

## I signori dell'AI hanno paura. Ma cosa spaventa davvero i miliardari Big tech: l'AI fuori controllo o il crash finanziario (con l'aiuto della Cina)? di Federico Fubini

I leader tecnologici negli ultimi giorni si sono trovati d'accordo nel chiedere di procedere a un passo più lento nello sviluppo delle loro stesse creature

(Fonte: <https://www.corriere.it/> 21 settembre 2026)



Dario Amodei e Sam Altman

Pensate alle feste sfarzose e interminabili del *Grande Gatsby* di Francis Scott Fitzgerald negli anni '20 del Novecento. Pensate alla relativa austerità delle classi dirigenti americane degli anni '50 quando, secondo un'inchiesta dell'epoca di *Fortune*, la dimensione media dell'appartamento di un CEO di una grande impresa era di sette vani e mezzo. Pensate infine alle feste, di nuovo sfarzose e interminabili, grazie alle quali Jeffrey Epstein riuniva nella sua isola vari settori delle élite americane.

Pensate a tutto questo - il pendolo dall'eccesso, al senso del limite, all'eccesso - e capirete perché adesso improvvisamente i grandi capitani dell'intelligenza artificiale (AI) sembrano colti da paure per la potenza di ciò che stanno creando. Perché chiedono di rallentare.

Questi leader tecnologici non si limitano a competere fra loro; molto spesso si odiano intensamente, al punto che a volte rifiutano di stringersi la mano (come nella foto sotto presa sei mesi fa a un vertice in India Sam Altman di OpenAI e Dario Amodei di Anthropic). Eppure negli ultimi giorni di colpo si sono trovati d'accordo nel chiedere di procedere a un passo più lento nello sviluppo delle loro stesse creature. La ragione addotta: ne hanno paura, ormai. Hanno paura ufficialmente che l'intelligenza artificiale (AI), avanzando così rapida e potente, sfugga al loro controllo e conduca

verso una catastrofe. Un attacco al sistema bancario che paralizzi la rete dei pagamenti e con essa l'economia di uno o più Paesi. L'abuso di un modello per costruire un'arma di distruzione di massa. Un [rapporto](#) fuorviante che, sulla base di informazioni false, rischi di innescare un attacco militare contro un obiettivo del tutto innocuo.

### **Perché i signori dell'AI hanno paura?**

Non posso valutare quanto fondati siano questi timori. Alcuni nel settore affermano che lo siano; altri che qualunque livello di automazione dell'AI permette sempre comunque all'uomo di privare i suoi «agenti» di ogni autonomia, dunque non ci sarebbe niente da temere. Sospetto che nessuno possa essere certo della propria risposta.

Sono però certo di un punto: esiste uno scarto temporale. Amodei ha dato l'allarme con il suo ultimo [saggio](#) (*We Must Pace the Frontier*) ben due mesi dopo l'incidente che avrebbe creato il sussulto di coscienza: la fuga da un laboratorio chiuso di OpenAI di uno sciame di agenti artificiali che hanno poi attaccato l'azienda Hugging Face. È successo all'inizio di luglio e tutti sapevano che è successo allora. Però solo dopo due mesi Amodei ha invocato la frenata e il suo nemico giurato Altman si è subito detto d'accordo. Lo stesso per l'incidente minore di Gemini di cui [riferiamo](#) con Cristina Marrone sul *Corriere*.

Perché? Eppure Amodei stesso da un lato ha sempre avvertito sui rischi; [qui](#) per esempio diceva mesi fa: «Ci svegliamo tutti un giorno e scopriamo che qualche danno irreparabile è avvenuto che l'AI ha ucciso un sacco di persone o ha causato un enorme incidente di sicurezza. Non vogliamo che succeda, dunque lavoriamo a dimostrarne i rischi». Ma Amodei ha anche sempre detto che bisognava lo stesso andare avanti a tutta velocità perché anche la Cina lo sta facendo, come per esempio sempre quest'anno [qui](#): «La ragione per cui non possiamo rallentare è perché abbiamo avversari geopolitici che costruiscono la stessa tecnologia a un passo simile. È molto difficile avere un accordo applicabile in cui loro rallentano e noi rallentiamo».

