

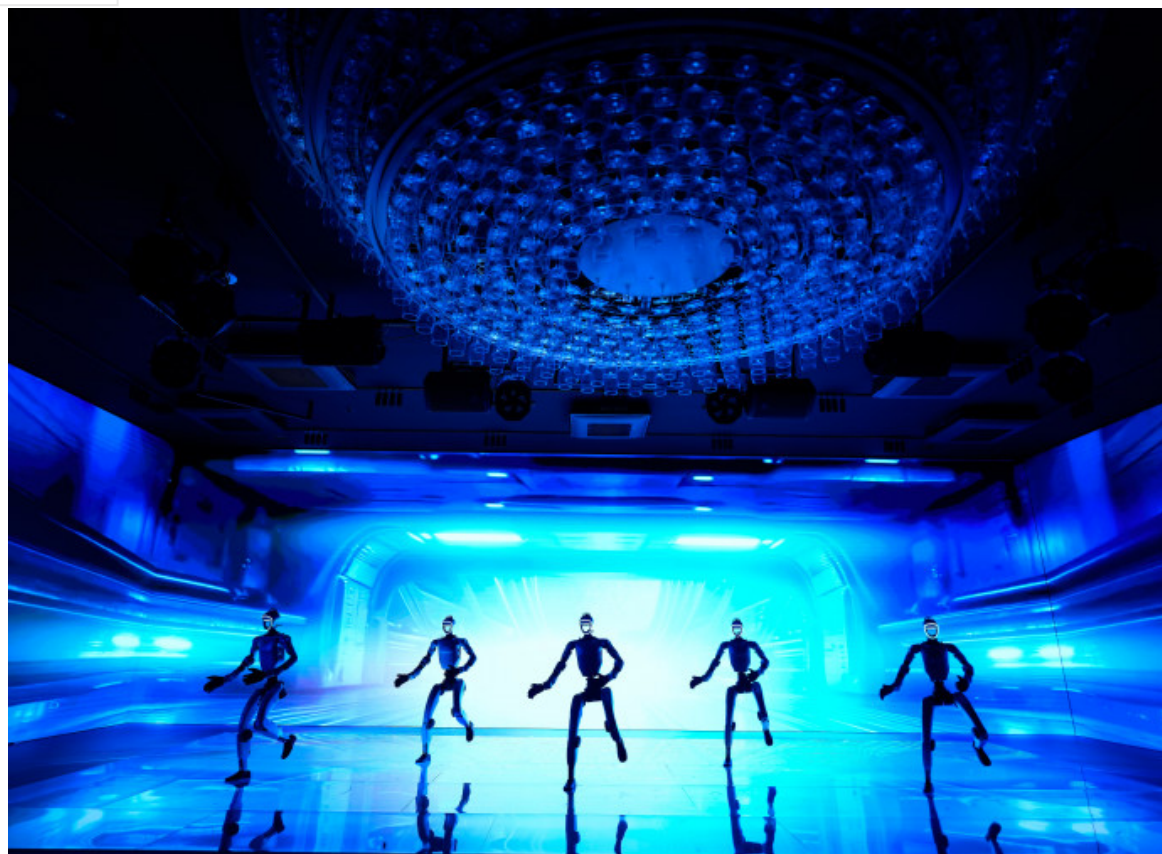
LA NUOVA PAURA

Perché, di colpo, tutti hanno paura dell'intelligenza artificiale

ECONOMIA

16_09_2026

**Daniele
Ciacchi**



Come abbiamo già avuto modo di trattare qui sulla *Nuova Bussola Quotidiana*, Jacob Coxon, ingegnere appena dimessosi da Anthropic, ha scritto a chiare lettere su X che chiunque sviluppi intelligenza artificiale abbia in cuore il dubbio che la stessa potrebbe

sterminarci "entro la fine del decennio". Evan Hubinger, anch'egli di Anthropic, ha confermato una probabilità superiore al 10%. Da lì è, giustamente, partito il tour mediatico, con titoli da fine del mondo ovunque, una nuova Guerra dei Mondi di Wells ma con un nemico nostrano, allarme estinzione imminente, e ancora l'umanità a essere carnefice di se stessa (gli algoritmi non si sono fatti da soli). Il libro di Eliezer Yudkowsky e Nate Soares, *If Anyone Builds It, Everyone Dies*, uscito a settembre 2025, è entrato nella classifica dei libri più letti del *New York Times* con una tesi abbastanza lineare sulla corrispondenza tra la creazione di una superintelligenza e la morte dell'umanità.

