



Intelligenza artificiale: scienziati e politici per una sicurezza globale

Descrizione

L'intelligenza artificiale (IA) è entrata ufficialmente nell'agenda diplomatica mondiale. Durante l'80ª sessione dell'Assemblea generale delle Nazioni unite è stato lanciato un [appello](#) globale per introdurre linee rosse sull'IA. La finalità è impedire che lo sviluppo tecnologico superi i confini della sicurezza, dell'etica e del controllo umano. Oltre duecento gli scienziati, premi Nobel e politici che hanno sottoscritto l'appello. Solo per ricordare alcuni tra i promotori: i premi Nobel per la fisica Giorgio Parisi e Geoffrey Hinton premio Turing; Joseph Stiglitz, John Hopfield e Daron Acemoglu, Nobel per l'economia; lo storico e filosofo Yuval Noah Harari; scienziati di OpenAI, Anthropic e Google; Yoshua Bengio e Joseph Sifakis, premi Turing.

Rischi inevitabili e controlli necessari

L'iniziativa, presentata nel corso della [High-Level Week](#), evidenzia da un lato l'enorme potenziale dell'IA per migliorare il benessere umano e rileva, altresì, che l'attuale traiettoria di sviluppo comporta pericoli senza precedenti. Che potrebbe presto superare di molto le capacità umane ed amplificare minacce come pandemie ingegnerizzate, disinformazione diffusa, manipolazione delle persone su larga scala inclusi i minori rischi per la sicurezza nazionale e internazionale, disoccupazione di massa e violazioni sistematiche di diritti umani.

I promotori sottolineano che alcuni sistemi di IA avanzata hanno già manifestato comportamenti ingannevoli e dannosi, eppure a questi sistemi viene concessa un'autonomia sempre maggiore nell'agire e nel prendere decisioni nel mondo reale. Ciò significa che molti esperti compresi coloro che sono in prima linea nel suo sviluppo avvertono che, se non adeguatamente governata, nei prossimi anni diventerà sempre più difficile garantire un controllo umano effettivo dell'IA.

Nell'appello si rivolge una esortazione ai governi affinché si raggiunga un accordo internazionale, intervenendo con decisione e garantendo l'effettiva applicazione attraverso solidi meccanismi di controllo e di attuazione. Necessario un accordo internazionale che stabilisca limiti chiari e verificabili per prevenire rischi universalmente inaccettabili. Senza escludere le imprese,

â??cosÃ– da assicurare che tutti i fornitori di IA avanzata siano soggetti a criteri condivisiâ?•.

Il tutto si fonda su un principio sostanziale: la *governance* deve essere umanocentrica. Ma la strada per tradurre lâ??appello in prescrizioni Ã– un percorso prevedibilmente complicato con diversi ostacoli. Tra questi, soprattutto le sovranitÃ– nazionali e gli interessi economici delle *big tech*. A cui si aggiunge il divario significativo tra rapiditÃ– delle innovazioni tecnologiche e tempi dei decisori politici (*paceing problem*). Altro aspetto, certo non secondario, Ã– che le *big tech*, pur sostenendo lâ??opportunitÃ– di limiti, avvertono che vincoli eccessivi potrebbero rallentare lâ??innovazione tecnologica.

Una governance flessibile

Una possibile prospettiva per dare risposte coerenti e fattibili a queste criticitÃ– potrebbe essere solo una *governance* flessibile, basata su criteri di rischio e valutazione indipendente in grado di coniugare sicurezza, competitivitÃ– e innovazioni.

Un appello, un progetto ambizioso che non puÃ² non avere come interlocutore la politica che, per dirlo con [Carlo Galli](#), â??contrapponga alla superficie la profonditÃ– e la trascendenza, al presente la storia, al progresso come destino lâ??avvenire come libera creazione, al formicaio la cittÃ–, allâ??oligarchia la democrazia, al virtuale il reale.â?•

In questa prospettiva il Nobel Giorgio Parisi e Pierluigi Contucci, docente dellâ??UniversitÃ– di Bologna Alma Mater, unitamente ad autorevoli studiosi, hanno lanciato il [Manifesto](#) per lâ??istituzione di un Centro europeo di ricerca sullâ??IA che, sul modello del Cern, consenta di lavorare fianco a fianco scienziati di discipline diverse ma complementari a un comune sforzo sia scientifico che umanistico. Dando prioritÃ– allo sviluppo *open-source*, affinchÃ© gli strumenti dellâ??IA siano accessibili a ricercatori, istituzioni, imprese e cittadini, rafforzando lâ??accesso democratico e promuovendo lâ??innovazione. Con una particolare attenzione ad integrare la riflessione etica e la valutazione dellâ??impatto anticipando le conseguenze sociali, economiche e culturali. Coinvolgendo decisori politici, societÃ– civile e opinione pubblica per allineare lâ??innovazione ai valori democratici, ai diritti umani e allâ??interesse collettivo.

Lâ??appello presentato allâ??Assemblea generale delle Nazioni unite ed il Manifesto di Giorgio Parisi evidenziano lâ??attualitÃ– di questioni che non possono non interrogarci nÃ© tantomeno essere sottovalutate. Tra queste, chi puÃ² governare e sono governabili i sistemi artificiali intelligenti? Intrecciandosi questioni tecniche, etiche e geopolitiche: chi stabilisce cosa Ã– accettabile? Chi decide quando un sistema di IA diventa una minaccia?

Un prima risposta Ã– stata data con [lâ??AI Act](#), il Regolamento del Parlamento europeo e del Consiglio che stabilisce regole armonizzate sullâ??IA. Tuttavia, lo sviluppo planetario dei sistemi artificiali intelligenti richiede un orizzonte ancor piÃ¹ ampio da prendere in considerazione. Sempre piÃ¹ necessario visti gli sviluppi di tecnologie che si alimentano con un ritmo incessante e veloce secondo una prospettiva di preponderanza dellâ??intelligenza statistico-computazionale rispetto a quella logico-simbolica dellâ??umano.

Democrazie o algocrazie?

Ã? in questione il rapporto tra democrazia e sistemi artificiali intelligenti ovvero lâ??oligopolio delle *big tech* che rappresentano dei veri e propri attori globali. Potenze computazionali che ampliano le

capacità umane ma che sono dotate di un enorme potere, non solo economico ma anche cognitivo e politico. Si definiscono i criteri con cui l'IA viene progettata, le priorità e i valori che possono orientare la vita collettiva. Gli algoritmi diventano nuove forme di autorità. Algorazie, per dirlo con p. Paolo Benanti, che sottraggono alla deliberazione pubblica questioni decisive per la democrazia in cui il sapere e il potere sono diffusi, plurali, sottoposti a controllo pubblico. Altrimenti, cittadinanze esposte a poteri opachi, capaci di influenzare comportamenti e opinioni.

Il rischio è la democrazia svuotata dall'interno e sostituita da processi automatizzati delegati a decidere al posto nostro con una concentrazione di potere che genera una dipendenza strutturale degli Stati. Senza escludere il ricorso a sistemi di IA da parte di democrazie illiberali e autocrazie populiste per sostenere i propri poteri.

Contrastare queste derive non significa fermare il progresso, ma restituire l'orizzonte politico del bene comune e investire sull'intelligenza sociale diffusa. Solo così l'IA potrà diventare strumento di libertà condivisa e non di dominio concentrato.

Crediti foto di [Markus Winkler](#) su [Unsplash](#)

Data di creazione

26 Ottobre 2025

Autore

lucio_romano