

I tre uomini che decidono per tutti

Come poche persone senza mandato democratico orientano la traiettoria dell'intelligenza artificiale per miliardi di persone.



CONTROPOTERE

un libro di Alessio Farinella



I tre uomini che decidono per tutti

Sam Altman, Dario Amodei ed Elon Musk: come poche persone senza alcun mandato democratico stanno orientando la traiettoria dell'intelligenza artificiale per miliardi di persone

Alessio Farinella

Indice

Prefazione	2
Capitolo 1 – Chi sono, cosa controllano davvero ...	3
Capitolo 2 – OpenAI: da no-profit a colosso commerciale	6
Capitolo 3 – Anthropic: nato da una scissione, ma nella stessa corsa	8
Capitolo 4 – xAI e il conflitto di interessi di Musk ..	11
Capitolo 5 – Governance aziendale, non democratica	14
Capitolo 6 – Il lobbying	17
Capitolo 7 – Le porte girevoli	19
Capitolo 8 – Quando i team di sicurezza se ne vanno	22
Capitolo 9 – Il potere del compute	25

Capitolo 10 — Cosa NON è mai stato dimostrato . .	27
Capitolo 11 — Le obiezioni, trattate con onestà	31
Capitolo 12 — Cosa si potrebbe cambiare	34
Capitolo 13 — Cosa puoi fare tu, adesso	36
Nota sulla trasparenza nell'uso dell'intelligenza artificiale	38

Prefazione

Nessuno di loro ha mai fatto una campagna elettorale. Nessuno di loro compare su una scheda che qualcuno, da qualche parte, ha potuto votare o respingere. Eppure le decisioni prese in tre consigli di amministrazione — quanto velocemente sviluppare un sistema di intelligenza artificiale, quali rischi accettare, quali regole chiedere ai governi di scrivere o non scrivere — arrivano a influenzare la vita di miliardi di persone più in fretta di qualsiasi legge parlamentare.

Questo libro non è un attacco personale a Sam Altman, Dario Amodei o Elon Musk. È un tentativo di descrivere, con fatti verificabili, uno squilibrio strutturale: il potere di decidere la traiettoria della tecnologia più trasformativa del nostro tempo si è concentrato in pochissime mani, dentro strutture societarie che rispondono a consigli di amministrazione e investitori, non a elettori. È lo stesso meccanismo che questa collana racconta in decine di altri ambiti — dalle banche centrali ai forum privati che orientano l'agenda economica mondiale — applicato qui alla tecnologia

che, secondo [[Libro_Ultima_Invenzione]] di questa stessa serie, potrebbe rappresentare la trasformazione più grande mai vissuta dall'umanità.

Un avvertimento sul metodo, importante quanto il contenuto: attorno a queste tre persone circolano online moltissime affermazioni false o mai dimostrate. Questo libro le tratta esplicitamente in un capitolo dedicato, perché la forza di una critica reale sta nel non avere bisogno di inventare nulla — i fatti verificabili, da soli, bastano a raccontare uno squilibrio di potere che nessuna esagerazione renderebbe più serio di quanto già sia.

Capitolo 1 — Chi sono, cosa controllano davvero

Tre aziende, tre storie diverse

Sam Altman guida OpenAI, l'azienda che ha reso l'intelligenza artificiale generativa un fenomeno di massa con il lancio di ChatGPT nel 2022. Dario Amodei guida Anthropic, fondata nel 2021 da un gruppo di ex dirigenti e ricercatori usciti proprio da OpenAI per divergenze sulla sicurezza. Elon Musk guida xAI, fondata nel 2023, oggi fusa con X (l'ex Twitter) sotto un'unica holding. Tre storie diverse, tre stili di comunicazione pubblica molto diversi tra loro — ma un'unica corsa comune, quella raccontata in

[[Libro_Ultima_Invenzione]], verso sistemi di intelligenza artificiale sempre più capaci.

Cosa significa «controllare» un'azienda di questo tipo

Il potere di ciascuno di loro non va misurato solo guardando quante azioni possiedono personalmente — un dato che, come vedremo, è spesso frainteso. Va misurato guardando chi ha davvero voce in capitolo sulle decisioni più importanti: quanto velocemente sviluppare un nuovo modello, quali test di sicurezza superare prima del lancio, quali accordi commerciali o politici stringere. Sono decisioni prese all'interno di consigli di amministrazione composti da poche persone, con un peso di voto che raramente viene reso pubblico nel dettaglio.

Un dato che smentisce un'idea diffusa

Vale la pena chiarirlo subito, perché è un errore comune: Sam Altman, CEO di OpenAI dal 2019, **non possiede alcuna quota azionaria diretta** nell'azienda che guida. Il suo potere non deriva, come si potrebbe pensare, da un enorme patrimonio personale legato a OpenAI, ma dal ruolo di guida operativa e dalla capacità di orientare pubblicamente il dibattito sull'intelligenza artificiale — un tipo di potere meno visibile nei bilanci societari, ma non meno reale.

Amodei e Musk, al contrario, mantengono partecipazioni dirette rilevanti nelle rispettive aziende —

un'asimmetria che, come racconta il resto di questo libro, rende ciascuno dei tre un caso a sé anche dal punto di vista degli incentivi economici personali in gioco.

