



Commenti al manifesto di Meta “The Future is for Everyone”

Agosto 15, 2026

Come gruppo AlgoRadar, sentiamo il dovere di commentare il recente manifesto di Mark Zuckerberg, “The Future is for Everyone: The Path to a Positive IA Future”, perché tocca direttamente molte delle tematiche che hanno portato alla nascita di AlgoRadar.

Amiamo il progresso autentico, ma temiamo un futuro distopico. Il nostro compito collettivo fondamentale è promuovere un monitoraggio continuo dei limiti tecnologici, degli annunci iperbolici e delle possibili conseguenze indesiderate.

Questa valutazione preliminare è organizzata in due sezioni: la prima parte è un'analisi generale, mentre la seconda affronta i principali punti del manifesto, esplorando aspetti più tecnici e specialistici. I temi sono, tuttavia, molto ampi, e ci riserviamo il diritto di tornarvi con future analisi approfondite.

CONSIDERAZIONI GENERALI

Sebbene il manifesto di Meta possa in larga misura essere letto come una mossa promozionale volta a tenere il passo con i leader dell'Intelligenza artificiale (IA) “di frontiera” — vale a dire Anthropic, OpenAI, Google DeepMind e diversi attori cinesi — mostra anche segni di quello che si potrebbe chiamare “transumanesimo di livello 1”.

Meta afferma: “Il futuro è per tutti. Le persone potranno utilizzare una superintelligenza oltre la capacità umana.” Stiamo certamente assistendo a una rivoluzione di portata storica, ma fin dall'inizio viene impostato un tono eccessivamente positivo e inequivocabilmente propagandistico. Bisogna allora chiedersi: è comparabile all'intelligenza umana (senza nemmeno parlare della “mente” umana, concetto più ampio) in modo tale che le due possano essere misurate sulla stessa scala, giustificando così l'uso della parola “oltre”?

La visione di un'IA personale, sempre più presente nella nostra vita quotidiana, viene presentata come il prossimo grande passo dell'umanità. Tuttavia, riteniamo utile fermarsi un momento e distinguere ciò che è davvero nuovo da ciò che già conosciamo.

Per anni abbiamo vissuto in un ecosistema digitale in cui gli algoritmi raccolgono dati, apprendono dalle nostre preferenze, costruiscono profili e influenzano ciò che vediamo, compriamo e scegliamo.

Non si tratta di una dinamica che riguarda solo Meta: è una trasformazione che ha progressivamente coinvolto l'intero mondo digitale.

L'idea degli "agenti intelligenti" era già oggetto di studio nel secolo scorso, ma l'infrastruttura computazionale e una base utenti di riferimento non erano ancora disponibili. La novità dell'IA agentica non è l'idea di agente intelligente in sé, bensì la combinazione di scala, infrastruttura e sostenibilità economica.

L'IA agentica consente di passare dal modello "chatbot" (comprendere il contesto e rispondere alle nostre richieste) al ricordare, pianificare ed eseguire azioni per nostro conto. Si tratta di una importante discontinuità tecnologica.

Ma dobbiamo comprendere pienamente che cosa comporti questo spostamento e quale grado di autonomia siamo davvero disposti ad affidare alle macchine. È qui che, per noi, inizia la questione etica.

La lettera parla di "empowerment" individuale — un obiettivo che condividiamo — ma essere più capaci grazie alla tecnologia non significa automaticamente essere più autonomi.

Un'IA che ci conosce sempre meglio può diventare uno straordinario strumento di supporto. Ma la stessa conoscenza può consentirle di anticipare i nostri bisogni, orientare le nostre scelte e, nel tempo, influenzare il nostro comportamento.

Il problema, quindi, non riguarda solo ciò che accade ai nostri dati. Riguarda anche il confine tra assistenza e influenza, tra delegare un compito e delegare una decisione, tra essere aiutati e diventare dipendenti da un sistema che conosce le nostre abitudini e le nostre vulnerabilità.

