

Intelligenza artificiale: tecnologia non neutrale e sfide etiche per l'uomo

Silvia Crafa

Università di Padova

Laboratorio Nazionale CINI Informatica e Società

L'epoca che stiamo vivendo è caratterizzata dai profondi cambiamenti dovuti al ruolo pervasivo che le tecnologie digitali hanno assunto in tutti gli ambiti della vita e delle strutture sociali. Nel chiedersi cosa renda questa tecnologia così rilevante rispetto ad altre, J.H. Moor nel 1985 sottolineò che ciò che rende rivoluzionari i computer non è la loro velocità, pervasività o potenza di calcolo, ma la loro malleabilità logica: possono essere programmati per eseguire qualsiasi compito rappresentabile in termini di input/output e operazioni logiche. Poiché la logica si applica ovunque, il computer è dunque simile ad uno strumento universale, e proprio come l'energia della macchina a vapore è stata la risorsa grezza della Rivoluzione Industriale, così la logica del computer è la risorsa grezza della Rivoluzione Digitale [1]. In altri termini, il digitale è una *tecnologia abilitante*, che permette cioè la nascita e lo sviluppo di nuove classi di strumenti, processi, applicazioni, ma anche nuove idee, nuova scienza, e dunque nuova politica e nuova cultura per comprendere e governare questo sviluppo. Allo stesso tempo, analizzare e valutare la rivoluzione che stiamo vivendo è particolarmente difficile perché, essendone immersi e vivendone in prima persona gli effetti e i rapidi cambiamenti, non è possibile avere uno sguardo completo e oggettivo. Possiamo però individuare alcune prospettive sulle tecnologie digitali che ci permettano di ancorare in modo efficace le riflessioni e mantenere un approccio pragmatico all'interno di un orizzonte di senso più ampio.

Non è solo una questione di uso corretto di strumenti

Un atteggiamento molto diffuso è quello di considerare la tecnologia come un insieme di strumenti di cui è importante fare un uso positivo, cioè un uso etico, critico e consapevole. Per quanto condivisibile, questo approccio non tiene in considerazione il fatto che la tecnologia e i suoi strumenti non sono mai neutri [2,3,4,5,6,7], ma, essendo il risultato di specifiche scelte progettuali e operative, nascono già orientati a scopi precisi e veicolano -anche in modo implicito- specifici valori. Consideriamo ad esempio un software che analizza il lavoro dei dipendenti di un'azienda e associa loro un indice di produttività espresso con un numero. Il solo fatto di aver scelto di usare come output dei numeri fornisce un ordinamento totale dei dipendenti, rendendo tutti gli individui direttamente comparabili tra loro (indipendentemente dalla questione se ciò abbia senso). Per quanto tale comparazione assoluta in molti casi non sia stata richiesta, la sola disponibilità di questo dato permette degli usi (e abusi) che non sarebbero possibili progettando un output diverso da un semplice numero. Come ulteriore esempio consideriamo i software scolastici Google Classroom o il registro elettronico come ClasseViva: sono certamente strumenti utili per la gestione delle attività scolastiche, ma le funzionalità che contengono o non contengono, la maggiore o minore semplicità e flessibilità nell'uso e perfino i colori e le forme usate nell'interfaccia orientano verso determinati modi di utilizzo che influenzano la relazione tra studenti, insegnanti e famiglie. Ad esempio, la progettazione dell'interfaccia può favorire oppure escludere persone con disabilità, la sospensione automatica delle notifiche e dei messaggi fuori da certe fasce orarie può influenzare la comunicazione dentro e fuori la classe; inoltre l'eccesso di elementi di *gamification* [8,9,10] e il focus visuale sui valori numerici dei voti può spingere a dare ai voti scolastici un ruolo eccessivo rispetto all'andamento dello studio e della preparazione personale. Si pensi infine alla discussione sollevata dalla recente introduzione nel registro elettronico di una componente, per quanto opzionale, di pubblicità e giochi. Prima di chiedersi come usare

correttamente questi strumenti, va presa consapevolezza che tutti questi elementi di progettazione sono non neutrali.