Perché «i tre uomini» e non un elenco più lungo

Esistono altre figure di primo piano nel settore — dirigenti di Google DeepMind, Meta, e altri laboratori. Questo libro si concentra su Altman, Amodei e Musk perché rappresentano, insieme, il nucleo più diretto della corsa raccontata nel resto della serie «Intelligenza Artificiale» di questa collana: OpenAI e Anthropic sono i due laboratori più avanti nella competizione descritta in [[Libro_Ultima_Invenzione]], mentre xAI di Musk aggiunge un elemento che nessuno degli altri due possiede — il controllo diretto, nella stessa persona, anche di una piattaforma social globale e di un'infrastruttura satellitare, come racconta il Capitolo 4.

I prossimi capitoli raccontano, azienda per azienda, come si è arrivati a questa concentrazione di potere.

Capitolo 2 — OpenAI: da no-profit a colosso commerciale

Una missione nata per essere «a beneficio dell'umanità»

OpenAI nasce nel 2015 come organizzazione senza scopo di lucro, con una missione dichiarata esplicitamente pubblica: garantire che l'intelligenza artificiale generale, quando arriverà, vada a beneficio di tutta l'umanità, non di un singolo gruppo di persone o di un'azienda. Per anni, questa struttura no-profit è stata parte integrante dell'immagine pubblica dell'azienda — una promessa implicita che la ricerca non sarebbe stata guidata dal solo profitto.

La lunga trasformazione

Quella struttura si è progressivamente allontanata dal proprio disegno originario. Già nel 2019 OpenAI crea un ramo «a profitto limitato» per attrarre gli investimenti necessari a competere nella corsa al calcolo e ai talenti. Nel corso degli anni successivi, la parte commerciale dell'azienda cresce fino a diventare, di fatto, il cuore operativo di OpenAI — con la struttura no-profit che mantiene un ruolo di supervisione formale, ma sempre più distante dal controllo quotidiano.

Il passaggio più netto arriva nell'**ottobre 2025**, quando OpenAI completa una ristrutturazione societaria profonda: la no-profit diventa la **OpenAI Founda-**

tion, mentre il ramo commerciale diventa una **Public Benefit Corporation** chiamata OpenAI Group PBC — una forma societaria che, per legge, deve bilanciare il profitto con l'interesse pubblico dichiarato nella missione, ma che resta comunque un'azienda for-profit a tutti gli effetti, in grado di attrarre investitori e distribuire utili.

Chi possiede cosa, dopo la ristrutturazione

La OpenAI Foundation mantiene una partecipazione di controllo nel ramo commerciale, stimata in circa **130 miliardi di dollari**, pari a circa il **26%** della società. Dipendenti attuali ed ex dipendenti insieme agli investitori ne detengono circa il **47%**. Il singolo maggiore azionista esterno è **Microsoft**, con circa il **27%**, per un valore stimato di circa 135 miliardi di dollari.

All'inizio del 2026, un round di finanziamento guidato da Amazon, Nvidia e SoftBank porta altri **110 miliardi di dollari** nelle casse dell'azienda, con una valutazione record di **730 miliardi di dollari** — il più grande finanziamento privato della storia fino a quel momento.

Cosa resta della missione originaria

Non è possibile misurare con precisione quanto della missione originaria di OpenAI sopravviva nella pratica quotidiana dell'azienda — è in parte una do-

manda di giudizio, non solo di fatti. Ma i numeri di questo capitolo raccontano una traiettoria chiara: un'organizzazione nata esplicitamente per non essere guidata dalla logica del profitto è oggi valutata centinaia di miliardi di dollari, con un azionista di maggioranza relativa esterno (Microsoft) che ha un interesse commerciale diretto e legittimo nel successo economico dell'azienda, non nella sua missione dichiarata in astratto. La struttura della OpenAI Foundation resta come garanzia formale — ma resta da vedere, nei prossimi anni, quanto quella garanzia riuscirà a contare davvero quando gli interessi economici e quelli della missione dovessero entrare in conflitto diretto.

Capitolo 3 — Anthropic: nato da una scissione, ma nella stessa corsa

Fondata per fare le cose diversamente

Anthropic nasce nel 2021 da un gruppo di ex dirigenti e ricercatori di OpenAI, tra cui Dario Amodei, usciti proprio per divergenze sulla velocità con cui l'azienda stava sviluppando i propri sistemi rispetto alla prudenza che ritenevano necessaria. La storia fondativa di Anthropic si costruisce esplicitamente in opposizione a quella traiettoria: un'azienda che dichiara di voler mettere la sicurezza al centro del proprio sviluppo, non come vincolo esterno ma come parte della propria missione.

Una crescita commerciale che rivalessa con l'azienda da cui è nata

Il Capitolo 3 di [[Libro_Ultima_Invenzione]] racconta già, in questa stessa serie, la crescita esplosiva del fatturato di Anthropic: da circa un miliardo di dollari annualizzato all'inizio del 2025 a decine di miliardi nel 2026, una delle traiettorie di crescita più rapide mai registrate da un'azienda tecnologica. È un dato che va letto insieme alla storia fondativa raccontata sopra: un'azienda nata per rallentare, nei fatti, corre quanto o più dei concorrenti che voleva superare in prudenza.

Dario Amodei: proprietario e voce pubblica

A differenza di Sam Altman, Dario Amodei mantiene una partecipazione azionaria diretta e rilevante in Anthropic — un allineamento più diretto tra il proprio patrimonio personale e il successo commerciale dell'azienda, rispetto al caso OpenAI raccontato nel capitolo precedente. Amodei è anche, tra i tre protagonisti di questo libro, quello che più spesso ha parlato pubblicamente dei rischi della tecnologia che contribuisce a costruire — definendo, come racconta [[Libro_Ultima_Invenzione]], la sicurezza dell'intelligenza artificiale avanzata «la singola minaccia più seria alla sicurezza nazionale» degli ultimi cento anni, pur continuando a guidare un'azienda che compete direttamente per arrivare prima dei concorrenti su quella stessa tecnologia.