Per questo motivo, riteniamo che con l'IA agentica dovremmo porci domande molto concrete: cosa dovrebbe sapere un sistema su di noi? Cosa dovrebbe poter inferire? Cosa dovrebbe poter fare autonomamente? E dove dovrebbe fermarsi?

Per AlgoRadar, parlare di IA etica significa anche cercare di fornire risposte verificabili a queste domande. Non basta dichiarare che l'IA debba essere "human-centered": dobbiamo poter valutare se una tecnologia aumenta davvero la capacità di una persona di comprendere, scegliere e decidere, oppure se finisce per eroderla.

Non è un compito semplice, ma è un problema che va affrontato.

Non siamo né contro l'innovazione né contro l'IA agentica. Tuttavia, pensiamo che una trasformazione di questa portata non dovrebbe essere definita esclusivamente da chi sviluppa la tecnologia. Il progresso può ampliare enormemente ciò che siamo in grado di fare. Ma spetta a noi decidere che cosa vogliamo che faccia per noi e che cosa vogliamo continuare a fare da soli.

La diffusione di un'IA più pervasiva e autonoma solleva anche una questione più profonda: in che modo l'architettura linguistica dei modelli odierni plasma ciò che possono rappresentare, interpretare e, in definitiva, fare?

Nei sistemi basati su “large language models”, il linguaggio non è soltanto il mezzo attraverso cui l’utente interagisce con il sistema: le assunzioni linguistiche, sintattiche e culturali incorporate nel modello possono contribuire a modellare il modo in cui le informazioni vengono rappresentate, interpretate e trasformate in risposte e, sempre più, in azioni.

Questo solleva un’importante questione etica che va oltre il contenuto delle risposte prodotte da un sistema di IA: fino a che punto l’architettura attraverso cui un agente rappresenta il mondo influenza già la gamma delle possibili interpretazioni, scelte e azioni?

Riteniamo che questo tema meriti ulteriori approfondimenti. Non intendiamo affrontarlo in modo esaustivo qui, ma piuttosto come una delle domande che richiederanno un’analisi più profonda nei futuri lavori di AlgoRadar.

Condividiamo pienamente anche questa affermazione di Alberto Carrara (Presidente dell’International Institute of Bioethics) nella sua lettera aperta a Mark Zuckerberg: “una persona può diventare più informata senza diventare più saggia; più produttiva senza sentirsi realizzata; più connessa senza diventare capace di relazioni più profonde; più potente senza diventare più responsabile; più autonoma nell’esecuzione dei compiti e tuttavia sempre più dipendente dai sistemi tecnologici per formulare giudizi.”

Infine, ma non meno importante, dovremmo anche tenere conto di teorie più ampie che mettono in discussione i presupposti dominanti sull’IA. Per esempio, Federico Faggin sostiene: “se siamo macchine, una macchina può superarci; ma se siamo qualcosa di più di una macchina e le sue proprietà non coincidono con le nostre, che vanno oltre quelle di una macchina, possiamo trovare un modo per lavorare insieme”. Come noi, Faggin non sta criticando la superintelligenza, ma sta dicendo che non dovrebbe essere l’obiettivo finale. È più importante fare prima progressi sul lato umano, e questo è strettamente legato all’etica.

Se l’IA deve servire l’umanità, il dibattito non può essere lasciato soltanto a prestazioni, scalabilità e competizione di mercato. Deve includere autonomia, responsabilità e una riflessione seria su che cosa significhi davvero il progresso umano.

APPROFONDIMENTO

Meta: “Proponiamo una filosofia basata sull’emancipazione individuale come fonte di prosperità, sull’invenzione come scopo primario della superintelligenza e sull’equilibrio dei poteri come fondamento della sicurezza.”

