

Governare l'Intelligenza Artificiale: il ruolo delle politiche pubbliche e l'esempio dell'Europa

17 Gennaio 2024

Il dibattito sull'intelligenza artificiale è esploso a seguito del rilascio di ChatGPT e recentemente ha fatto notizia l'accordo raggiunto tra il Consiglio e il Parlamento europeo sull'AI Act, la prima normativa sull'intelligenza artificiale che ne regolerà l'utilizzo in Europa, imponendo obblighi, restrizioni e divieti a seconda di una classificazione in livelli di rischio. A discutere di IA non sono più solo informatici e ingegneri, ma anche filosofi, giuristi, economisti e sociologi perché l'IA è un tema trasversale che solleva questioni etiche, legali e sociali oltre che tecniche.

L'intelligenza artificiale ha enormi potenzialità di fare del bene: iniziative come la piattaforma "[AI for Good](#)" o il Think Thank "[AI4SDGs](#)" si impegnano nello studio e nella diffusione di applicazioni di IA per il benessere sociale. Tra gli esempi più immediati ci sono sicuramente le applicazioni in medicina che possono aiutare nella rilevazione precoce, per esempio, di patologie tumorali, accelerare lo sviluppo di nuovi farmaci, o trovare soluzioni che arrivano fino alla telemedicina. L'IA può essere utilizzata anche per analizzare i dati satellitari, al fine di monitorare fenomeni legati ai [cambiamenti climatici](#), identificare problemi di [malnutrizione](#) o pratiche di [pesca illegali](#). Ci sono inoltre numerose applicazioni in ambito educativo, agricolo, finanziario, di assistenza ai disabili, di ottimizzazione del traffico, gestione dei rifiuti o efficientamento degli edifici.

Tuttavia, l'IA porta con sé anche dei rischi che non possono essere ignorati: il riconoscimento di immagini potrebbe essere utilizzato ad esempio per identificare in modo molto rapido e preciso determinati gruppi religiosi in base a caratteristiche peculiari dell'abbigliamento.

Un'altra questione che preoccupa il dibattito pubblico è l'impatto di queste tecnologie sul lavoro: l'IA sostituisce alcune mansioni e ne crea di nuove o cambia il modo di svolgere vecchi lavori. Per questo è essenziale la formazione continua e giocano un ruolo centrale la scuola, gli ammortizzatori sociali e le politiche attive del lavoro. C'è inoltre un importante tema legato alle

diseguaglianze, poiché non è chiaro come questi benefici e costi verranno distribuiti.

I rischi legati all'IA sono spesso problematiche preesistenti, che vengono però aggravate dalla potenza delle nuove tecnologie: gli algoritmi replicano e amplificano ciò che trovano nei dati e da questo derivano innumerevoli questioni legate all'utilizzo dei dati, alla privacy, al diritto d'autore, alla disinformazione, alle distorsioni (*bias*) e alle discriminazioni. Il problema è che il punto di forza di queste tecnologie risiede proprio nella loro capacità di essere applicate su larga scala con rapidità ed efficienza. Perciò se i dati utilizzati sono distorti, contengono errori e pregiudizi umani, se i fini non sono etici o gli obiettivi non sono specificati correttamente, la potenza di calcolo degli algoritmi accrescerà esponenzialmente questi difetti. Diventa particolarmente pericoloso nel caso delle armi autonome, se alla scalabilità si aggiunge l'estrema precisione nel riconoscimento degli obiettivi.

Questi temi sono politici più che tecnici e spetta ai governi trovare soluzioni e promuovere il dialogo. Ne parla Stuart Russell^[1] nel suo libro "[Human Compatible: AI and the Problem of control](#)", in cui esplora le sfide etiche e di governance dell'IA. Russell solleva, tra le varie questioni, il problema del disallineamento dei valori, ovvero il fatto che le macchine massimizzano l'obiettivo specificato per loro, ma non comprendono né condividono i valori degli umani con cui interagiscono. Pertanto, un obiettivo mal definito può portare a conseguenze catastrofiche. Russell si rifà al mito di Re Mida che chiese a Dioniso il dono di trasformare in oro tutto ciò che toccasse, finché si accorse che anche il cibo diventava oro appena i suoi denti lo sfioravano e, distrutto dalla fame, dovette chiedere al dio di riprendersi il suo dono. Adattando la storia ai tempi moderni, l'autore immagina di chiedere a una macchina di contrastare la rapida acidificazione degli oceani dovuta all'aumento dei livelli di anidride carbonica. La macchina sviluppa un nuovo catalizzatore che ripristina i livelli di pH degli oceani. Purtroppo, un quarto dell'ossigeno presente nell'atmosfera viene consumato nel processo lasciando asfissiare l'umanità.