La stessa «trappola multipolare» raccontata altrove in questa serie

Il Capitolo 3 di [[Libro_Ultima_Invenzione]] descrive la dinamica competitiva che spinge ogni laboratorio ad accelerare anche quando i propri stessi dirigenti esprimono preoccupazione: nessuno può permettersi di rallentare da solo, per il timore che il concorrente arrivi prima a un vantaggio decisivo. Anthropic è, forse, il caso più chiaro di questa dinamica in azione: un'azienda fondata esplicitamente per fare le cose diversamente, che si trova comunque dentro la stessa corsa, con la stessa pressione competitiva, degli altri due protagonisti di questo libro.

Perché questo non è, di per sé, un'accusa di ipocrisia

Vale la pena essere onesti su un punto: non è chiaro se un'Anthropic che avesse davvero rallentato, restando fedele alla propria missione originaria alla lettera, avrebbe avuto un'influenza reale sulla sicurezza del settore nel suo complesso, o se sarebbe semplicemente rimasta ai margini di una corsa che gli altri due laboratori avrebbero continuato comunque. È il dilemma centrale raccontato nel Capitolo 3 di [[Libro_Ultima_Invenzione]], e Anthropic ne è, oggi, il caso di studio più concreto: restare nella corsa, sperando di poterla influenzare dall'interno, o restarne fuori, rischiando l'irrilevanza. Nessuna delle due scelte è priva di conseguenze, ed è un giudizio che questo

libro lascia al lettore, dopo avergli fornito i fatti necessari per formarselo.

Capitolo 4 — xAI e il conflitto di interessi di Musk

Un impero che si allarga, non si divide

Mentre OpenAI e Anthropic restano aziende dedicate quasi esclusivamente all'intelligenza artificiale, il caso di Elon Musk è diverso per natura: xAI non è un'azienda isolata, ma un pezzo di un insieme molto più ampio di attività controllate dalla stessa persona — Tesla, SpaceX, Starlink, e dal **28 marzo 2025** anche X, l'ex Twitter.

Quel giorno, Musk annuncia che xAI ha acquisito X in un'operazione interamente azionaria che valuta xAI **80 miliardi di dollari** e X **33 miliardi di dollari**, per un'entità combinata da circa **113 miliardi di dollari** sotto un'unica holding, xAI Holdings Corp. Vale la pena notare un dettaglio spesso omissso da chi racconta l'operazione come un puro successo: i 33 miliardi attribuiti a X rappresentano un calo netto rispetto ai 44 miliardi che lo stesso Musk aveva pagato per acquistarla nel 2022.

Perché la fusione tra un'IA e un social network non è un dettaglio tecnico

Unire in un'unica azienda un laboratorio di intelligenza artificiale e una delle principali piattaforme social del mondo crea una combinazione che nessuno degli altri due protagonisti di questo libro possiede: l'accesso diretto a un flusso enorme di dati generati in tempo reale da centinaia di milioni di utenti, utilizzabile per addestrare i propri modelli, insieme al controllo diretto sulla piattaforma dove quei modelli possono essere distribuiti e promossi. È una integrazione verticale — dal dato grezzo, all'addestramento, alla distribuzione al pubblico — che concentra in un'unica persona un controllo su più fasi della filiera dell'informazione digitale raramente visto nella storia recente.

Il caso DOGE: quando il regolatore e il regolato sono la stessa persona

Nel **febbraio 2025**, la sottocommissione investigativa permanente del Senato statunitense solleva formalmente un problema specifico: Musk ricopriva contemporaneamente un ruolo di guida in xAI e un incarico all'interno del **DOGE** (Department of Government Efficiency), l'organismo voluto dall'amministrazione Trump per tagliare la spesa pubblica federale. La preoccupazione, espressa per iscritto dai senatori, era che questa doppia posizione potesse permettere a xAI di ottenere vantaggi rispetto ai concorrenti e di influen-

zare in modo improprio proprio le agenzie federali che avrebbero dovuto regolarla.

Musk operava come «**special government employee**», uno status che per legge limita l'incarico a 130 giorni all'anno. Il termine viene raggiunto il **28 maggio 2025**, giorno in cui Musk annuncia pubblicamente la propria uscita dal ruolo, in un mandato segnato da controversie legali e da stime di risparmio pubblico contestate da diversi osservatori indipendenti.

Starlink e la stessa domanda, in un altro settore

Un problema simile riguarda **Starlink**, la rete di satelliti per internet controllata da SpaceX, altra azienda di Musk: le attività di lancio spaziale di SpaceX sono regolate dalla Federal Aviation Administration (FAA), la stessa agenzia con cui Starlink ha stretto accordi per l'integrazione dei propri terminali nella rete di gestione dello spazio aereo statunitense. Anche qui, chi vigila e chi è vigilato finiscono per intrecciarsi attraverso la stessa persona.

Una nota di chiarezza necessaria

Per accuratezza, vale la pena chiarire un punto su cui circola spesso confusione: **Neuralink**, l'azienda di Musk che sviluppa impianti cerebrali, è una società separata da xAI, con un'attività completamente diversa. Le due aziende non vanno confuse quando si valuta

la reale portata dell'influenza di Musk sul settore dell'intelligenza artificiale specificamente.

