



Diritto e innovazione

Quali impatti avranno su di noi i *Large Language Models* e CHAT GPT (e quali saranno le conseguenze per il mondo del diritto)

di [Ercole Aprile](#)

27 febbraio 2023

1. Introduzione

I Large Language Model, e su tutti il più conosciuto tra essi, Chat GPT, hanno cambiato in brevissimo tempo il modo in cui concepiamo la tecnologia.

In un mondo in cui la quantità di informazioni e di dati disponibili è in costante crescita, i Large Language Model propongono al pubblico un nuovo modo per accedere e utilizzare queste informazioni: dialogare in linguaggio naturale con un sistema che fornisca l'illusione di una relazione con un'entità senziente, in grado di rispondere alle nostre domande e di elaborare testi, più o meno complessi, nella forma che più ci aggrada.

Ai Large Language Model, infatti, possiamo chiedere di scrivere un'e-mail su un argomento specifico, di produrre un articolo con un determinato numero di battute su un tema a nostra scelta, possiamo anche indicargli di scrivere una poesia imitando lo stile dell'autore che più ci aggrada.

In tutti questi casi, il risultato potrà essere più o meno completo e dettagliato, ma nella gran parte delle occasioni fornirà la sensazione di trovarsi davvero di fronte a un prodotto dell'ingegno umano.

Al di là delle comprensibili tentazioni di rendere complessa ed eccessivamente tecnica la definizione, il funzionamento e i possibili campi di fruizione di tali modelli, tentazioni che possono riguardare sia gli operatori del diritto che quelli delle tecnologie, per comprendere appieno quanto dirompente possa essere stata la diffusione nell'utilizzo di questi strumenti, non solo nel definire il nostro rapporto con l'universo digitale, ma anche nell'esperienza quotidiana, è utile, se non necessario, riflettere su come l'avvento della diverse rivoluzioni digitali che si sono susseguite nell'ultimo quarto di secolo, a partire dalla prima, legata al World Wide Web, abbia modificato in maniera radicale il nostro rapporto con l'idea stessa di conoscenza.

2. Internet e la trasformazione della conoscenza

La diffusione dell'accesso a Internet, infatti, ha rappresentato una svolta epocale nella composizione e nella struttura delle fonti del sapere, una svolta tale da portare il filosofo Silvano Tagliagambe a definirla *una seconda rivoluzione copernicana*[\[1\]](#).

Il passaggio è stato rapido e repentino.

Se per lungo tempo, a partire dal *Fedone* di Platone, passando dal *Novum Organum* di Francis Bacon, fino al contemporaneo Albert Einstein, per rappresentare la conoscenza umana e le sue molteplici ramificazioni ci si è riferiti alla metafora dell'*Albero*, le tecnologie digitali hanno reso tale metafora non più applicabile su larga scala.

Nell'*Albero*, le radici formano le basi della conoscenza, il fusto simboleggia la sua struttura centrale e le diverse ramificazioni disegnano le molteplici branche del sapere.

Con Internet, tutto è cambiato, si è passati da una struttura gerarchica a una struttura distribuita e interconnessa, rappresentata proprio dalla nuova metafora della Rete.

Nella Rete, infatti, non esiste una base o un centro e i nodi sono tutti ugualmente importanti nella trasmissione del sapere. Ancora più importanti sono le interconnessioni tra i nodi, perché più la Rete è fitta, più è efficiente[\[2\]](#).

Se al tempo della prima diffusione di questa idea la natura democratica che sottendeva la metafora era apparsa come il preludio a una libertà totale nella costruzione del sapere,

svincolata da oligopoli e concentrazioni, questa promessa appare, oggi, in buona parte disattesa.

Non solo.

Se da un lato la disponibilità planetaria ed orizzontale delle fonti ci ha garantito l'accesso, in potenza, a un sapere sconfinato, dall'altro, per le modalità con cui si è realizzata, ha influito su molti aspetti del nostro rapporto con l'assimilazione dei contenuti, prospettando finanche una modifica dei processi che poniamo alla base della ricerca e della memorizzazione.