È davvero una “filosofia”? Siamo sicuri che il termine “emancipazione individuale” sia una definizione neutrale? La categoria dell’“individuo”, in contrapposizione a quella di “persona”, rischia di recidere le dimensioni relazionali e sociali che sono cruciali per la costituzione dell’essere umano. In questo modo potrebbe renderci ciechi di fronte a dinamiche quali “nuove dipendenze, esclusioni, manipolazioni e disuguaglianze” [Magnifica Humanitas, n. 95].

Lo stesso termine “empowerment” affonda le sue radici nella nozione di potere — un concetto tutt’altro che problematico: “Più potere non implica necessariamente qualcosa di migliore. In questo senso restano attuali le parole di Romano Guardini: ‘L’uomo contemporaneo non è stato educato a usare bene il potere’.”

Meta: “le opportunità e le sfide saranno probabilmente maggiori di qualsiasi cosa abbiamo visto nel corso delle nostre vite”

In linea di principio, si tratta di un’affermazione difendibile, perché riconosce opportunità e sfide (anche se non le definisce come rischi) e invita a confrontarsi seriamente con l’intera questione.

Meta: “Mettere il potere nelle mani delle persone”

Questo evoca, non troppo sottilmente, “potere al popolo” (detto da un miliardario della Silicon Valley). È comunque una posizione che possiamo in larga misura condividere: la concentrazione di controllo e potere nelle mani di pochi è giustamente motivo di preoccupazione. Promuove inoltre l’iniziativa privata e la realizzazione personale, che verrebbero rafforzate dalla superintelligenza.

Meta: “Invenzione, non automazione”

L’evoluzione dell’IA da “conversazionale” ad “agentica” è in effetti già in corso nelle applicazioni reali (industriali e non); ciò che riteniamo più problematico è la componente “creativa” (inventiva), che — pur ammettendone la piena fattibilità tecnologica — implica attrito con gli aspetti cognitivi più distintamente umani.

Meta: “lasciare che le persone decidano ciò che conta nella propria vita... distribuirla ampiamente e dare a ogni persona la capacità di dirigerla”

In linea di principio, questo è un punto che possiamo condividere, poiché riafferma il valore personale.

Meta: “Potrai interagire con il tuo agente attraverso qualsiasi dispositivo, compresi i tuoi occhiali, per restare presente nel momento con le persone a cui tieni”.

Sei davvero “presente” in una stanza con “le persone a cui tieni” quando interagisci contemporaneamente con “il tuo agente” tramite un dispositivo indossabile?

Anche nei contesti sociali ordinari, l’IA indossabile può rendere le interazioni meno trasparenti: a differenza di un laptop o di uno smartphone, un paio di occhiali con IA può mediare, registrare o elaborare una conversazione senza rendere tale mediazione visibile agli altri.

Va detto che, quando sei con le persone a cui tieni, ti senti in un ambiente amichevole in cui ciò che dici liberamente è “protetto” e non verrà usato contro di te. Che cosa accade se ciò che dici viene catturato ed elaborato dall’agente IA incorporato negli occhiali del tuo amico, e successivamente condiviso con altri, persino senza che chi li indossa ne sia pienamente consapevole?

Qui emerge un grande problema di privacy e riservatezza, che rischia di sconvolgere completamente i rapporti di fiducia tra le persone.

Meta: “Tutti avranno un agente personale eccezionalmente capace che comprende te, i tuoi obiettivi e tutto ciò a cui tieni”

Pur sembrando entusiasmante, i temi fondamentali della privacy e dell’abitudine tecnologica non emergono.

Meta: “Tutti avranno straordinari strumenti di creazione per esprimere le proprie idee. Tutti presto avranno superpoteri di invenzione.”

Tuttavia, questi superpoteri resterebbero in larga misura esterni, tranne che per poche persone con capacità, cultura ed esperienza sufficienti a usarli per arricchirsi davvero interiormente.