Capitolo 5 — Governance aziendale, non democratica

Dove si prendono davvero le decisioni

Nessuna delle scelte più importanti raccontate in questa serie — quanto velocemente sviluppare un nuovo modello, quali test di sicurezza superare prima di renderlo pubblico, come reagire a un incidente — passa da un parlamento, da un referendum, o da qualunque altro processo che permetta ai cittadini comuni di avere voce in capitolo. Passa da consigli di amministrazione composti, nella maggior parte dei casi, da meno di dieci persone: dirigenti dell'azienda, rappresentanti dei principali investitori, e un numero variabile di membri «indipendenti» scelti dagli stessi soci fondatori o dai grandi azionisti.

Il caso OpenAI: un consiglio che può licenziare (e reintegrare) il proprio CEO

Il caso più pubblico di questa dinamica resta l'episodio del **novembre 2023**, quando il consiglio di amministrazione di OpenAI licenzia improvvisamente Sam Altman, per poi reintegrarlo pochi giorni dopo sotto la pressione di dipendenti, investitori e dello stesso Microsoft, azionista di peso descritto nel Capitolo 2. L'episodio mostra, in modo plastico, dove risiede

davvero il potere di decisione su un'azienda di questa portata: non in un processo pubblico e trasparente, ma in un negoziato interno tra poche persone, in gran parte nascosto al pubblico fino a conclusione avvenuta.

Public Benefit Corporation: una garanzia formale, non una garanzia reale

La struttura di **Public Benefit Corporation** adottata da OpenAI nella ristrutturazione del 2025, descritta nel Capitolo 2, impone per legge al consiglio di amministrazione di bilanciare il profitto con l'interesse pubblico dichiarato nella missione. È una garanzia più forte di quella di una società tradizionale, ma resta comunque una garanzia interpretata e applicata dallo stesso consiglio di amministrazione che deve rispettarla — non da un organismo esterno indipendente con potere di veto reale.

Anthropic e il «Long-Term Benefit Trust»

Anthropic ha adottato un meccanismo di governance più elaborato: il **Long-Term Benefit Trust**, un organismo indipendente composto da **cinque persone senza interessi finanziari diretti** nell'azienda — tra i nomi che vi hanno fatto parte, l'ex presidente della Federal Reserve **Ben Bernanke** e l'ex giudice della Corte Suprema della California **Mariano-Florentino Cuéllar**, insieme a esperti di sicurezza nazionale e

politiche pubbliche. Il Trust ha il potere formale di scegliere e rimuovere una quota crescente del consiglio di amministrazione, fino a diventarne, nel tempo, la maggioranza — un traguardo già raggiunto, con i membri nominati dal Trust che oggi costituiscono la maggioranza del board.

È un meccanismo più solido, sulla carta, di quello di OpenAI raccontato nel Capitolo 2, proprio perché prevede un'evoluzione automatica verso un controllo esterno crescente. Ma resta comunque un sistema in cui cinque persone, per quanto competenti e senza interessi finanziari diretti, sono state scelte da chi ha fondato l'azienda — non da un processo pubblico o elettivo di alcun tipo. La sua efficacia reale, inoltre, non è mai stata messa alla prova in una crisi in cui gli interessi commerciali e quelli di sicurezza siano entrati in conflitto aperto.

Perché questo capitolo conta per il resto del libro

Il punto centrale, che collega questo capitolo al resto del libro, è semplice da enunciare ma importante da tenere a mente: ogni meccanismo di governance descritto qui — consigli di amministrazione, trust, strutture ibride no-profit/for-profit — resta comunque un sistema di controllo interno, deciso da chi ha fondato o finanziato l'azienda. Nessuno di questi meccanismi, per quanto sofisticato, equivale a un processo demo-

cratico in cui le persone effettivamente toccate dalle decisioni — cioè, in pratica, tutti — abbiano una voce reale nel decidere come vengono prese.

Capitolo 6 — Il lobbying

Chi scrive le regole che dovrebbero limitarli

Se il Capitolo 5 ha mostrato come si prendono le decisioni dentro queste aziende, questo capitolo guarda a un secondo livello, altrettanto importante: quanto queste stesse aziende investono per orientare le regole esterne — leggi, regolamenti, standard tecnici — che dovrebbero limitarne il comportamento.

Numeri in forte crescita, anno dopo anno

Nel 2025, la spesa complessiva di Anthropic in attività di lobbying federale negli Stati Uniti ha raggiunto circa **3,1 milioni di dollari**. Nel primo trimestre del 2026, OpenAI ha speso **1 milione di dollari** — il doppio dell'anno precedente — mentre Anthropic ha speso **1,6 milioni di dollari**, un aumento del **344%** rispetto ai 360.000 dollari dello stesso periodo del 2025.

Nel secondo trimestre del 2026, le due aziende insieme hanno speso **3,17 milioni di dollari**, il **23% in più** rispetto ai tre mesi precedenti. Anthropic da sola ha speso quasi 2 milioni di dollari, superando per la prima volta OpenAI nella spesa trimestrale — un sorpasso significativo, perché segnala quanto l'azienda nata per

«fare le cose diversamente», raccontata nel Capitolo 3, stia oggi investendo nella stessa arena di influenza politica dei concorrenti che diceva di voler distinguere.

Solo nel primo semestre del 2026, la spesa di Anthropic ha già superato il totale dichiarato per l'intero 2025.

Su cosa fanno pressione, esattamente

I temi dichiarati ufficialmente dalle due aziende in questa attività di lobbying includono cybersicurezza, diritto d'autore, cloud computing e appalti legati alla difesa — aree dove le nuove regole possono avere un impatto diretto sui costi operativi e sulle opportunità di mercato di entrambe le aziende. Non è possibile, dai soli dati di spesa dichiarata, ricostruire nel dettaglio quali proposte specifiche di legge siano state sostenute o osteggiate — le norme sulla trasparenza del lobbying negli Stati Uniti impongono di dichiarare importi e temi generali, non le posizioni precise assunte in ogni incontro.

