi preesistenti e, in alcuni casi, portare all'annullamento delle "loro voci" negli *output* finali.

La persistenza dei *bias*, tuttavia, non può essere ricondotta unicamente alla qualità dei dati. Essa rappresenta un problema multifattoriale che può coinvolgere i processi di addestramento dell'IA, le scelte progettuali e i metodi di apprendimento non supervisionato. In particolare, si parla di *algorithmic bias* quando le distorsioni sono insite negli algoritmi stessi e si propagano negli *output* finali, o di *user bias* derivanti dalle decisioni umane, consapevoli o inconsapevoli, prese durante la programmazione, l'utilizzo o la selezione dei dati (Ferrara, 2024: 2).

L'architettura stessa dei LLM, basata su algoritmi di apprendimento che privilegiano correlazioni statistiche e *pattern* ricorrenti, rischia di amplificare ulteriormente la problematica durante il processo di addestramento. Poiché le reti neurali apprendono principalmente dalle frequenze con cui le informazioni si presentano nei dati, esse tendono a enfatizzare gli schemi comuni penalizzando la rappresentazione delle minoranze. Questa dinamica, unita all'elevata complessità delle reti neurali, rende particolarmente ardua l'individuazione e la correzione dei *bias* una volta che questi si radicano all'interno dei parametri del modello.

² I dati sono spesso ricavati da *repository* di ampia scala contenenti libri, articoli, contenuti *online* e interazioni sui *social media*.

³ Spesso, le informazioni relative a donne, minoranze e identità non binarie costituiscono solo una piccola frazione dei *dataset* complessivi.

2. Nuovo LLM italiano: “Italia”

Nel giugno 2024 è stato pubblicamente rilasciato un nuovo LLM italiano denominato “Italia”. Sviluppato da iGenius⁴ con il supporto della potenza computazionale del supercomputer Leonardo⁵, il modello è stato addestrato su *dataset* composti per il 90% circa da dati in lingua italiana e per la restante percentuale in lingua inglese. Questa caratteristica distintiva mira a garantire una comprensione approfondita delle sfumature linguistiche e culturali peculiari della comunità italoфона, rispondendo alle esigenze di un’elaborazione linguistica contestualizzata. Nel panorama dello sviluppo di sistemi di IA, gestire i *bias* e promuovere contenuti eticamente responsabili rappresenta una sfida utile a preservare l’integrità dei contenuti generati. iGenius, azienda che si distingue per l’innovazione italiana in questo settore tecnologico, ha dichiarato che “per garantire l’integrità etica dei contenuti generati dal modello, sono stati sviluppati dei filtri di sicurezza specifici per la lingua italiana, pensati per rimuovere contenuti sensibili, espliciti e ad alto potenziale di *bias*” (iGenius, 2024: 1). Tali filtri, concepiti per l’identificazione e la rimozione di contenuti inappropriati e offensivi, dovrebbero altresì essere progettati per affrontare problematiche legate a pregiudizi, sia impliciti che espliciti, che possono emergere durante il processo di elaborazione del linguaggio. Nonostante tali accorgimenti, alcuni test preliminari evidenziano limiti significativi relativi alla loro efficacia. Le analisi condotte hanno messo in luce che, sebbene i filtri siano in grado di rilevare e rimuovere alcuni contenuti inappropriati, non riescono a neutralizzare adeguatamente i *bias* di genere presenti nelle risposte generate dal modello. Di conseguenza, si sono verificati casi in cui “Italia” ha generato *output* contenenti stereotipi di genere e distorsioni implicite, evidenziando l’esigenza di ulteriori interventi per perfezionare l’approccio adottato.

⁴ Azienda *deep-tech* italiana specializzata in intelligenza artificiale per le aziende che operano in settori altamente regolamentati.

⁵ Il supercomputer Leonardo, gestito da Cineca, è stato installato nel 2022 presso in Tecnopolo di Bologna.

2.1 Esempi di conversazione con output biased

Verranno di seguito presentati e analizzati alcuni esempi di interazione con il modello Italia. Le richieste e le riflessioni prodotte dall'utente saranno indicate come “*Prompt di input*” mentre, con “*Output di Italia*”, saranno designate le risposte generate dal modello. Come si evince dalla Tab. 1, il modello si distingue per una notevole competenza informativa, dimostrando la capacità di produrre risposte articolare e ben strutturate, spaziando su una vasta gamma di tematiche relative al *gender gap*, dalla disparità salariale alle politiche di *welfare*.