Meta: “Tutti avranno strumenti potenti per creare nuove attività imprenditoriali”

Questo è in parte vero, ma costruire un’attività di successo dipende da molto più che da conoscenze e algoritmi. Richiede anche qualità autentiche come intuizione, pensiero laterale, disponibilità a correre rischi, capacità comunicative e interpersonali, oltre a... un po’ di fortuna.

Meta: “Tutti avranno un tutor e un coach personalizzati con un PhD in ogni disciplina”

Questi sistemi possono essere utili in alcuni contesti, ma non sono indispensabili e possono diventare dannosi se incoraggiano una dipendenza crescente dall’IA.

Ciò che viene descritto qui non è paragonabile alla dinamica docente-discente che è al centro della supervisione dottorale. Né un buon insegnante deve necessariamente avere “pazienza illimitata”.

Un membro del nostro team ricorda che l’impazienza del suo relatore per la bozza del capitolo successivo fu un forte stimolo intellettuale e contribuì in ultima analisi al completamento del suo dottorato.

Equiparare un agente IA a un tutor umano o a un insegnante è, di per sé, una distorsione anomala e pericolosa. Significa sostituire figure che — pur non essendo onniscienti — sono in grado di condividere la propria esperienza, con intelligenze diverse dall’intelligenza umana e prive di un fondamento cognitivo fisico diretto nel mondo reale.

Ciò incoraggia implicitamente l’idea errata che l’IA funzioni in modi paragonabili al cervello umano, cosa fuorviante sia dal punto di vista scientifico sia da quello filosofico.

Meta: “Tutti beneficeranno dei progressi scientifici e potranno contribuire al progresso della scienza”

Potrebbe esserci un ampliamento dei contributi, ma affermare che chiunque possa contribuire in modo significativo ai progressi scientifici semplicemente usando l’IA sembra eccessivo.

Meta: “Tutti avranno accesso libero o economico a questi strumenti”

Non bisogna farsi illusioni: l’accesso gratuito sarà disponibile solo in modalità ridotta e per catturare l’attenzione del grande pubblico; i sistemi più capaci resteranno probabilmente a pagamento, mantenendo così le disuguaglianze. Del resto, Meta e organizzazioni simili sono aziende, non enti di beneficenza.

Piuttosto, tutto ciò è più probabile che ponga le basi per un nuovo divario strutturato lungo fratture geografiche ed economiche.

La formulazione attentamente scelta — “Per chi desidera pagare per usare più capacità di calcolo, ci sarà un meccanismo di asta dinamica” — non maschera forse l’asimmetria di potere di mercato tra produttori e consumatori di questi servizi, dalla quale una gran parte dell’umanità potrebbe rimanere esclusa? Qui vi sono almeno due questioni chiave. In primo luogo, l’accesso a Internet rimane profondamente diseguale tra i Paesi: nei Paesi ad alto reddito l’uso di Internet si avvicina alla copertura universale, con il 94% della popolazione online, mentre nei Paesi a basso reddito solo il 23% della popolazione usa Internet. In secondo luogo, all’interno dei Paesi, il divario tra chi può permettersi l’accesso a servizi IA più potenti e costosi e chi è confinato a modelli gratuiti più semplici e meno potenti potrebbe diventare sempre più marcato.

Meta: “Allo stesso tempo, le nuove tecnologie portano anche nuovi rischi”

È positivo che ciò venga evidenziato, anche se solo alla fine dei numerosi benefici promessi nel manifesto. Inoltre, non viene presentata alcuna strategia per affrontare correttamente questo problema.

Meta: “La visione convenzionale è che... se prendiamo abbastanza tempo, allora possiamo perfezionare o allineare la tecnologia per produrre una singola superintelligenza benevola. Penso che questa visione dell’allineamento sia fundamentalmente errata. Invece, proponiamo che sia più produttivo guardare a ciascuna di queste preoccupazioni attraverso la lente del raggiungimento del giusto equilibrio di potere.”