Perché questa cifra, per quanto ancora piccola, conta

In termini assoluti, queste cifre restano modeste se confrontate con altri settori — l'industria farmaceutica o quella dei combustibili fossili spendono ordini di grandezza in più ogni anno. Ma il ritmo di crescita — più che raddoppiato in un solo trimestre in più di un'occasione — segnala un settore che sta rapida-

mente imparando le stesse regole del gioco politico già consolidate altrove: più un'industria diventa centrale per l'economia e la sicurezza nazionale, più investe per assicurarsi che le regole che la riguardano vengano scritte tenendo conto della propria prospettiva, prima che di quella dei cittadini che quelle regole dovrebbero proteggere.

Capitolo 7 — Le porte girevoli

Un fenomeno vecchio, applicato a un settore nuovo

Il termine «porte girevoli» descrive un fenomeno ben noto in altri settori raccontati da questa collana — quello bancario, quello farmaceutico, quello energetico: persone che passano da un ruolo di regolatore pubblico a un ruolo dentro le aziende che quello stesso regolatore dovrebbe controllare, e viceversa. È un meccanismo che non richiede nulla di illegale per essere un problema: basta che chi ha lavorato per anni fianco a fianco con un settore porti con sé, nel nuovo ruolo, relazioni, informazioni e sensibilità che un funzionario esterno non avrebbe.

Il settore dell'intelligenza artificiale non fa eccezione

Negli ultimi anni, diverse figure con ruoli di primo piano nelle agenzie federali statunitensi che si occupano di regolamentazione tecnologica e di sicurezza

nazionale sono passate a ruoli dirigenziali o consultivi dentro OpenAI, Anthropic o aziende collegate — e, simmetricamente, dirigenti di queste stesse aziende hanno assunto ruoli di consulenza o incarichi formali all'interno di amministrazioni pubbliche, incluso il caso di Elon Musk raccontato nel Capitolo 4, esempio più visibile in assoluto di questo fenomeno negli anni recenti.

Perché è difficile da misurare, e proprio per questo importante da segnalare

A differenza della spesa di lobbying, raccontata nel capitolo precedente, il fenomeno delle porte girevoli non è soggetto a un obbligo di dichiarazione pubblica sistematica e dettagliata: emerge soprattutto da inchieste giornalistiche, cambi di ruolo annunciati pubblicamente sui social professionali, o segnalazioni di organizzazioni indipendenti che monitorano questi passaggi. Questa scarsità di dati sistematici non significa che il fenomeno sia raro: significa che è strutturalmente più difficile da quantificare con precisione rispetto alla spesa di lobbying dichiarata.

Cosa comporta, in pratica

Il rischio concreto non è tanto la corruzione in senso stretto — un favore esplicito scambiato con denaro — quanto una forma più sottile di cattura regolatoria: chi ha passato anni della propria carriera dentro un settore tende, anche in perfetta buona fede, a condivi-

derne il linguaggio, le priorità, la visione di cosa sia un rischio accettabile. Quando queste stesse persone si trovano, in un ruolo successivo, a scrivere o applicare le regole per quel settore, la linea tra «conoscenza tecnica utile» e «prospettiva di parte» diventa difficile da tracciare — anche per la persona stessa coinvolta, non solo per chi osserva da fuori.

Un problema strutturale, non solo individuale

Come per il capitolo precedente sul lobbying, il punto di questo capitolo non è puntare il dito contro singole persone che si spostano legittimamente tra settore pubblico e privato — è una pratica comune e non illegale in molti sistemi democratici. Il punto è segnalare che, in un settore dove le decisioni tecniche hanno conseguenze enormi e dove la competenza specialistica è concentrata in pochissime aziende, il bacino da cui i governi possono attingere consulenti davvero indipendenti si restringe pericolosamente — lasciando, di fatto, che chi scrive le regole del gioco e chi le rispetta condividano spesso lo stesso background professionale.

Capitolo 8 — Quando i team di sicurezza se ne vanno

Un segnale che vale più di mille dichiarazioni pubbliche

Le dichiarazioni pubbliche di un'azienda sulla propria attenzione alla sicurezza sono facili da scrivere e difficili da verificare dall'esterno. Un segnale molto più concreto, e più difficile da controllare nella narrazione pubblica, è osservare cosa fanno le persone che lavoravano proprio su quella sicurezza, quando decidono di andarsene.

Il caso OpenAI: un esodo che dura da due anni

Il Capitolo 2 di *[[Libro_Ultima_Invenzione]]* racconta già la storia di Daniel Kokotajlo, uscito da OpenAI nell'aprile 2024. Non è stato un caso isolato. Nel **maggio 2024**, si dimettono insieme **Ilya Sutskever**, Chief Scientist e co-fondatore dell'azienda, e **Jan Leike**, co-responsabile del team «Superalignment» dedicato alla sicurezza a lungo termine. Leike dichiara pubblicamente che «la cultura e i processi di sicurezza sono passati in secondo piano rispetto ai prodotti scintillanti». Il team Superalignment viene sciolto subito dopo — e secondo quanto riportato pubblicamente, i suoi membri erano stati ripetutamente privati dell'hardware necessario per condurre i propri

test, con quei chip destinati invece a far funzionare ChatGPT per gli utenti.