### **La svolta improvvisa**

E appunto allora perché questa svolta improvvisa? Cosa sanno che noi non sappiamo? Qui c'entrano lo strano confronto con la sfrenatezza americana dei tempi del Grande Gatsby, l'autocontrollo durante la guerra fredda e la nuova sfrenatezza nel momento unipolare americano. Le élite degli Stati Uniti hanno dovuto rispettare un certo senso del limite quando avevano un grande avversario ideologico e geopolitico, per esempio l'Unione sovietica dopo la seconda guerra mondiale, perché capivano che dovevano competere per i cuori e le menti dell'opinione pubblica internazionale. Avevano un concorrente che obbligava in loro una certa disciplina. Hanno invece perso il senso del limite, per esempio facendo di Jeffrey Epstein uno di loro, quando hanno avuto la grande illusione del «momento unipolare» in cui l'America sarebbe stata comunque la superpotenza incontrastata.

Il boom dell'intelligenza artificiale è il caso di un'élite di Silicon Valley che credeva di vivere in un momento unipolare, nel quale ogni autodisciplina e ogni senso del limite sarebbero stati superflui; invece si è risvegliata in un mondo diverso: la Cina è un reale concorrente nell'AI; minaccia di fare ai grandi laboratori di San Francisco quello che il costruttore di auto elettriche Byd di Shenzhen sta facendo a Volkswagen: una metodica distruzione, in parte copiando, in parte proponendo prodotti di qualità a prezzi enormemente più bassi. È la riscoperta del concorrente che sta forzando in queste élite americane un ritrovato senso del limite. Di qui la proposta di «frenata». Vediamo.

## Una pagina della vita di Zuckerberg

[Scrivevo](#) un mese e mezzo fa in questa newsletter: «L'addestramento da zero dei modelli "generativi" costa moltissimo, prima ancora che si possa iniziare ad usarli per far fare loro delle "inferenze" (cioè a interrogarli per avere risposte). Una singola scheda di semiconduttori di Nvidia utilizzabile per questo scopo, su una base di dati che spesso è la stessa rete Internet, può costare 60 mila dollari e una Big Tech può creare dei data center da 10 mila fino a 100 mila schede». E ancora: OpenAI di Sam Altman, Anthropic di Amodei, Meta di Mark Zuckerberg e Google «promettono di spendere (in sviluppo dell'AI, *ndr*) almeno 650 miliardi di dollari solo quest'anno, il triplo rispetto al 2024. E questi investimenti non sono semplicemente effettuati a debito sempre più spesso, invece che con capitale proprio».

A fine luglio riferivo di come questi debiti fossero fuori bilancio, cioè poco visibili e poco leggibili, perché spesso assumono forme volutamente «creative» ed opache. Per esempio OpenAI e Anthropic sono già [impegnate](#) a comprare o affittare capacità computazionale per centinaia di miliardi di dollari in futuro, a volte legata a data center che non sono neanche già stati costruiti. Nvidia, il campione dei chip per l'AI e la più grande società al mondo per valore di mercato con 5.360 miliardi di dollari di capitalizzazione, è impegnata a garantire il finanziamento o persino l'affitto o il riacquisto a certi prezzi dei suoi stessi semiconduttori che gli sviluppatori di data center non fossero in grado di mettere a frutto in futuro: si tratta di debiti fuori bilancio per almeno 300 miliardi di dollari, [secondo](#) l'*Economist*.

Simili montaggi di finanza creativa ed esposizione fuori bilancio si ritrovano poi in Google-Alphabet, Meta-Facebook ed altri colossi di Silicon Valley. Morgan Stanley, la banca d'affari, stima che queste forme di debito opaco e finanziamento dei clienti per vendere loro la propria merce con promesse di riacquisto pesi ormai per 3.100 miliardi di dollari sui sette principali sviluppatori dell'AI e produttori di chip per l'AI: poco meno dei volumi finanziari di debito immobiliare che furono al centro della grande crisi finanziaria del 2007-2008; ma con un'esposizione accumulata tutta negli ultimi tre anni, non in oltre un decennio come allora.