Questo approccio convince solo in parte, ma è almeno relativamente originale.

L’IA sta evolvendo in una direzione in cui l’imposizione tecnologica di barriere interne diventa sempre più complessa e, in ogni caso, insufficiente. È quindi necessario agire anche sul contesto politico, economico, sociale ed etico globale in cui l’IA opera. Dobbiamo, ovviamente, confrontarci in modo pragmatico con scenari geopolitici e logiche di mercato, che nessuna istituzione o legislazione è in grado di governare efficacemente. Tuttavia, questo non significa affidarsi completamente a tali logiche, come sembra proporre Zuckerberg; al contrario, richiede un monitoraggio costante e una funzione pubblica indipendente di informazione.

Crediamo davvero che “persone e istituzioni” con interessi concorrenti si controllino e bilancino naturalmente in modi che portino a esiti positivi? È sorprendente che un altro attore cruciale in questo campo venga involontariamente omissso dall’elenco: le aziende. Più in profondità, questa affermazione non trascura forse persino la condizione più elementare affinché verifiche e contrappesi funzionino, ossia che gli attori coinvolti dispongano di forme di potere sufficientemente bilanciate perché tali controlli siano efficaci? Possiamo seriamente assumere che le grandi aziende hi-tech e i singoli governi nazionali — per non parlare dei singoli utenti — abbiano un potere comparabile?

Da qui nasce l'urgenza per l'Europa (e, naturalmente, per chiunque altro sia interessato a mantenere la possibilità di prendere decisioni basate sul pensiero critico) di sviluppare una propria infrastruttura digitale pubblica indipendente, senza affidarsi interamente alle aziende Big Tech.

Meta: “Ci sono diverse ragioni per essere ottimisti sul fatto che ci sarà un'abbondanza di posti di lavoro in futuro.”

Queste affermazioni non sono supportate da sufficienti evidenze empiriche. Questo è uno dei punti discutibili del manifesto. Proprio perché l'IA è una rivoluzione con un tasso di crescita senza precedenti e sviluppi persino imprevedibili, riteniamo che i rischi dal punto di vista occupazionale (sia in termini di numero di posti di lavoro sia di qualità del lavoro) siano elevati. Del resto, non ci si può basare soltanto sul contesto della Silicon Valley; il mondo va valutato in tutta la sua diversità e complessità. Basta, per esempio, conoscere la situazione in Cina, che sta diventando un leader tecnologico nell'IA.

Meta: “Costruire l'infrastruttura IA con le comunità”

Intenzioni eccellenti, almeno in apparenza. Le aziende cercano il sostegno delle comunità locali per consentire l'ulteriore espansione dei data center in cambio di vari benefici. Oltre al consumo di suolo, il peso ambientale ed energetico resta una questione strutturale nell'intero settore del calcolo su larga scala.

Inoltre, studi aggregati a livello nazionale negli Stati Uniti indicano che la crescita storica dei data center fino al periodo 2019–2023 era in realtà correlata a tariffe medie al dettaglio più basse, perché le vendite di elettricità ad alto volume aiutavano i servizi pubblici a distribuire i costi fissi di manutenzione della rete su una base più ampia. Tuttavia, i dati emergenti del 2025–2026 mostrano che questa dinamica di stabilizzazione si sta frammentando in aree ad alta concentrazione, dove i vincoli locali di capacità e gli onerosi aggiornamenti infrastrutturali stanno iniziando a far salire i prezzi all'ingrosso e quelli delle aste di capacità.

Meta: “Proteggersi dagli abusi dell'IA nella cybersecurity, nel bioterrorismo e altro ... idealmente dovremmo accelerare la scienza mitigando al contempo i rischi.”

La questione è estremamente importante ma viene trattata in modo molto debole e generico.