A far sorgere ulteriori complicazioni è il comportamento umano: una macchina che non ne conosce le preferenze tenderà a desumerle dalle azioni osservate, ma le scienze comportamentali mostrano che le scelte umane spesso deviano dalla razionalità per fattori psicologici e contestuali.

Per affrontare queste sfide è necessario, perciò, un approccio multidisciplinare che metta in dialogo l'informatica, la statistica e la robotica con l'etica, il diritto e le scienze sociali fin dalla fase di progettazione di tali sistemi e durante l'intero il ciclo di vita.

Giocano pertanto un ruolo centrale le politiche pubbliche, la cui sfida è stare al passo con le tecnologie che avanzano molto più rapidamente rispetto ai tempi della legislazione. Nel 2019 L'OCSE ha pubblicato i [principi sull'IA](#) e nel 2021 ha adottato la [raccomandazione per una governance regolamentare dell'innovazione](#) in cui si è sottolineata la necessità di adottare norme flessibili in grado di adattarsi all'evoluzione tecnologica e di integrare strumenti di valutazione di impatto della regolamentazione. Diversi Paesi hanno adottato strategie nazionali per l'intelligenza artificiale. L'osservatorio OCSE per le politiche sull'IA ([OECD AI Policy Observatory](#)) fornisce un quadro completo delle politiche nazionali in materia di intelligenza artificiale e sottolinea una crescente tendenza nei governi a riconoscere il proprio ruolo in quest'ambito.

L'Unione Europea si è mossa in modo deciso su questa tematica e il 9 dicembre 2023 Parlamento e Consiglio hanno raggiunto l'accordo su un progetto di legge che introduce obblighi e restrizioni basati su diverse categorie di rischio. Saranno vietati in particolare gli utilizzi classificati come rischi inaccettabili per la sicurezza e per i diritti fondamentali delle persone, tra cui il riconoscimento biometrico in luoghi pubblici, i sistemi di social scoring e di polizia predittiva o il riconoscimento delle emozioni sul posto di lavoro e negli istituti scolastici – salvo ragioni mediche o di sicurezza

(ad esempio, il monitoraggio dei livelli di stanchezza di un pilota). I sistemi classificati ad alto rischio comporteranno invece specifici obblighi di conformità, quali la valutazione preliminare dell'impatto, sistemi di gestione del rischio e della qualità, obblighi di supervisione umana, di governance dei dati e corretta informazione degli utenti. Per i sistemi ritenuti meno pericolosi la normativa prevede solo requisiti di trasparenza e invita le aziende a impegnarsi volontariamente in codici di condotta.

Il testo dell'AI Act, che entrerà in vigore entro la fine del 2024, uniformerà le modalità di regolamentazione dell'IA dei 27 Stati membri portando con sé importanti implicazioni extraterritoriali, perché la regolamentazione si applicherà a tutti i sistemi di IA che possono avere un impatto sulle persone nei Paesi UE, ma a prescindere dal luogo di sviluppo o distribuzione di questi sistemi.

L'AI Act segue altre importanti normative digitali dell'UE, come il regolamento generale sulla protezione dei dati (GDPR), la legge sui servizi digitali (DSA) e sui mercati digitali (DMA) la legge sui dati e la legge sulla cyber-resilienza. La normativa europea rivela la tendenza ad adottare un approccio di cautela verso le nuove tecnologie, privilegiando la mitigazione dei rischi rispetto alla corsa all'innovazione, che caratterizza invece le strategie per l'IA di Stati Uniti e Cina. Questo approccio può comportare però un ritardo tecnologico, mentre i potenziali effetti negativi delle tecnologie sviluppate altrove potrebbero comunque propagarsi attraverso fenomeni di spillover internazionale. . Attualmente esiste infatti un vuoto nella governance globale dell'IA colmato solo in parte dalle iniziative dell'OCSE, come la definizione dei Principi sull'IA e l'Osservatorio sull'IA che cercano di promuovere informazione, dialogo e collaborazione internazionale.

Mentre l'IA continua a stupire con la sua capacità di trasformare interi settori e semplificare molti compiti, solleva anche questioni complesse di natura etica, legale e sociale. Il ruolo della regolamentazione e delle politiche pubbliche diventa sempre più cruciale per bilanciare innovazione e rischi e per orientare il potenziale dell'IA in una direzione che sia socialmente vantaggiosa ed eticamente sostenibile. L'AI Act è un primo passo importante, ma resta da valutare quanto sarà efficace nel raggiungere questo delicato equilibrio.

[1] *Stuart Jonathan Russell è professore di informatica all'Università della California, Berkeley e noto per i suoi contributi nel campo dell'intelligenza artificiale*

LINK

<https://magazine.fbk.eu/it/news/governare-lia-il-ruolo-delle-politiche-pubbliche-e-lesempio-delleuropa/>

TAG

- #artificial intelligence
- #Europa
- #Intelligenza artificiale
- #politiche pubbliche

AUTORI

- Greta Sofia Lampis