Più in generale, gli strumenti tecnologici hanno sempre un ruolo di *mediatori delle nostre relazioni*: le relazioni con le altre persone, con gli spazi che ci circondano, la relazione con il nostro lavoro, perfino la relazione con noi stessi passa sempre di più attraverso strumenti tecnologici. E dunque in quanto mediatori di relazioni, sono strumenti che per loro natura influenzano -magari in modo positivo- i nostri comportamenti, i nostri modi di lavorare, di pensare e di essere; pertanto hanno conseguenze morali a prescindere dal modo in cui sono usati. È fondamentale quindi imparare ad individuare gli elementi di non neutralità insiti negli strumenti che usiamo, perché questo sguardo meno ingenuo permetterà di orientare meglio la decisione su quale sia un loro uso positivo, oppure la decisione di evitare, o addirittura vietare o rendere illegale, l'uso di strumenti che hanno intrinsecamente elementi problematici.

La non neutralità della tecnologia dell'Intelligenza Artificiale

L'intelligenza artificiale è una tecnologia profondamente non neutrale: lo è nella natura sostanzialmente statistica del suo funzionamento, come vedremo in dettaglio nel seguito, ma lo è anche per via dell'ecosistema politico e finanziario che, insieme alla ricerca, ne co-determina lo sviluppo e la distribuzione. Anche il nome "intelligenza artificiale" non è neutrale, ma è scelto appositamente per veicolare una specifica percezione di quelle che sono le effettive capacità di questa tecnologia. Al di là dei dibattiti su quanto siano avanzate queste capacità, l'antropomorfizzazione e la cattura culturale [11,12] di termini come "intelligenza", "apprendimento", "ragionamento" o "allucinazione" (che si riferisce in realtà ad un malfunzionamento) non aiutano la comprensione scientifica [13] e, puntando l'attenzione sulle macchine, opacizzano il ruolo e la responsabilità delle persone e delle concentrazioni di potere che guidano lo sviluppo e la diffusione di queste tecnologie [5,6,14]. Il primo passo dunque per capire cos'è l'intelligenza artificiale è rendersi conto che non si tratta solo di tecnica o di strumenti

disponibili all'uso, ma di un complesso *sistema socio-tecnico* che sta riconfigurando le infrastrutture sociali (incluse quelle educative), il mercato, le istituzioni e i sistemi di potere [15,16].

Rimanendo su un piano più matematico, invece del termine generico “intelligenza artificiale”, possiamo innanzitutto parlare con più precisione di algoritmi di “apprendimento automatico”, o *machine learning*. Si tratta di una classe di sistemi probabilistici che, con tecniche di inferenza induttiva anche molto sofisticate, riconoscono modelli statistici all'interno di enormi quantità di dati di esempio su cui sono “allenati”. La natura statistica di questi sistemi porta con sé delle assunzioni e dei limiti di funzionamento ben noti in matematica, tra cui la qualità e la rappresentatività dei dati, la differenza tra i dati di allenamento e quelli di uso effettivo, la difficoltà di testare in modo robusto la correttezza e la generalizzabilità dei risultati probabilistici [8,17]. Le caratteristiche intrinseche di questo tipo di funzionamento dovrebbero limitare le sue applicazioni a contesti con elevata tolleranza al rischio. Ad esempio, all'interno di una fabbrica robotizzata un prodotto malfunzionante per un errore di produzione può essere scartato prima di essere immesso nel mercato, mentre un errore nel software di ammissione al credito bancario o nel software di sorveglianza pubblica può cambiare la vita di una persona. I regolamenti europei GDPR e AI Act [18] mirano proprio a prevenire e controllare questi rischi, ma è importante essere consapevoli che essi sono intrinseci alla natura statistica del machine learning, diversamente ad esempio dal software del pilota automatico di un aereo, che, per quanto comunque soggetto ad errori e malfunzionamenti, ha invece natura deterministica ed è pienamente sotto il controllo di una logica esplicita di funzionamento.

Lo scarto tra i dati e la realtà

Molte delle criticità del machine learning sono riconducibili all'irriducibile scarto che c'è tra la realtà (su cui si intende operare) e i dati (su cui effettivamente si opera). I dati elaborati dagli algoritmi sono sempre rappresentazioni informatiche di natura numerica: ad