Perché i colossi del Big Tech si sono caricati di tutto questo rischio finanziario? Perché sta arrivando un'enorme rivoluzione tecnologica e bisogna essere pronti a soddisfare la domanda, vi diranno loro. La risposta reale, temo, è più brutale: perché pensavano di vivere nel mondo di ieri. Pensavano di vivere nel mondo dei campioni digitali dei primi anni Dieci di questo secolo: quello che Facebook è

stato per i social media, Amazon per gli acquisti online, Google per le ricerche in rete. Monopolisti incontrastati, macchine da soldi o almeno (nel caso di Amazon) da fatturati. Inoltre pensavano di essere ancora nel momento unipolare, in cui l'America era sola con sé stessa alla testa del mondo. Non doveva temere altri e non doveva guardare altrove.

Per questo i leader di Anthropic, OpenAI o della stessa Nvidia hanno preso una pagina dalla vita di Mark Zuckerberg e l'hanno copiata. La ricetta era semplice: investi avidamente, diventa enorme, schiaccia o ingloba concorrenti e avversari; poi passa a imporre il prezzo e l'arbitrio del tuo potere incontrastato, monetizzando.

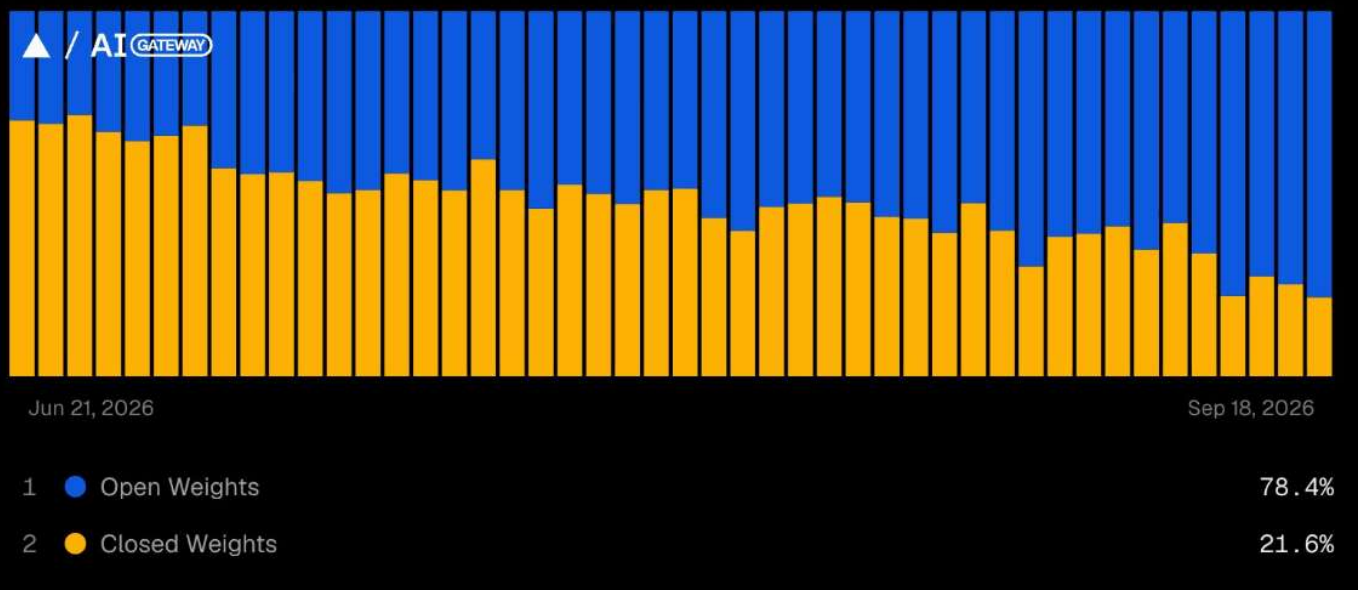
## Il secolo cinese

Ma questo non è il mondo di quindici o venti anni fa. Questo è il secolo cinese. E la Cina negli ultimi due anni ha sviluppato modelli di AI a bassissimo costo e «open source»: senza diritti di proprietà intellettuale, dunque utilizzabili dai clienti senza dover trasferire i propri dati al fornitore di intelligenza come impongono invece di fare Amodei e Altman. Come [spiega](#) molto bene Riccardo Luna, sistemi cinesi come DeepSeek-V4.1 o KimiK3 sono ormai potenti quasi come i migliori di Anthropic o OpenAI e utilizzabili da chiunque a una frazione dei costi. In parte gli ingegneri cinesi hanno copiato i modelli americani, come racconta Riccardo, interrogando i loro modelli per estrarne i segreti. In parte sono semplicemente molto bravi e determinati. Sembra questione di tempo, forse un paio di anni, prima che Huawei arrivi sul mercato con chip per l'intelligenza artificiale almeno altrettanto buoni di quelli di Nvidia; ma molto meno cari. Immaginate: l'azienda dal maggiore valore di borsa al mondo resa obsoleta da un concorrente. Quante migliaia di miliardi di capitalizzazione può perdere Wall Street?