Meta: “La privacy è una base importante per l'emancipazione individuale e la libertà. Crediamo che le persone debbano avere una modalità completamente privata per gli agenti personali, nella quale nemmeno Meta o qualsiasi altro fornitore di servizi possa vedere o concedere accesso alle tue informazioni.”

Qui viene finalmente affrontato uno dei temi cruciali: la privacy. Si promette una modalità completamente privata per gli agenti personali; il problema è come venga garantita, chi la verifichi e come la superintelligenza possa operare con dati privati interagendo al tempo stesso con migliaia di altri sistemi e agenti.

Meta: “Garantire la leadership americana”

Non entriamo in questioni politiche riguardanti gli Stati Uniti o altre nazioni. Dal punto di vista dell'equilibrio dei poteri, promosso dallo stesso Zuckerberg, riteniamo utile che la conoscenza tecnologica sia distribuita in modo equo e che l'Europa acceleri i propri investimenti nel settore dell'IA.

Meta: “Allineamento con le persone e affrontare il rischio esistenziale”

L'allineamento è il processo volto a garantire che i sistemi di IA perseguano in modo affidabile obiettivi coerenti con i valori umani, le intenzioni e i principi etici. Il manifesto afferma l'importanza di questo aspetto ma non indica alcuna soluzione concreta. Qualcuno crede davvero che gli utenti siano sempre in grado di decidere quale bene/servizio consumare e in quale quantità?

Sappiamo che gli individui spesso lottano con problemi socialmente rilevanti come alcolismo, abuso di sostanze e obesità. Un fattore comune in molti di questi fenomeni è la dipendenza, cioè un disturbo della scelta in cui il soggetto diventa, poco alla volta, incapace di compiere la scelta migliore per sé. Se si aggiunge, come in questo caso, che diversi allineamenti saranno in competizione tra loro, siamo sicuri che agenti che, nell'esempio di Zuckerberg, “rifiutano di aiutare a redigere una lettera ai futuri genitori in una scuola” prevarranno su quelli meno “eticamente rigorosi”?

La teoria evolutiva dei giochi mostra che l'evoluzione di comportamenti “altruistici” rispetto a quelli “egoistici” all'interno di una popolazione può produrre risultati molto diversi e spesso non ottimali dal punto di vista sociale.

Meta: “Mantenere il controllo della superintelligenza”

Questo tema è trattato in modo confuso; in sostanza, ciò che si auspica è il bilanciamento del carico computazionale dell'IA tra gli obiettivi di auto-miglioramento tecnologico delle stesse aziende di IA e gli esiti utili per gli utenti finali.

Meta: “L'approccio open source è una forza positiva e importante per l'emancipazione delle persone”

C'è ambiguità nella terminologia usata, tra modelli “open source” e “open weight”; l'apertura dei modelli è certamente un fattore positivo, anche se oggi la tendenza è costruire applicazioni agentiche composite in cui i modelli sono solo una parte di un sistema di IA.

ALGORADAR

AlgoRadar è un gruppo indipendente di monitoraggio e ricerca nel settore dell'IA. Riunisce scienziati ed esperti provenienti dal mondo accademico e industriale, sviluppando nuovi paradigmi e quadri etici per l'intelligenza artificiale in sanità, istruzione, economia e giustizia.

www.linkedin.com/company/algoradar

E-MAIL: research.staff@algoradar.org

FONDATORI:

Mirella MASTRETTI	IA Scientist
Maria Pia ABBRACCHIO	Neuroscientist & Pharmacologist
Stefano ANTONIAZZI	Computing Paradigms and Technologies Scientist
Ariela BENIGNI	Biotechnology Scientist
Michela CINQUILLI	Canonic Lawyer
Elio FRANZINI	Philosopher
Raul GABRIEL	Artist, Writer & IA Editorialist
Mario A. MAGGIONI	Economist
Daniela MARI	Medical Scientist
Fabrizio PECORARO	Public Health and Social Research Scientist