esempio un'immagine è un insieme di numeri, ciascuno che descrive il colore di una piccola porzione dell'immagine (un pixel). Il viso di una persona diventa un dato assimilando il viso ad un'immagine con in più le misure delle distanze tra alcuni punti di riferimento come occhi e naso. Analogamente, diventa un insieme di dati numerici anche la vita professionale di una persona, il suo comportamento di acquisto online, perfino il suo carattere e la sua propensione politica [20]: tutto viene tradotto in un insieme di dati numerici, su cui poi agiscono gli algoritmi di intelligenza artificiale. Chiaramente questo processo di *datificazione* diventa più problematico quanto più complessa e sfumata è la realtà che si vuole rappresentare numericamente. Non tutto è misurabile, ma la pressione sociale verso l'uso di algoritmi di machine learning e la disponibilità di enormi quantità di dati fanno sì che tutto sembri datificabile tramite un numero sufficiente di dati, finendo per dimenticare ciò che si è perso nello scarto tra la realtà e la sua rappresentazione. Un esempio concreto è la differenza tra la *valutazione* della preparazione scolastica di uno studente e la *misurazione* della sua preparazione tramite i suoi voti. Qualsiasi analisi, predizione, raccomandazione fornita dagli algoritmi di intelligenza artificiale si basa esclusivamente su dati, che pertanto incorporano necessariamente decisioni non neutrali rispetto a cosa viene scartato della realtà effettivamente sotto analisi.

La questione della datificazione è il cuore di numerose criticità delle scienze dei dati come la statistica e il machine learning: da essa discendono problemi come il campionamento degli esempi di allenamento, la formulazione del problema (proxy bias), la scelta dei dati da considerare outlier, la valutazione dei risultati in termini di accuratezza numerica. Per una trattazione molto ricca di come questi aspetti prettamente scientifici si traducano in problematiche concrete e dilemmi etici, rimandiamo al testo di K.Martin [19], menzionando qui solo un paio di aspetti.

La natura statistica dell'IA: correlazioni e outliers

Come già detto, l'apprendimento automatico si basa sull'estrazione di pattern statistici da enormi quantità di dati. Tra questi pattern uno dei più ricorrenti è l'estrazione automatica di correlazioni. È noto che inferire dai dati relazioni di causalità è estremamente difficile, mentre l'identificazione di semplici correlazioni è molto facile. Pertanto gli algoritmi di intelligenza artificiale usano le correlazioni come surrogati della causalità, sotto l'assunzione che un numero sufficientemente elevato di dati di input assicurerà l'assenza di correlazioni spurie lasciando nei dati solo quelle che effettivamente tracciano dipendenze causali. In questo modo ad esempio sono costruiti i software che valutano i cv e raccomandando l'assunzione di alcuni candidati ad un posto di lavoro, sulla base (anche) di specifiche correlazioni individuate nei migliori cv analizzati in fase di allenamento. Software di questo genere sono largamente in uso in tutto il mondo, ma funzionano sostanzialmente nello stesso modo della frenologia, la pseudo-scienza che nel XIX secolo usava la misura delle protuberanze del cranio delle persone per predirne, sulla base di correlazioni, i tratti mentali e di personalità. Al giorno d'oggi sembra impossibile utilizzare ancora questo approccio, ma l'enorme disponibilità di dati lo ha effettivamente reso pragmaticamente efficace: la qualità dei risultati degli algoritmi di intelligenza artificiale è sorprendente pensando che è ottenuta tramite correlazioni, senza alcuna comprensione della vera natura delle relazioni tra i dati analizzati. Ciò succede in particolare nel caso dei Large Language Models, che producono testi sensati e di qualità pur senza avere internamente alcuna rappresentazione semantica o grammatica del linguaggio: questi algoritmi sono stati allenati a riconoscere le probabilità di co-occorrenza delle parole (e di frazioni di parole, dette token) all'interno delle frasi, e a riprodurre queste probabilità nei testi che generano in output. Quindi, così come le correlazioni sono usate come "surrogati a basso costo" della causalità, nei LLMs la probabilità di co-occorrenza di più parole è usato come efficace "surrogato a basso costo" della semantica di quelle parole. Internamente non c'è alcun modello di conoscenza

linguistica e alcuna comprensione di significati; i testi generati sono piuttosto delle “estruzioni di stringhe” che un osservatore umano riconosce come sensate e ben scritte. Va anche menzionato che per i principali LLMs molta della qualità dei testi prodotti dipende anche dall’enorme quantità di *micro-lavoratori umani* che rivedono, ripuliscono e raffinano sia i testi di input che quelli di output, nella maggior parte dei casi lavorando in situazioni di precarietà, sfruttamento e invisibilità generale [21,22,23,24].