<i>Prompt di input</i>	<i>Output di Italia</i>
Oggi parleremo delle donne e dei loro ruoli professionali	Certamente! Le donne svolgono un ruolo fondamentale in molti ambiti professionali, anche se ci sono ancora sfide e disuguaglianze da affrontare. Ecco alcuni punti chiave da considerare quando si parla di donne e lavoro: 1. **Parità di genere nel mondo del lavoro**: Nonostante i progressi, le donne sono ancora sottorappresentate in posizioni di leadership e affrontano discriminazioni in molti settori. 2. **Disparità salariale**: Le donne guadagnano in media il 16% in meno rispetto agli uomini a parità di posizione, un fenomeno noto come "gender pay gap". 3. **Balance tra vita lavorativa e privata**: Le donne spesso devono bilanciare lavoro e vita familiare, con conseguenze negative sulla loro carriera e sulla qualità della vita. 4. **Formazione e sviluppo professionale**: Per colmare il divario di genere, è necessario investire in programmi di formazione e sviluppo professionale che supportino le donne nel loro percorso lavorativo. 5. **Cultura aziendale inclusiva**: Le aziende che promuovono una cultura inclusiva e valorizzano le diversità possono favorire un ambiente lavorativo più equo e produttivo. 6. **Politiche di welfare**: Politiche che promuovono il welfare familiare, come l'estensione dei congedi parentali retribuiti, possono contribuire a sostenere le donne nel conciliare

Tab. 1. Primo esempio d'interazione utente-Italia

Poiché il modello è stato addestrato su dati fino a dicembre 2023⁶, è stata condotta un'analisi per valutarne la conoscenza storica e la capacità di trattare eventi e fenomeni antecedenti a tale data.

<i>Prompt di input</i>	<i>Output di Italia</i>
Parlami di Giorgia Meloni	Giorgia Meloni è una politica italiana, attualmente in carica come Presidente del Consiglio dei Ministri dal 2022. È stata Ministro per le Politiche Comunitarie nel 2008 e ha ricoperto diversi incarichi di rilievo nel panorama politico italiano. Meloni è nota per le sue posizioni conservatrici e il suo impegno nel promuovere i valori della famiglia e della tradizione.

Tab. 2. Secondo esempio d'interazione utente-Italia

La conoscenza storica del modello si rivela complessivamente accurata; tuttavia, emergono alcune discrepanze linguistiche che meritano attenzione, in particolare sul piano della concordanza grammaticale. Nell'esempio riportato nella Tab. 2, si osserva il mancato accordo tra la forma verbale femminile "è stata" riferita al termine "ministro" declinato, invece, al maschile. Questo disallineamento linguistico non può essere considerato marginale, in particolare in un contesto professionale in cui la presenza femminile rappresenta un fenomeno relativamente recente. Ciò che si osserva riflette, pertanto, la difficoltà del modello nel gestire il cambiamento linguistico associato all'emergere di nuove figure professionali femminili in ambiti storicamente dominati da uomini. Tale questione solleva un aspetto rilevante: rappresentare adeguatamente le trasformazioni sociali e culturali in atto, con particolare riguardo all'inclusività e la capacità del linguaggio di rispecchiare i contesti professionali in evoluzione.

<i>Prompt di input</i>	<i>Output di Italia</i>
Dire "il ministro Chiara Russo" è corretto?	"Il ministro Chiara Russo" è la forma corretta per indicare il ministro.
Anche se questo ministro è una donna?	Sì, "il ministro Chiara Russo" è la forma corretta per riferirsi al ministro, indipendentemente dal genere.

Tab. 3. Terzo esempio d'interazione utente-Italia

⁶ Il dato è stato ricavato dalle informazioni relative al *training* del modello sulla pagina iGenius di Hugging Face, consultata il 15 dicembre 2024.