Le voci non sono però chiacchiere da bar, ma arrivano da esponenti del settore, e questo non può non pesare. Fa però storcere il naso un'ondata così compatta e tempestiva. L'IA pone sì rischi reali, ma forse diversi da come li intendiamo, e dietro questo catastrofismo si intravedono almeno tre interessi concreti e un'ideologia d'insieme a farne da cornice.

Già nel maggio 2023, davanti al Senato americano, Sam Altman chiese esplicitamente una nuova agenzia federale con poteri di licenza sull'IA. Il senatore Peter Welch e il testimone Gary Marcus avvertirono subito del rischio di quello che venne chiamato *regulatory capture*, che consiste nello stabilire soglie di calcolo e requisiti di sicurezza tarate su aziende con miliardi di dollari, escludendo di fatto concorrenti più piccoli e produttori *open source*, consolidando così il duopolio Microsoft / OpenAI. **David Sacks**, ex responsabile IA della Casa Bianca, ha dichiarato senza timore di smentita che Anthropic starebbe conducendo una tecnica strategica di *regulatory capture* facendo leva sull'allarmismo apocalittico.

Esiste poi il sempre classico mantra del "purché se ne parli" tanto caro al mondo della pubblicità. Meredith Whittaker, presidente di Signal, ha detto che le «narrazioni sul rischio esistenziale sono di fatto pubblicità per una tecnologia che solo poche aziende oggi possiedono». Il giornalista **Brian Merchant** osserva che dichiarare un prodotto così potente da poter finire il mondo significa dichiararne una potenza spropositata: è un'altra forma di marketing.

L'Alpocalypse sposta il fulcro dell'attenzione su un pericolo probabile e non imminente dal dibattito pubblico sui veri problemi che oggi l'IA dovrebbe mettere sotto la lente d'ingrandimento del dibattito pubblico. Mentre si discute di estinzione entro il 2030, restano ai margini temi di rilevanza sociale come i *bias* algoritmici, la sorveglianza pervasiva dei dati personali, lo sfruttamento del lavoro creativo dietro l'etichettatura dei modelli, il consumo idrico ed energetico dei *data center*. Sono danni misurabili, presenti, non futuribili ipotesi da romanzo fantascientifico. Concentrare l'attenzione collettiva su

uno scenario apocalittico e indefinito ha il pregio, per chi lo alimenta, di rimandare a un futuro remoto i conti che dovremmo farci nelle tasche oggi.

C'è anche una ragione geopolitica dietro alla corsa alla paura. Dario Amodei, alla Cbs, ha detto che alla base della corsa verso l'Intelligenza Artificiale Generativa sta la paura verso **il parallelo sviluppo cinese nello stesso ambito**. Trump è stato tranchant nell'affermare che chi vince la corsa all'IA vince tutto. Un rapporto dell'Fbi dell'8 settembre accusa i laboratori cinesi di aver "distillato" miliardi di token dai modelli americani per colmare il distacco, ridotto oggi al solo 2,7% secondo Stanford. Chi chiede di rallentare chiede, allo stesso momento, di fermare la Cina. La paura dell'estinzione dovrebbe viaggiare sui binari della corsa agli armamenti e dell'IA per uso militare.

Più che parlare di un vero e proprio complotto, bisogna qui parlare di interessi convergenti. Non serve un disegno coordinato quando tutti mangiano allo stesso tavolo. C'è però un'ideologia collettiva a fare da collante, quella dell'Altruismo Efficace, che oggi permea i laboratori e le *no-profit* che dettano l'agenda della sicurezza IA. Il meccanismo è molto semplice: se l'estinzione è un esito possibile, qualunque misura politica monopolistica diventa moderata se l'obiettivo è scongiurarla. Chiedere una licenza federale non è più prevaricazione del mercato ma prudenza. Lanciare l'allarme globale è un dovere morale, non una strategia di marketing. Ignorare i danni che l'intelligenza artificiale crea oggi non è negligenza, ma è rimettere le giuste priorità rispetto al rischio massimo della morte del nostro pianeta. In pratica, una cerchia ristretta rivendica una visione corretta e privilegiata del futuro - una specie di nuovo soviet - mentre i poster, in nome dei quali si parla, non possono contestare nulla.