L'esodo continua: nell'**ottobre 2024** si dimette **Miles Brundage**, capo consulente per la «AGI Readiness» (la preparazione dell'azienda all'arrivo di un'intelligenza artificiale generale), per fondare un'organizzazione indipendente sulla sicurezza. Alla fine del 2025, **Andrea Vallone** lascia OpenAI per raggiungere proprio Anthropic, seguendo lo stesso percorso di Leike. Nel **luglio 2026**, anche il responsabile sicurezza **Johannes Heidecke** lascia l'azienda, che integra i team di sicurezza rimasti sotto la responsabile della ricerca, eliminando la linea di segnalazione indipendente che questi team avevano fino a quel momento — la **sesta uscita di un dirigente senior della sicurezza in due anni**. Di circa trenta persone che lavoravano su temi legati alla sicurezza dell'AGI, ne restano oggi circa sedici.

Cosa significa, messo tutto insieme

Un singolo addio si può sempre spiegare con motivi personali. Sei uscite di alto profilo in due anni, con dichiarazioni pubbliche esplicite sul motivo — «la sicurezza è passata in secondo piano» — non sono più un caso isolato: sono un pattern. E un pattern di questo tipo, proveniente da persone che avevano accesso diretto alle decisioni interne dell'azienda, pesa più di qualunque comunicato stampa aziendale sulla sicurezza.

Non solo OpenAI, ma con intensità diversa

Anthropic, nata proprio dalla scissione raccontata nel Capitolo 3, non ha vissuto un esodo di sicurezza paragonabile per intensità — anzi, ha accolto alcune delle persone uscite da OpenAI, come Jan Leike e Andrea Vallone. È un dato che, insieme al Long-Term Benefit Trust raccontato nel Capitolo 5, rende Anthropic il caso relativamente più solido tra i tre sul piano della governance della sicurezza — pur restando, come mostra il resto di questo libro, pienamente dentro la stessa corsa competitiva degli altri due.

Cosa dovrebbe fare un lettore di questo capitolo

Questo capitolo non prova che il disastro raccontato in [[Libro_Ultima_Invenzione]] sia certo. Prova qualcosa di più modesto ma altrettanto importante: che le persone con l'accesso più diretto alle decisioni interne di questi laboratori, quando hanno dovuto scegliere tra restare in silenzio o andarsene dichiarando pubblicamente perché, in numero crescente hanno scelto la seconda strada. È un segnale che, per onestà, nessun libro su questi tre uomini può permettersi di ignorare.

Capitolo 9 — Il potere del compute

La barriera all'ingresso più forte di qualunque brevetto

Nei settori industriali tradizionali, una nuova azienda può entrare in competizione con i leader di mercato investendo in ricerca, assumendo talenti, trovando capitali. Nel settore raccontato da questo libro, esiste una barriera aggiuntiva, più difficile da superare di tutte le altre messe insieme: l'accesso al calcolo (il «compute») necessario ad addestrare sistemi di intelligenza artificiale avanzati.

Il Capitolo 6 di `[[Libro_Guerra_Chip]]`, altro titolo di questa stessa serie, racconta nel dettaglio quanto sia concentrata la produzione dei chip più avanzati e quanto costino i data center necessari a farli funzionare insieme su scala. Questo capitolo guarda alla stessa concentrazione da un angolo diverso: cosa significa, per il potere relativo di Altman, Amodei e Musk, avere accesso privilegiato a quella stessa risorsa.

Chi ha già vinto la gara prima di cominciare

OpenAI, Anthropic e xAI non competono su un terreno neutro. Ciascuna delle tre aziende ha costruito, negli ultimi anni, accordi enormi con i pochi fornitori mondiali di calcolo — Microsoft per OpenAI, Amazon e Google per Anthropic, l'infrastruttura propria co-

struita da xAI attraverso ingenti investimenti diretti di Musk. Una nuova azienda che volesse competere con loro partendo da zero oggi non incontrerebbe solo il problema di trovare buone idee o buoni ricercatori: incontrerebbe l'impossibilità pratica di accedere, a condizioni ragionevoli, alla stessa quantità di calcolo che questi tre laboratori hanno già assicurato attraverso anni di accordi e investimenti.

Un circolo che si autoalimenta

Più un'azienda ha accesso al calcolo, più può addestrare modelli capaci, più genera ricavi e attira nuovi investimenti, più può permettersi di comprare o costruire ancora più capacità di calcolo. È un circolo che tende, per sua stessa natura, a concentrare il potere invece che distribuirlo — l'opposto di quello che accade in molti mercati dove nuovi concorrenti riescono, nel tempo, a insidiare i leader consolidati. Nel settore raccontato da questo libro, il vantaggio di chi è arrivato prima, e ha già assicurato accesso al calcolo necessario, tende a consolidarsi anno dopo anno invece di ridursi.

Perché questo capitolo chiude la prima parte del libro

I capitoli precedenti hanno mostrato come il potere di questi tre uomini si eserciti attraverso la struttura societaria (Capitoli 2-3), il controllo diretto su più settori contemporaneamente (Capitolo 4), la governance

interna (Capitolo 5), l'influenza sulle regole esterne (Capitoli 6-7), e la gestione interna del dissenso sulla sicurezza (Capitolo 8). Questo capitolo aggiunge l'ultimo tassello, forse il più strutturale di tutti: anche se domani emergesse un quarto o un quinto attore con idee migliori e intenzioni più prudenti delle attuali, la barriera del compute renderebbe estremamente difficile per quell'attore competere davvero con i tre protagonisti di questo libro. Il potere descritto in questo libro, in altre parole, non è solo concentrato oggi: è costruito in un modo che tende a restare concentrato anche domani.