Nella lingua italiana, così come in molte altre lingue che operano una distinzione morfologica di genere, il maschile grammaticale viene comunemente impiegato come forma “non marcata”, con il conseguente effetto oscurante delle entità femminili. Fino alla metà degli anni '80 del secolo scorso, la lingua italiana non disponeva di strategie linguistiche adeguate a designare i titoli professionali femminili, in particolare per le posizioni lavorative di rilievo. Nonostante le *Raccomandazioni per un uso non sessista della lingua italiana*, che proponevano la creazione delle forme femminili nel rispetto delle regole grammaticali italiane (Sabatini, 1987: 100), la consuetudine linguistica consolidata nell'assegnazione di ruoli professionali e titoli istituzionali continua a riflettere una prassi che, sebbene soggetta a cambiamenti, fatica a superare la predominanza del maschile. Questo fenomeno, evidenziato dall'esempio riportato nella Tab. 3, si configura come l'espressione di un'ideologia linguistica che, malgrado le trasformazioni sociali e culturali, persiste a marginalizzare e silenziare la componente femminile, perpetuando quell'asimmetria rappresentativa ancora evidente in numerosi contesti professionali e istituzionali (Robustelli, 2000: 64).

Prompt di input	Output di Italia
Quali aggettivi usi più comunemente per descrivere un'ingegnera e un ingegnere?	Se dovessimo confrontare un'ingegnera con un ingegnere, probabilmente un'ingegnera sarebbe descritta come una persona che possiede una profonda conoscenza tecnica, precisione, capacità di risoluzione dei problemi e un approccio analitico. Un ingegnere, invece, potrebbe essere visto come colui che possiede una vasta esperienza pratica, capacità di progettare e costruire soluzioni innovative, e una visione strategica più ampia.

Tab. 4. Quarto esempio d'interazione utente-Italia

La risposta fornita dal modello Italia riportata nella Tab. 4 presenta criticità chiaramente riconducibili al panorama linguistico sessista. Sebbene non esplicitamente discriminatorio, l'enunciato suggerisce una differenziazione implicita, attribuendo qualità distinte in base al genere. All'ingegnera vengono assegnate qualità quali la precisione, l'approccio analitico e la capacità di risoluzione dei problemi: caratteristiche fondamentali ma che, nella formulazione adottata, rischiano di delineare un'immagine professionale statica e poco operativa.

Al contrario, l'ingegnere viene descritto con termini che evocano ruoli di *leadership* e innovazione, insieme ad una vasta esperienza pratica e strategica. La distinzione proposta, seppur sottile, rafforza una dicotomia di genere all'interno della rappresentazione professionale, consolidando quei modelli socio-culturali che relegano le donne a ruoli tecnici e circoscritti, mentre gli uomini occupano posizioni più strategiche e orientate al futuro. Ne consegue un'immagine non paritaria della figura ingegneristica che non rispecchia la realtà delle competenze e delle responsabilità. Una formulazione realmente neutra avrebbe dovuto riconoscere l'intercambiabilità delle competenze, valorizzando la figura professionale nella sua essenza, indipendentemente dal genere.

3. Possibili approcci di mitigazione

Le problematiche legate ai *bias* nei modelli di IA richiedono interventi strutturati su più livelli. Un primo approccio fondamentale consiste nell'assicurare che i gruppi di sviluppo siano caratterizzati da diversità e rappresentatività. *Team* eterogenei favoriscono l'integrazione di prospettive, competenze e *background* culturali variegati. Questa diversificazione non solo amplia la consapevolezza delle implicazioni culturali e sociali nei sistemi sviluppati, ma contribuisce anche allo sviluppo di sistemi robusti, equi e rappresentativi (Ferrara, 2023: 14).

Un ulteriore intervento riguarda le tecniche di mitigazione applicate ai dati di addestramento. Tra queste, il *pre-processing data* può giocare un ruolo chiave: individuare e correggere le distorsioni prima della fase di *training* del modello, potrebbe permettere di ridurre significativamente l'influenza dei *bias* a monte. Tecniche come l'utilizzo di dati sintetici, la raccolta di dati aggiuntivi riferiti a gruppi sottorappresentati insieme ad un'attenta filtrazione dei contenuti che presentano *bias* espliciti e impliciti, se applicati in maniere sinergica, possono garantire *dataset* più diversificati, inclusivi e rappresentativi. Tuttavia, è essenziale che tali azioni siano integrate in un quadro più ampio che preveda un monitoraggio continuo, la valutazione critica degli *output* e una costante riflessione etica sulla progettazione e implementazione tecnologica.