3. L'universo digitale ci sta cambiando...e probabilmente ci cambierà ancora

Oggi, grazie a ricerche e studi che hanno esaminato la natura di questa metamorfosi^[3], appare sempre più plausibile che l'evidente cambiamento nella proposizione e fruizione del sapere non riguardi solo un approccio funzionale e di natura pratica, ma investa anche la sfera fisiologica e le connessioni sinaptiche che definiscono i percorsi attraversati dai dati che elaboriamo nei nostri processi mentali.

Il nostro cervello, infatti, per adattarsi alla realtà che lo circonda, modifica i collegamenti neuronali e il modo in cui reagiscono agli stimoli esterni, generando i nuovi percorsi cerebrali che questi stimoli possono percorrere.

L'utilizzo del digitale, in forma consistente e costante, può amplificare queste modulazioni, sia per la sua natura di nuova realtà con cui socialmente ci troviamo a interagire, sia per le modalità con cui questa fruizione si realizza: rapida, breve, frequente.

Quel che emerge dai citati studi è un impatto, in generale, sulle *capacità metacognitive*, ovvero su quei processi mentali che ci permettono di organizzare e controllare le attività cognitive e di pensiero. Appare plausibile che le capacità metacognitive stiano parzialmente cedendo il passo alla componente più istintuale e affettiva, legata, nel panorama di contesto, ai processi di piacere e di ricompensa. Questo ribilanciamento sarebbe estremamente rilevante dal momento che i processi di razionalizzazione del pensiero contribuiscono in maniera essenziale all'acquisizione e al consolidamento di ciò che impariamo e lo fanno utilizzando l'attesa come strumento fondamentale, soprattutto nel caso di processi di lunga durata. In pratica, quando dobbiamo conservare memoria di elementi, ragionamenti, riflessioni che richiedono un impegno sia in termini di applicazione che di tempo necessario all'apprendimento, le pause di attesa

predispongono alcune condizioni necessarie affinché i processi di assimilazione si svolgano correttamente e con successo.

La frenesia, la scarsa profondità e la bulimia informativa che caratterizzano molte delle nostre interazioni digitali quotidiane confliggono con questa architettura di processi e generano condizioni diverse alle quali le nostre menti provano ad adattarsi, portando come conseguenza una progressiva riduzione delle capacità di attenzione e concentrazione.

Appare evidente, oggi, una diffusa e pervasiva incapacità di affrontare l'attesa, sia essa quella relativa al caricamento di una pagina Web, sia quella per una coda inaspettata nello svolgimento di una commissione, sia, ancora, un allungamento del tempo di accesso per la fruizione di un contenuto.

In tutti questi casi, e in altri a essi assimilabili, è percepibile e palpabile un aumento diffuso del nervosismo, del disagio e del fastidio.

Se l'osservazione empirica ci permette di valutare gli impatti sociali di questi fenomeni, quel che ai più, probabilmente, sfugge è la conseguenza che ne deriva: se non siamo capaci di aspettare, non riusciamo a imparare.

E sempre nell'osservazione empirica della quotidianità, nei dialoghi e nelle interazioni relazionali emerge come, di pari passo con il progredire nella frequenza e rilevanza degli accessi al digitale, si diffonde la scarsa capacità di scatenare il meccanismo di ricompensa, il più efficace tra i *driver* dei nostri processi mentali, nel caso di elaborazioni che richiedano una quantità di tempo più consistente; a questo corrisponde una costante ricerca di ricompense a breve o a brevissimo termine, un esempio su tutti lo scrolling compulsivo di notizie o post sui social network, non in grado di fornirci il grado di benessere e di piacere di cui abbiamo bisogno.