Capitolo 10 — Cosa NON è mai stato dimostrato

Perché questo capitolo è necessario quanto tutti gli altri

I capitoli precedenti hanno descritto uno squilibrio di potere reale, documentato con fatti verificabili: struttura societaria, quote di controllo, spesa di lobbying, uscite di dirigenti della sicurezza. Non c'è bisogno di aggiungere altro per rendere quello squilibrio serio. Eppure, online, circolano su Altman, Amodei e soprattutto Musk affermazioni molto più estreme di quelle raccontate finora — e altrettanto spesso, prive di prove verificabili. Questo capitolo esiste per un motivo preciso: chi diffonde affermazioni false su queste tre

persone, anche in perfetta buona fede, indebolisce involontariamente la credibilità delle critiche reali e documentate, dando a chi le difende un facile pretesto per liquidarle tutte insieme come «disinformazione».

«Sam Altman vuole impiantare micro-chip obbligatori nella popolazione»

Questa affermazione, molto diffusa online, nasce da una confusione o da un'esagerazione attorno a **Worldcoin**, un progetto reale co-fondato da Altman: un dispositivo chiamato «Orb» che scansiona l'iride dell'occhio per creare un'identità digitale univoca, su base **volontaria**, pensato per distinguere le persone reali dai sistemi automatizzati online. Non esiste alcun impianto fisico nel corpo, e non esiste alcuna prova di un obbligo imposto a chicchessia di sottoporsi alla scansione. Il progetto solleva legittime domande sulla privacy e sulla raccolta di dati biometrici — argomenti seri, che non hanno bisogno di essere esagerati in «microchip obbligatori» per essere presi sul serio.

«Elon Musk sta impiantando chip cerebrali attraverso xAI»

Anche questa affermazione nasce da una confusione, in questo caso tra due aziende diverse: **Neuralink**, che sviluppa effettivamente impianti cerebrali per scopi medici dichiarati (in primo luogo per persone con gravi disabilità motorie), è una società separata da **xAI**, l'azienda di intelligenza artificiale al centro

di questo libro. Le due aziende condividono lo stesso fondatore, ma non la stessa attività, gli stessi prodotti, né lo stesso management. Attribuire a xAI i progetti di Neuralink è un errore fattuale, non un'interpretazione critica legittima.

«Il licenziamento di Altman nel 2023 fu causato da una AGI segreta già raggiunta»

Nel novembre 2023, il consiglio di amministrazione di OpenAI licenzia improvvisamente Altman, per poi reintegrarlo pochi giorni dopo — un episodio reale, raccontato nel Capitolo 5. Ha circolato ampiamente online la teoria secondo cui il licenziamento sarebbe stato causato dalla scoperta interna di un sistema di intelligenza artificiale generale già raggiunto in segreto (talvolta chiamato con il nome in codice «Q*»), troppo pericoloso per essere reso pubblico. Questa teoria **non è mai stata confermata da alcuna fonte primaria verificabile**. Le dichiarazioni ufficiali del consiglio dell'epoca hanno parlato genericamente di un«collasso di comunicazione e fiducia» con Altman — un motivo vago, che ha alimentato le speculazioni proprio per la sua vaghezza, ma che non equivale a una conferma della teoria dell'AGI segreta.

«Queste aziende ricevono finanziamenti governativi segreti»

Circolano periodicamente affermazioni secondo cui OpenAI, Anthropic o xAI riceverebbero finanziamenti diretti da agenzie governative senza alcun contratto pubblico — al di là dei rapporti commerciali legittimi e in parte pubblici che queste aziende hanno con agenzie della difesa e dell'intelligence, raccontati indirettamente nel Capitolo 6 a proposito del lobbying su «appalti della difesa». Senza un contratto pubblico documentato o una fonte giornalistica verificabile con prove dirette, queste affermazioni restano voci non dimostrate, e questo libro non le riprende come fatti.

Perché la disciplina conta più della provocazione

Un libro che si intitola «I tre uomini che decidono per tutti» potrebbe essere tentato di rendere la propria tesi più forte aggiungendo le voci più sensazionali che circolano online. Questo libro fa la scelta opposta, per lo stesso motivo per cui l'ha fatta ogni altro titolo di questa collana che tratta figure pubbliche specifiche: i fatti verificabili raccontati nei capitoli precedenti sono già, da soli, una critica seria e sufficiente. Aggiungere affermazioni false non la rafforzerebbe. La renderebbe più facile da respingere.

Capitolo 11 — Le obiezioni, trattate con onestà

«Qualcuno doveva pur farlo, ed è meglio che siano loro»

È l'obiezione più diffusa tra chi difende questi tre laboratori: se la tecnologia raccontata in questa serie sta comunque per essere costruita — da qualcuno, da qualche parte nel mondo — è meglio che a farlo siano aziende occidentali con una qualche forma di supervisione, per quanto imperfetta, piuttosto che attori senza alcun vincolo dichiarato di sicurezza. È un argomento che ha un fondo di verità: il Capitolo 4 di [[Libro_Guerra_Chip]] racconta gli sforzi massicci della Cina per costruire una propria capacità autonoma in questo campo, con vincoli di sicurezza dichiarati diversi da quelli occidentali. Ma l'obiezione non risponde alla domanda centrale di questo libro: anche accettando che qualcuno debba farlo, perché quel «qualcuno» deve restare così concentrato in tre sole aziende, senza un controllo democratico reale?

