



## Diritto e innovazione

# I.A.: il “complesso di Frankenstein” e l’urgenza di regolare.

*Intelligenza artificiale. Quali regole* di G. Finocchiaro

di [Antonio Scalera](#)

18 maggio 2024

---

Il 13 marzo scorso, il Parlamento Europeo ha approvato il regolamento sull’intelligenza artificiale (c.d. “AI ACT”), frutto dell'accordo raggiunto con gli Stati membri nel dicembre 2023<sup>[1]</sup>.

Il Regolamento entrerà in vigore venti giorni dopo la pubblicazione nella Gazzetta ufficiale dell'UE e inizierà ad applicarsi 24 mesi dopo l'entrata in vigore, salvo alcune eccezioni.

Il Regolamento contiene, all’art. 3, la seguente definizione di “sistema di IA”: *“un sistema automatizzato progettato per funzionare con livelli di autonomia variabili e che può presentare adattabilità dopo la diffusione e che, per obiettivi espliciti o impliciti, deduce dall’input che riceve come generare output quali previsioni, contenuti, raccomandazioni o decisioni che possono influenzare ambienti fisici o virtuali”*.

Questo atto normativo del legislatore europeo costituisce una prima risposta, a livello unionale, a quell’ “urgenza di regolare” il fenomeno dell’intelligenza artificiale, di cui parla nella sua recente pubblicazione, *“Intelligenza artificiale. Quali regole?”* Giusella Finocchiaro, ordinaria di diritto privato e di internet all’Università di Bologna, una dei maggiori esperti italiani della materia.

“Urgenza di regolare” che – osserva l’autrice – nasce dalla paura che la tecnologia ci possa sopraffare e dominare.

In poche e lucide pagine, rivolte non solo ad un pubblico di giuristi, la Finocchiaro affronta essenzialmente tre temi.

Anzitutto, la paura che l’intelligenza artificiale suscita (“*il complesso di Frankenstein*”[\[2\]](#)) e che, come si è detto, porta a invocare nuove regole, per assicurare e placare l’ansia.

Poi, il linguaggio utilizzato, che costruisce sapientemente una nuova narrazione e un nuovo mito, denso di parole seducenti, ma ingannevoli: si parla di “*intelligenza*”, appunto; ma anche di “*oracolo*”, per riferirsi all’attività dei sistemi di I.A. di fornire risposte alle domande; e ancora di “*allucinazioni*”, ovvero gli esiti errati delle elaborazioni di ChatGPT; di “*incantesimi*”, per indicare le inserzioni nei comandi impartiti ai sistemi di I.A., i c.d. prompts, inserzioni che consentono di aggirare le limitazioni e i divieti imposti a ChatGPT e alle applicazioni analoghe e di condizionarne il comportamento.

Infine, la retorica che avvolge l’intelligenza artificiale e che, certamente, influenza anche il ragionamento giuridico. Comprendere il fenomeno è il primo indispensabile passo per poi decidere come procedere, per valutare quali siano le regole applicabili e come applicarle e se debbano esserne create di nuove.

Accanto alla definizione di I.A., di natura descrittiva ed estremamente generale, contenuta nel regolamento, ve ne sono molte altre, tra le quali è ancora attuale quella elaborata da A. Turing nel 1950, secondo cui l’intelligenza artificiale può essere definita come “*la scienza di far fare ai computer cose che richiedono intelligenza, quando vengono fatte dagli esseri umani*”[\[3\]](#).

La ricerca di una definizione porta a interrogarsi sulla soggettività dell’intelligenza artificiale.

La domanda, indubbiamente suggestiva – “*l’intelligenza artificiale è cosciente?*” – prelude necessariamente alla costruzione di un modello di responsabilità

L’autrice è dell’opinione che, *de iure condendo*, si debba formulare un modello giuridico adeguato a dare risposte sulla responsabilità per i danni cagionati dai sistemi di I.A.

A questo riguardo, vi è chi propone, sulla scia di quanto indicato dalla Risoluzione del Parlamento Europeo del 16.2.2017, l’istituzione di uno *status* giuridico specifico per i robot, di modo che almeno i robot autonomi più sofisticati possano essere considerati come persone elettroniche responsabili di risarcire qualsiasi danno da loro causato.

Questa soluzione è, a ben vedere, solo apparente e cede al fascino retorico della soggettività delle applicazioni di I.A.[\[4\]](#)

La soggettività del sistema di I.A. – osserva la Finocchiaro – sembra essere una costruzione che aumenta la complessità giuridica, piuttosto che diminuirla. Il nodo della questione non è tanto l’attribuzione della soggettività giuridica all’I.A. quanto l’individuazione dei criteri di allocazione della responsabilità.

Secondo alcuni, un principio che potrebbe rivelarsi di grande utilità è quello dell’*accountability* [5], cioè la responsabilizzazione del soggetto che trae vantaggio dall’applicazione di I.A.

Utilizzando questo criterio, il soggetto che trae maggior vantaggio risponde adottando esso stesso le misure necessarie ad evitare il rischio.

Sul punto, l’AI ACT non è di aiuto, in quanto si limita a prevedere che il fornitore di un sistema di I.A. ad alto rischio[6] è chiamato a garantire che lo stesso sia conforme a determinati requisiti che ne consentano la certificazione.