Un altro elemento caratteristico del ragionamento statistico che vogliamo qui citare è la gestione degli *outliers*, cioè quei dati anomali, numericamente distanti dalla maggioranza dei dati raccolti, che vengono scartati perché interpretati come casi isolati che deviano dal pattern principale. Spesso dipendono da qualche errore isolato nella raccolta dei dati, o da casi molto marginali, quindi è preferibile non considerarli perché le statistiche che derivano da campioni contenenti outlier possono essere fuorvianti. Chiaramente è molto difficile -e non neutrale- la scelta di quali dati considerare come outlier e quali invece mantenere nel campione da cui estrarre pattern statistici. La scienza offre una metodologia per gestire questa scelta, che varia a seconda del tipo di tecnica statistica adottata. D’altra parte però, il pattern individuato statisticamente in un gruppo di dati di esempio (detto anche “modello dei dati”) non si adatta quasi mai perfettamente a tutti i casi di esempio ma solo alla larga maggioranza (a meno che non siano dati che rappresentano un fenomeno effettivamente molto regolare, e dunque pienamente descrivibile tramite un pattern preciso). Per loro natura quindi, gli algoritmi di analisi dei dati (che, ricordiamo, sono di tipo induttivo) faticano a gestire ciò che devia dalla regolarità statistica, e gli esempi che deviano in modo sostanziale sono scartati come outliers. Se pensiamo però che i dati-esempio manipolati dal machine learning possono rappresentare curriculum o volti di persone, comprendiamo come lo scarto di un outlier sia nella pratica lo scarto di una persona che ha un comportamento o caratteristica che deviano dalla media. Vediamo qui bene come un aspetto puramente tecnico-scientifico

diventa nella pratica la sorgente di discriminazioni e decisioni che non riteniamo etiche; e questo non succede per intenzioni malevole dei programmatori o per errori dell'algoritmo, ma proprio per la sua natura statistica.

Al contrario degli algoritmi, gli uomini sono in grado di gestire le anomalie e le eccezioni in modo più robusto, comprendendo quando ci si trova di fronte ad un caso che richiede un supplemento di riflessione. È proprio qui che opera infatti l'intelligenza, cioè quella capacità squisitamente umana di guardare alle regole e ai protocolli tenendo insieme rigore ed elasticità. A differenza di una macchina, un professionista è in grado di *apprendere dall'esperienza*, che differisce dall'"apprendimento dai dati" perché l'esperienza è la capacità di *riflettere* sulle osservazioni e non semplicemente estrarre pattern statistici da loro rappresentazioni numeriche. L'intelligenza umana infatti è molto più che la capacità di calcolo, ma include aspetti come la capacità di giudizio, l'intuizione e la saggezza.

Essere consapevoli di ciò a cui vogliamo dare valore

Dopo aver spostato l'attenzione dall'uso degli strumenti al riconoscimento dei loro elementi di non neutralità, serve fare un altro passo: serve avere dei criteri per determinare se la non neutralità identificata è positiva e desiderabile (come ad es. una protesi innovativa per persone con disabilità) oppure è addirittura pericolosa (ad es. gli algoritmi dei social network progettati per creare dipendenza negli adolescenti). Tale discriminazione entra nel campo del giudizio morale, dei principi normativi (ad es. il principio di precauzione) e delle norme legali. È importante ricordare che l'intreccio complesso tra morale, etica e diritto è un tema non certo nuovo per la filosofia e la giurisprudenza, che da secoli affrontano e gestiscono la regolamentazione della tecnologia. Pertanto, anche l'intelligenza artificiale più innovativa non arriva nel vuoto normativo e, anche se certamente richiede adattamenti e innovazioni, non ha tutti quegli elementi di eccezionalità che spesso accompagnano la sua narrazione [5,8].

Senza entrare nella complessità della questione, citiamo nuovamente le riflessioni di J.H. Moor [1], che ricorda come, sempre in analogia con la rivoluzione industriale, anche la rivoluzione digitale ha attraversato due fasi. Nella prima fase le innovazioni sono introdotte, testate, migliorate e industrializzate in alcuni segmenti economici, mentre nella seconda fase la tecnologia diventa parte integrante delle attività umane e delle istituzioni sociali, fino a trasformarle. Pertanto, nella prima fase la domanda di fondo è “Quanto bene un computer riesce a svolgere questa attività?”, mentre nella seconda fase la domanda diventa “*Qual è la natura ed il valore di questa attività*, ora che anche un computer sembra in grado di farla?”. In altre parole, uno sguardo consapevole sui cambiamenti portati dalle tecnologie digitali non può sfuggire alle nuove *domande di senso* che si pongono: ad esempio, ora che la didattica in classe si avvale di strumenti digitali, qual è la natura ed il valore di una lezione in classe? Ora che la comunicazione dei risultati di una verifica può avvenire mediante notifica automatica, qual è la natura ed il valore della comunicazione di un voto? Ora che la scuola non è più l’unico luogo in cui si imparano nuovi contenuti, qual è la natura ed il valore della scuola? E domande analoghe si possono porre per tutti gli ambiti della vita personale e sociale.