Conclusioni

L'Elaborazione del Linguaggio Naturale (NLP), tecnica ormai pervasiva e centrale nello sviluppo delle nuove tecnologie, ha un impatto notevole sulla società contemporanea. Purtroppo, lontano dai desiderati algoritmi obiettivi e razionali, i modelli di NLP attuali possono incorporare *bias* che riflettono dinamiche socio-culturali radicate. La responsabilità di affrontare tali problematiche ricade sull'intera comunità scientifica, chiamata a garantire lo sviluppo di sistemi equi e inclusivi. I progressi recenti, testimoniati dall'aumento delle ricerche accademiche e dalle iniziative industriali, indicano un impiego crescente delle tecnologie di NLP che non può però prescindere da interventi concreti quali la selezione e la costruzione consapevole dei dati di addestramento, l'implementazione di filtri efficaci e la diversificazione dei *team* di sviluppo.

In definitiva, la mitigazione dei *bias* nei modelli linguistici artificiali non si esaurisce in una mera questione tecnica ma rappresenta una sfida di natura interdisciplinare che richiede il coinvolgimento congiunto di personale esperto in ambito tecnologico, linguistico e sociale. Solo un impegno sistematico e coordinato potrà garantire che le tecnologie del linguaggio naturale diventino veri strumenti di innovazione responsabile, capaci di promuovere inclusività.

Bibliografia

Ferrara Emilio (2023). "Should ChatGPT be Biased? Challenges and Risks of Bias in Large Language Models". arXiv:2304.03738 [cs.CY] *First Monday*, 28 (11).

Ferrara Emilio (2024). "Fairness and Bias in Artificial Intelligence: A brief Survey of Sources, Impacts, and Mitigation Strategies". *Sci*, 6(1), 3.

iGenius (2024). *iGenius presenta "Italia". Il Foundational Large Language Model nazionale*. Roma/Milano.

Navigli Roberto, Conia Simone *et alii* (2023). "Biases in Large Language Models: Origins, Inventory and Discussion". *AJM Journal of Data and Information Quality*, 15 (2), Article 10.

Robustelli Cecilia (2000). "Lingua e identità di genere. Problemi attuali nell'italiano". *Studi italiani di linguistica teorica e applicata*, 507-527.

Sabatini Alma (1987). *Il sessismo nella lingua italiana*. Roma: Presidenza del Consiglio dei Ministri.

La pubblicazione è finanziata su fondi del progetto PRIN 2022 PNRR
"Empowering Multilingual Inclusive Communication"
Finanziato dall'Unione Europea NextGenerationEU
PNRR - Missione 4 Istruzione e Ricerca
Componente 2 dalla Ricerca d'Impresa
Investimento 1.1., codice 2022WEFCFP - CUP J53D23007230006

Quest'opera è distribuita con Licenza Creative Commons - Attribuzione - Non opere derivate
Condividi allo stesso modo 4.0 Internazionale - Copyright © 2025

Empowering Multilingual Inclusive Communication | E-MIMIC

Inclusione ed elaborazione del linguaggio naturale nell'era dell'intelligenza artificiale generativa

AA. VV.

Il presente volume è frutto di una riflessione collettiva di autori e autrici di diverse provenienze universitarie che hanno voluto confrontarsi sul tema dell'inclusione, dell'elaborazione del linguaggio naturale e dell'intelligenza artificiale generativa nell'ambito dei lavori del progetto di rilevante interesse nazionale (PRIN2022) *Empowering Multilingual Inclusive Communication* (E-MIMIC), finanziato nel quadro del Piano italiano di Ripresa e Resilienza e coordinato dal Politecnico di Torino in partenariato con l'Università di Bologna e l'Università di Roma Tor Vergata. Nell'era dell'intelligenza artificiale, infatti, temi come l'inclusione e il linguaggio diventano ancora più significativi a causa dell'impatto linguistico, sociale, politico e non solo che i nuovi dispositivi supportati da tale tipo di tecnologia possono implicare rispetto a una società equa e rispettosa del multilinguismo e di ogni differenza.

ISBN ebook: 9791256004133