L'osservazione correlata evidenzia come, se in passato la possibilità di utilizzare informazioni in maniera efficace era strettamente legata alla possibilità e alla capacità del singolo di riuscire a memorizzarle, oggi ci si sia socialmente affidati a una memoria "di accesso" e non di "di contenuto": non avendo più la necessità di memorizzare le informazioni in sé, conserviamo memoria dei percorsi da utilizzare per ricercarle, e ciò accade, nella maggior parte dei casi, attraverso un rapido utilizzo, e una conseguente rapida consultazione, degli strumenti digitali a nostra disposizione.

4. La delega ai sistemi digitali e la “consapevolezza” come strumento per affrontare il futuro

Possiamo condividere, quindi, la visione di una progressiva delega al contesto tecnologico di talune parti che hanno da sempre caratterizzato la natura del processo di acquisizione del sapere, delega che include il processo di ricerca delle fonti, spesso demandato alla elaborazione algoritmica dei motori di ricerca, e della memoria.

Possiamo persino spingerci a paventare una ridotta capacità di imparare che esuli dalle motivazioni sociali e si spinga a includere impatti fisiologici.

E proprio per l'esperienza che abbiamo accumulato negli ultimi cinque lustri, sebbene brevissima dal momento che in termini di evoluzione sociale venticinque anni possono essere considerati poco più di un battito di ciglia, è fondamentale che l'avvento dei nuovi Large Language Model sia accolto, sì, con entusiasmo e fiducia, ma anche con profonda consapevolezza.

È necessario comprenderne le caratteristiche, interrogarsi sui contesti in cui possano costituire valido supporto e, soprattutto, riflettere sui rischi di diversa natura che un loro utilizzo poco informato possa comportare. È necessario, oltre che auspicabile, inoltre, che una riflessione di questo tipo coinvolga anche le componenti legislativa e giudiziaria, dal momento che, mai come oggi, è urgente non rimanere indietro nella valutazione di impatto su entrambi i versanti, considerata la velocità estrema con cui le innovazioni digitali si impongono nella società contemporanea.

Se è vero, come è vero, che la consapevolezza nasce dalla conoscenza, dobbiamo, come primo passo, interrogarci su cosa siano davvero i Large Language Model e come quello che tra essi a oggi appare il più utilizzato ed efficiente, Chat GPT, abbia origine e funzionamento.

5. Cosa sono i LLM e come funzionano

Un Large Language Model è un tipo di modello di intelligenza artificiale addestrato per elaborare e generare testo in modo simile a come lo farebbe un essere umano. Semplificando, un Large Language Model è un sistema in grado di elaborare il linguaggio naturale e di produrre testo coerente e significativo.

Per addestrare un Large Language Model, si utilizzano enormi quantità di dati (articoli di giornale, libri, pagine Web, messaggi di social media, ecc.) in modo che, attraverso la loro analisi e una programmazione parallela, possa apprendere e progressivamente affinare le regole del linguaggio, la grammatica, le regole di costruzione delle frasi, le parole che si usano più frequentemente.

Per avere un'idea di cosa si intenda con “enormi quantità di dati”, basti pensare che, durante la fase di addestramento, al modello CHAT GPT sono stati dati in pasto una versione filtrata dell'intero Web, raccolta tra il 2011 e il 2021, e collezioni di libri in formato digitale composte da milioni di volumi differenti sugli argomenti più disparati.

Una volta addestrato, il modello può produrre nuovo testo in base a un input iniziale.

Se si pone una domanda al modello, esso può generare una risposta basata sulla elaborazione dei termini che compongono la domanda e del contesto di riferimento, il che rende gli Large Language Model utili in una vasta gamma di applicazioni: generazione di testo creativo, assistenza all'elaborazione del linguaggio naturale, traduzione automatica, sintesi del testo e molto altro ancora.

Quel che è importante sottolineare, però, è che questo processo di elaborazione non include una comprensione semantica del testo.