Per quanto complesse, queste domande di senso non possono essere eluse, perché solo avendo chiaro a cosa vogliamo dare valore, come individui e come cittadini, possiamo davvero governare i cambiamenti portati dalla rivoluzione digitale. È bene sottolineare che non si tratta di domande astratte e filosofiche, ma di un discernimento alla base di scelte quotidiane molto concrete, ad esempio su come usare le applicazioni di intelligenza artificiale. Per orientarsi tra le difficoltà di queste riflessioni etiche, menzioniamo qui semplicemente due aspetti. Il primo è che il modo migliore di usare l’intelligenza artificiale si trova quando non ci si sente obbligati a usarla, sottraendosi cioè alla pressione sociale contemporanea che spesso vede nell’innovazione tecnica un valore positivo aprioristico invece che una semplice opzione aggiuntiva. Il secondo aspetto è che, in particolare nell’ambito della didattica scolastica, una bussola

efficace è data dal chiedersi in quale ruolo è posto l'uomo, lo studente, l'insegnante, in una data tecnologia o applicazione, e in che modo ne è influenzato. Ad esempio, invece di porre domande come "WhatsApp è buono o cattivo?" oppure "Usare ChatGPT in questa situazione è giusto o sbagliato?", è preferibile chiedersi "L'interazione tramite WhatsApp come influisce sul mio rapporto con i compagni?" o "L'uso di ChatGPT ha aggiunto qualità al mio lavoro?", riformulando cioè le riflessioni in modo che il soggetto centrale sia sempre la persona e non lo strumento [25].

Conclusioni

Come abbiamo visto, l'intelligenza artificiale non è solo una questione di tecnologia, ma il modo in cui funziona, in cui è progettata, è distribuita ed è usata è un complesso intreccio di aspetti tecnici e di scelte umane intrise di valori. Ogni scelta inoltre è la scelta di qualche persona, inclusa la scelta di delegare alle macchine le decisioni. È importante quindi che la responsabilità per le più rilevanti tra queste scelte non resti nascosta all'interno dei processi di sviluppo o nelle dinamiche economiche e geopolitiche dell'IA.

Anche da un punto di vista didattico, la complessità delle sfide socio-tecniche poste dall'intelligenza artificiale non può dunque essere affrontata se non con un approccio interdisciplinare. La frammentazione del sapere e la specializzazione scientifica portano a focalizzare l'attenzione sulle applicazioni, i loro usi e i loro miglioramenti possibili, ma conducono a perdere lo sguardo d'insieme, fino al rischio di una deriva scienziata. La scienza può analizzare e conoscere un fenomeno, ma non può indagarne il senso, che richiede invece il linguaggio e il metodo delle discipline umanistiche [26,27]. È dunque attraverso la costruzione del dialogo, tra discipline e tra persone, che passa la sfida di mantenere l'intelligenza artificiale al servizio dei valori umani.

Bibliografia/Sitografia

- [1] J. H. Moor (1985) *What is computer ethics?* Metaphilosophy, 16(4). <https://web.cs.ucdavis.edu/~rogaway/classes/188/spring06/papers/moor.html>
- [2] L. Winner (1980) *Do artifacts have politics?* Daedalus, 109(1).
- [3] K. Crawford (2018) *You and AI – Just an Engineer: The Politics of AI*. Royal Society. <https://royalsociety.org/topics-policy/projects/machine-learning/you-and-ai/>
- [4] K. Crawford (2021) *Atlas of AI: Power, Politics, and the Planetary Costs of Artificial Intelligence*. Yale University Press
- [5] D. Tafani (2023) *Sistemi fuori controllo o prodotti fuorilegge*. Bollettino telematico di filosofia politica, maggio 2023. <https://btfp.sp.unipi.it/it/2023/05/sistemi-fuori-controllo-o-prodotti-fuorilegge/>
- [6] D. Tafani (2024) *Omini di burro. Scuole e università al Paese dei Balocchi dell'IA generativa*. Bollettino telematico di filosofia politica, ottobre 2024. <https://btfp.sp.unipi.it/it/2024/10/omini-di-burro-scuole-e-universita-al-paese-dei-balocchi-dellia-generativa/>
- [7] E. Ratti e F. Russo (2024) *Science and values: a two-way direction*. European Journal for Philosophy of Science 14:6 <https://doi.org/10.1007/s13194-024-00567-8>
- [8] *Automi e Persone, Introduzione all'etica dell'IA e della robotica*. A cura di F.Fossa, V.Schiaffonati, G.Tamburrini. Carocci editore 2021
- [9] S. Crafa e M. Rizzo (2019) *Tecnologie digitali e manipolazione del comportamento*. Mondo Digitale n. 86 http://mondodigitale.aicanet.net/2019-7/Articoli/MD86_04_Tecnologie_digitali_e_manipolazione_del_comportamento.pdf
- [10] A. Trocchi *Giocare o essere giocati?*. Capitolo del libro Internet mon amour. Ed. Circe. <https://ima.circex.org/storie/2-relazioni/h-gamificazione.html>