Il libro si conclude con alcune considerazioni critiche sul regolamento europeo di recente approvazione.

In primo luogo, l’approccio regolatorio è di tipo orizzontale: si norma in generale, l’I.A. e non, invece, applicazioni specifiche della stessa.

Inoltre, vi è il rischio che il sistema regolatorio non sia sufficientemente dinamico da adattarsi ai successivi sviluppi dell’I.A. Il sistema tracciato appare rigido. La classificazione dei sistemi di I.A. nell’ambito delle diverse tipologie di rischio sarà soggetta a revisione. Nuovi sistemi verranno sviluppati e nuovi metodi per implementare i sistemi già esistenti verranno creati, modificando il livello di rischio.

Ulteriore criticità è quella di prevedere, a fronte di una generalissima definizione di I.A., in capo a ogni tipo di imprenditore i medesimi obblighi, a prescindere dalle dimensioni dell’impresa. Un modello di questo genere ha all’origine il difetto di trattare tutte le applicazioni di I.A. allo stesso modo, senza considerare che esse possono essere assai diverse tra loro e declinate in maniera differente.

Vi è, poi, il rischio che l’approccio adottato dal regolatore europeo comporti una burocratizzazione del mercato europeo dell’I.A. senza che a ciò corrisponda una maggiore tutela dei diritti.

Infine, il monito dell’autrice.

La normativa europea non deve avere l’effetto di isolare l’Europa.

Piuttosto, è importante che questi temi siano affrontati in un’ottica di cooperazione internazionale, anche normativa, tanto più in questa fase in cui Stati Uniti<sup>[7]</sup> e Cina<sup>[8]</sup> sembrano orientarsi su principi non lontani da quelli europei: “*Se vogliamo costruire delle fortezze, dobbiamo poi ricordarci di costruire anche i ponti che ci consentano di collegarle ad altri sistemi*”.

---

<sup>[1]</sup> Il Regolamento, che stabilisce regole armonizzate sull’intelligenza artificiale e modifica i regolamenti (CE) n. 300/2008, (UE) n. 167/2013, (UE) n. 168/2013, (UE) 2018/858, (UE) 2018/1139 e (UE) 2019/2144 e le direttive 2014/90/UE, (UE) 2016/797 e (UE) 2020/1828, è stato approvato con 523 voti favorevoli, 46 contrari e 49 astensioni.

<sup>[2]</sup> Si veda M. Ciardi, *Mary Shelley, Isaac Asimov e il complesso di Frankenstein*, in *Anime Positroniche*, a cura di M. Ciardi e P. Gaspa, Quaderni Cicap, 2021.

<sup>[3]</sup> Nel famoso saggio *Computing Machinery and Intelligence*, A. Turing ideò un test al fine di determinare se una macchina sia in grado di pensare. Il test si basa sulla valutazione delle capacità di un computer di imitare il comportamento umano; in caso di esito positivo si vede ritenere che la macchina sia in grado di pensare in modo equivalente o, comunque indistinguibile da un essere umano.

<sup>[4]</sup> Il termine “intelligenza” è condizionante e induce a pensare a un essere intelligente. In questo caso si utilizza una metafora: l’applicazione di intelligenza artificiale si comporta “come se” fosse intelligente. Ma occorre essere consapevoli dei vantaggi così come dei limiti nell’utilizzo delle metafore, in modo da evitare il sopravvento della metafora sulla realtà. Sull’utilizzo delle metafore nel linguaggio giuridico si rinvia a F. Galgano, *Le insidie del linguaggio giuridico*, il Mulino, 2010.

<sup>[5]</sup> Il principio di *accountability*, introdotto dal Regolamento sui dati personali (GDPR), prevede che sia il titolare del trattamento dei dati personali a determinare le misure più adeguate al al trattamento stesso. In questo caso il legislatore affida al titolare l’onere di individuare in che

modo adempiere alle prescrizioni dettate dalla norma, calandole nella fattispecie concreta, assumendosi la responsabilità non solo dell’implementazione, ma anche della valutazione.

**[6]** È noto che il modello adottato dal legislatore europeo è basato sul rischio e differenzia gli usi dell’intelligenza artificiale secondo il livello di rischio, che può essere inaccettabile, alto o minimo.

**[7]** Con l’Executive Order emanato il 30.10.2023 dal Presidente Biden sullo sviluppo e sull’uso dell’I.A., gli Stati Uniti intendono non solo rimanere all’avanguardia nel campo dell’I.A., ma anche garantire che questa tecnologia sia utilizzata in modo responsabile, a beneficio del pubblico e della sicurezza nazionale; al primo punto, infatti, si prevede che gli sviluppatori dei sistemi di I.A. condividano i risultati dei loro test ed altre informazioni critiche con il governo degli Stati Uniti.

**[8]** La strategia cinese ha preso forma nel 2017, con la *Cybersecurity Law*, e ha condotto all’emanazione dei *Beijing AI Principles*, pubblicati nel 2019 dalla Beijing Academy of Artificial Intelligence e dei *Principles to Develop Responsible AI for the New Generation Artificial Intelligence*, pubblicati nel 2019 dal New Generation AI Governance Expert Committee. Dal 15.8.2023 sono in vigore anche i nuovi principi sull’I.A.: *Interim measures for the Administration of Generative Intelligence Services*.