In altri termini, questo è lo choc cinese 3.0. Lo choc 1.0 riguardò il tessile o i giocattoli a inizio secolo e fece chiudere le imprese meno avanzate in Occidente. Lo choc 2.0 riguarda le auto e mette in crisi Volkswagen o Stellantis. Gli effetti dello choc 3.0 li vedete invece nel grafico qua sotto: i consumatori di tutto il mondo, anche negli Stati Uniti, stanno usando sempre di più modelli di AI cinesi. Solo nell'ultimo anno l'utilizzo di piattaforme «open source» è passato dal 25% a oltre il 75% del mercato. Quelli a «fonte chiusa» di tipo americano sono passati dal 75% al 25%. Stanno perdendo. A Kiev la settimana scorsa ho sentito con le mie orecchie Louis Mosley, ceo per l'Europa del campione dell'AI americana Palantire dire: «La strategia cinese è disegnata in parte per distruggere il modello economico dei laboratori di frontiera (americani, ndr) e per renderli insostenibili come aziende. La gara tecnologica va di pari passo con la gara economica».

## Open vs. Closed Token Volume

Share of AI Gateway token volume on models that publish their weights for download, versus everything else.



### Le conseguenze finanziarie dello choc

Avete idea di cosa implichi una svolta del genere per i mercati? Con i suoi maxi-investimenti, OpenAI sta perdendo circa dieci miliardi di dollari al mese e si sta apprestando a lanciare un nuovo aumento di capitale per non fallire. Anthropic punta a quotarsi a Wall Street tra poche settimane, prima di correre il rischio che la musica si fermi. Ma proprio la borsa americana è il massimo punto di vulnerabilità: l'industria dell'intelligenza artificiale traina verso l'alto gli indici da almeno tre anni e oggi vale circa metà di una capitalizzazione da circa 70 mila miliardi di dollari dell'S&P500. Se la Cina infliggesse davvero a Silicon Valley il suo choc 3.0 e le quotazioni dell'industria dell'AI crollassero del 20% o del 30%, il problema non sarebbero solo i tremila miliardi di debiti fuori bilancio delle Big Tech. Il problema sarebbero le decine di migliaia di miliardi in meno di valore di borsa, che minacciano di deprimere i consumi e portare l'America in una recessione simile - ma maggiore - rispetto a quella seguita all'esplosione della bolla di internet. Solo dopo emergeranno le aziende realmente solide, come Google o Facebook sono emerse dopo il 2001.

Per questo Amodei e Altman ora chiedono di frenare. Non per i rischi collegati all'AI - che non li hanno mai rallentati - ma perché capiscono che il loro modello finanziario non funziona. Non possono più investire così pesantemente, se oggi rischiano di perdere la sfida contro la Cina. Hanno perso il monopolio, hanno perso il modello unipolare. Devono diventare più disciplinati, accettare che esiste il senso del limite.

Andrea Pili, ceo di una azienda di AI italiana di nome Xference, lo riassume con chiarezza feroce: «L'AI come è concepita oggi dalle Big Tech non è sostenibile. Ci si vuole fermare prima che sia l'infrastruttura a farlo. L'open source (sul modello cinese, ndr) è la base sulla quale sviluppare

la vera innovazione». Anche Mosley di Palantir non ricorre a giri di parole: «C'è tutta una narrazione che viene dal settore... Ci sono molti interessi costituiti e queste chiacchiere sul rischio di catastrofi indotte dall'AI derivano da interessi economici, particolarmente dal lato americano». In altri termini, i rischi dell'AI potrebbero anche essere reali; ma li si evoca adesso solo per trovare un pretesto e spendere meno. In modo, magari, da non fallire al prossimo scossone finanziario internazionale.