L'enorme capacità di calcolo di cui questi sistemi dispongono permette loro di eseguire una valutazione statistica sulla sequenza di parole e frasi legate ai termini presenti nella domanda. Nessun ragionamento, quindi, nessuna valutazione di merito, solo una ricerca, e in continua evoluzione, della probabilità che la parola successiva inserita nel testo sia la più appropriata e significativa rispetto alle parole che la precedevano.

6. Problemi effettivi e potenziali legati all'utilizzo dei LLM

Questa circostanza porta con sé due ordini di problemi.

Il primo è di natura fattuale. L'elaborazione statistica delle parole non garantisce alcuna attendibilità che vada oltre la costruzione di uno scritto verosimile nella forma. I meccanismi di risposta si basano sulle occorrenze, sui lemmi, sulle informazioni di cui parlavamo, e queste informazioni comprendono di tutto, elementi validi ed elementi che non lo sono. I Large Language Model sono sprovvisti di un qualsivoglia vaglio critico che permetta di approcciare in

maniera razionale o semantica il testo progressivamente costruito e i comprensibili tentativi operati dai ricercatori per porre un filtro agli argomenti che appaiono più scomodi o rischiosi sono stati, fino a oggi, aggirati o elusi, a volte finanche dagli stessi Large Language Model.

Va da sé che un qualsiasi ausilio alla produzione di testo da parte di questi sistemi va ipotizzato tenendo sempre e comunque presente la necessità di una attenta e approfondita analisi ex post, e questo vale in particolar modo per i contesti in cui la parola scritta ha consistenti impatti sul piano pratico, andando a influenzare o intervenire su aspetti rilevanti della vita dei singoli.

Il secondo è di natura elaborativo\sociale e riguarda l'accesso alle informazioni e il riferimento alle fonti e, anche in questo caso, può essere utile ricordare quanto già accaduto in passato per dotarsi di una cornice d'esperienza in cui inquadrare gli eventi correnti e le prospettive future.

Come scrivevamo in precedenza, la promessa della prima era Internet di fornire un'informazione libera da vincoli e condizionamenti è stata parzialmente disattesa. Se è vero, da un lato, che in Rete è possibile trovare ogni tipo di notizia, opinione, orientamento, è altrettanto vero che, oggi, gli utenti dedicano, mediamente, poco tempo e ancor meno attenzione alla valutazione di quanto gli viene proposto. Nella maggior parte dei casi, si affida il proprio quesito a un motore di ricerca o a un'applicazione social lasciando che siano i sistemi algoritmici a indicarci le fonti che è più utile proporci.

L'utilità cui ci riferiamo è duplice perché comprende, sì, l'interesse dell'utente alla esplorazione del risultato, dal momento che questo è uno dei passi fondamentali per la sua fidelizzazione, ma, sull'altro versante, persegue l'interesse economico del motore stesso che non chiede alcun pagamento monetario agli utilizzatori e utilizza la pubblicità mirata come fonte di guadagno.

Questo sistema si è trasformato, progressivamente, nella concretizzazione pratica, come scrivevamo, di una delega agli algoritmi da parte degli utenti, delega avente ad oggetto la scelta delle fonti attraverso cui informarsi, con le conseguenti distorsioni che un affidamento totale della dieta informativa a un soggetto terzo, non super partes, può portare con sé.

Eppure, con i Large Language Model, nell'attuale forma di interazione dialogica, si fa un passo verso una delega ancora più forte.

A oggi, nelle risposte previste da Chat GPT, ad esempio, non è presente alcuna indicazione sull'origine dei dati cui la produzione della risposta o del testo fa riferimento; la limitata possibilità di scelta della fonte offerta dai motori di ricerca è completamente sparita.

Utilizzando i Large Language Model non solo ci affidiamo a loro nella scelta sul tipo di informazioni da reperire e sul dove recuperarle, gli stiamo affidando anche la loro analisi, elaborazione e presentazione.