- [11] S. Vallor (2023) *The Danger Of Superhuman AI Is Not What You Think*. Noema. <https://www.noemamag.com/the-danger-of-superhuman-ai-is-not-what-you-think/>
- [12] S. Vallor (2024) *The AI Mirror: How to Reclaim Our Humanity in an Age of Machine Thinking*. Oxford University Press 2024. Lecture: https://media.ed.ac.uk/media/Shannon+Vallor+-+Inaugural+Lecture+23.10.24+/1_p9lrnfj5
- [13] Z. Lipton e J. Steinhardt (2019) *Troubling Trends in Machine Learning Scholarship: Some ML papers suffer from flaws that could mislead the public and stymie future research*. Queue 17(1) <https://queue.acm.org/detail.cfm?id=3328534>
- [14] J. Pantaleoni (2023) *The quickest revolution. An insider's guide to sweeping technological change, and its largest threats*. Mimesis International ed.
- [15] K. Brennan, A. Kak, and S. Myers West (2025) *Artificial Power: AI Now 2025 Landscape Report*. AI Now Institute, <https://ainowinstitute.org/2025-landscape>
- [16] H. Farrell (2024) *The Map is Eating the Territory: The Political Economy of AI*. <https://www.programmablemutter.com/p/the-political-economy-of-ai>
- [17] E. Nardelli (2022) *La rivoluzione informatica. Conoscenza, consapevolezza e potere nella società digitale.*, Themis ed.
- [18] D. Tafani (2024) *GDPR could protect us from the AI Act. That's why it's under attack*. Bollettino telematico di filosofia politica. <https://commentbfp.sp.unipi.it/gdpr-could-protect-us-from-the-ai-act-thats-why-its-under-attack/>
- [19] K. Martin (2022) *Ethics of Data and Analytics Concepts and Cases*. Auerbach Pub.
- [20] S. Crafa e A. Zangari (2019) *Il business della vendita dei dati, pratiche e conseguenze nell'etica sociale*. Mondo Digitale n. 86 http://mondodigitale.aicanet.net/2019-7/Articoli/MD86_05_Il_business_della_vendita_di_dati.pdf

- [21] M. Miceli, P. Tubaro, A. Casilli, et al. (2024) *Who Trains the Data for European Artificial Intelligence?* European Parliament; The Left, pp.1-40. <https://hal.science/hal-04662589v1>
- [22] B.Perrigo (2023) *OpenAI Used Kenyan Workers on Less Than \$2 Per Hour to Make ChatGPT Less Toxic.* TIME. <https://time.com/6247678/openai-chatgpt-kenya-workers/>
- [23] A.Casilli, T. LeBonniec, J. Posada (2025) *The human cost of Deepseek.* DiPLab Policy Memo 1(1). Inst. Polytechnique Paris <https://diplab.eu/diplab-releases-first-policy-memo-the-human-cost-of-deepseek/>
- [24] RaiPlay (2025) *Codice - La vita è digitale. Sfruttati per l' AI* 22/08/2025. <https://www.raipplay.it/video/2025/08/Sfruttati-per-l-AI---Codice-22082025-9d1db3ef-7c04-4897-90f4-97bac05ed46e.html>
- [25] S. Vallor (2016) *Technology and the Virtues: A Philosophical Guide to a Future Worth Wanting.* Oxford Academic
- [26] *The Digital Humanism Initiative.* <https://caiml.org/dighum/>
- [27] Crafa S. (2019) *Artificial Intelligence and Human Dialogue* Journal of Ethics and Legal Technologies, 1(1), 44-56. <https://jelt.padovauniversitypress.it/2019/1/3>