7. Conoscere per comprendere: predisporre ai futuri impatti nel mondo del diritto

Le conseguenze e i temi che questo cambiamento epocale può lasciar presagire sono molteplici e riguardano sia il versante teorico che l'applicazione pratica del diritto.

Sul piano civile, ad esempio, si pongono questioni riguardanti la responsabilità civile dei proprietari dei modelli nel caso di eventuali danni causati dal loro utilizzo da parte di terzi, come potrebbe accadere, ad esempio, nelle circostanze tutelate dal diritto d'autore.

Sul piano penale, invece, i Large Language Model possono essere utilizzati per commettere crimini, come la diffusione di contenuti illegali o pericolosi, il furto di dati sensibili o la frode. In questi casi, il legislatore e i giudici sono chiamati a considerare la questione della responsabilità penale dei proprietari dei modelli di linguaggio, nonché quella degli utenti che li utilizzano per commettere crimini.

Un aspetto particolarmente delicato, inoltre, è quello che riguarda l'impatto sulla privacy degli utilizzatori o di soggetti terzi.

Abbiamo scritto di come l'uso dei Large Language Model richieda la disponibilità di grandi quantità di informazioni per l'addestramento del modello e ciò comprende anche la raccolta e l'elaborazione di dati personali. Questi, possono essere raccolti tramite diversi canali, ad esempio social media, siti Web, e-mail e chat.

Appare chiaro, a tal proposito, che sia quanto meno da sollevare la questione relativa al consenso sulla loro raccolta ed elaborazione in un contesto che potrebbe andare ben oltre i termini di utilizzo sottoscritti dagli utenti delle diverse piattaforme di riferimento. Si consideri, ad esempio, la profilazione automatizzata, che si traduce nell'elaborazione delle informazioni personali per ottenere informazioni sul comportamento, le preferenze o le caratteristiche dell'utente.

Oggi, è difficile immaginare tutti gli impatti e le conseguenze sostanziali che queste nuovissime tecnologie potranno avere nel futuro prossimo e in quello più remoto.

Quel che, però, sappiamo di certo è che andranno ben oltre questi scenari, esemplificativi e non certo esaustivi; abbiamo, tuttavia, più di due decenni di esperienza alle spalle su come innovazioni digitali di questa portata abbiano avuto un impatto rilevante sul nostro modo di vivere la quotidianità e, di conseguenza, sul nostro sviluppo personale e sociale.

Partendo proprio da questa esperienza, sarà fondamentale utilizzare il, probabilmente breve, tempo che ci separa dal momento in cui l'adozione dei sistemi di intelligenza Artificiale di questo tipo sarà diffusa su larga scala, perché gli operatori del diritto imparino a conoscerli quanto più a fondo possibile. In questo modo potranno predisporre, se non a prevedere tutti gli impatti sul piano giuridico e normativo, quantomeno a normare, sul versante legislativo, e a trattare, sul versante giudiziario, scenari che includano la presenza di queste e altre modalità di utilizzo delle intelligenze artificiali. Avremo la possibilità, così, di far tesoro di quanto osservava Randy Pausch, celebre informatico e professore all'University of Virginia, che ricordava come *l'esperienza è ciò che otteniamo quando non otteniamo ciò che vogliamo*.

[1] Tagliagambe S., 1998, *Rete, paradigma della conoscenza*, Repubblica.it, online, <https://www.repubblica.it/online/internet/mediamente/tagliagambe/tagliagambe.html>, 14 febbraio 2023.

[2] AA.VV., 1997, *The Father of the WEB*, Wired.com, online, <https://www.wired.com/1997/03/ff-father/>, 14 febbraio 2023.

[3] Leonardi G., 2021, *L'attesa nella neuropsicologia: uno sguardo ai processi cognitivi sottostanti nell'era delle nuove tecnologie* in *DNA – Di Nulla Academia Rivista di studi camporesiani*, vol. 2, n.1, pagg 249-260.