«Sono aziende private, hanno diritto di operare secondo le proprie regole»

È vero, in senso stretto: OpenAI, Anthropic e xAI sono entità private, non enti pubblici, e operano secondo le leggi societarie dei paesi in cui sono registrate. Ma questo argomento vale meno quanto più le conse-

guenze delle loro decisioni superano i confini di un mercato normale. Una banca privata che fallisce può trascinare con sé l'intera economia di un paese — motivo per cui viene regolata più severamente di un piccolo negozio. Il Capitolo 9 di questo libro, sul potere del compute, e l'intera serie «Intelligenza Artificiale» di questa collana raccontano perché le conseguenze delle decisioni di queste tre aziende potrebbero, nello scenario descritto in [[Libro_Ultima_Invenzione]], superare di gran lunga quelle di qualunque banca. Il principio che vale per le banche — più le conseguenze sono sistemiche, più serve un controllo pubblico proporzionato — non sembra applicato con la stessa intensità a questo settore, almeno non ancora.

«Amodei e altri dirigenti parlano apertamente dei rischi — non è ipocrisia, è trasparenza»

È un'obiezione che riguarda soprattutto Dario Amodei, il più esplicito dei tre nel parlare pubblicamente dei rischi della propria stessa tecnologia. Va riconosciuto con onestà: parlare apertamente dei rischi, invece di minimizzarli, è meglio del silenzio — ed è un comportamento diverso da quello, per esempio, di altri settori industriali storicamente più reticenti sui propri rischi (l'industria del tabacco, per fare un esempio classico). Ma restare nella corsa competitiva, continuando ad accelerare lo sviluppo mentre si descrivono pubblicamente rischi enormi, resta comunque una

scelta — non un'azione neutra. Riconoscere il merito della trasparenza non equivale a considerare risolto il problema di fondo raccontato in questo libro.

«Il potere di queste aziende è comunque limitato dalla concorrenza tra loro»

Un'ultima obiezione sostiene che, proprio perché OpenAI, Anthropic e xAI competono ferocemente tra loro, nessuna delle tre può davvero abusare del proprio potere — la concorrenza stessa farebbe da freno. È un argomento parzialmente vero per il rapporto tra le tre aziende, ma non spiega cosa succede al rapporto tra queste tre aziende insieme e il resto della società: la concorrenza tra loro non produce automaticamente un controllo democratico esterno, semplicemente sposta la domanda su chi, tra i tre, arriva prima a un vantaggio decisivo — esattamente la «trappola multipolare» descritta nel Capitolo 3 di [[Libro_Ultima_Invenzione]].

Cosa resta, dopo aver ascoltato queste obiezioni

Nessuna di queste obiezioni cancella il quadro centrale di questo libro. Ma insieme aiutano a definirlo con più precisione: il problema non è che queste tre persone siano malintenzionate — nessuna prova raccolta in questo libro lo sostiene. Il problema è strutturale: un settore le cui decisioni potrebbero avere conseguenze enormi per l'umanità intera resta oggi governato da

meccanismi di responsabilità pensati per aziende ordinarie, non per attori di questa portata.

Capitolo 12 — Cosa si potrebbe cambiare

Obblighi di trasparenza sulla governance interna, non solo sui prodotti

Oggi le uniche informazioni pubbliche affidabili sulla governance interna di queste tre aziende arrivano da comunicati aziendali volontari o da inchieste giornalistiche indipendenti. Una prima proposta concreta, discussa da diversi esperti di regolamentazione tecnologica, è imporre per legge obblighi di trasparenza specifici sulla composizione dei consigli di amministrazione, sui meccanismi di veto interno sulla sicurezza, e sulle motivazioni reali delle uscite di dirigenti chiave — non lasciate alla sola volontà comunicativa dell'azienda.

Regole più severe sui conflitti di interesse nei ruoli pubblici

Il Capitolo 4 ha raccontato il caso di Musk e del DOGE, e il Capitolo 7 il fenomeno più ampio delle porte girevoli. Regole più stringenti sul divieto, o quantomeno sulla piena divulgazione pubblica, di doppie posizioni tra ruoli dirigenziali in aziende di intelligenza artificiale e incarichi pubblici di regolamentazione dello

stesso settore ridurrebbero almeno la parte più visibile di questo problema, anche se non lo eliminerebbero del tutto.

Estendere protezioni reali a chi segnala problemi di sicurezza

Il Capitolo 8 ha mostrato quante persone con accesso diretto alle decisioni interne di questi laboratori abbiano scelto di andarsene, spesso pubblicamente, per segnalare la propria preoccupazione. Il Capitolo 13 di [[Libro_Ultima_Invenzione]] racconta già la proposta di un «diritto di avvertire» per chi lavora nel settore — protezioni legali per chi vuole segnalare rischi gravi senza dover rinunciare a tutto, incluso lo stipendio e la propria carriera, per farlo. È una proposta che risponde direttamente al problema descritto in questo libro: rendere meno costoso, per chi vede un rischio dall'interno, dirlo apertamente.

Regole antitrust pensate per il compute, non solo per i mercati tradizionali

Il Capitolo 9 ha mostrato come il controllo sull'accesso al calcolo costituisca oggi una barriera all'ingresso più forte di qualunque brevetto tradizionale. Le leggi antitrust esistenti, pensate per mercati di prodotti fisici o servizi tradizionali, faticano ad applicarsi a una risorsa di questo tipo. Alcuni giuristi ed economisti propongono di ripensare esplicitamente il diritto della concorrenza per includere l'accesso al compute

come una risorsa da monitorare e, se necessario, regolamentare — sul modello di come si regolano oggi le infrastrutture essenziali come l'energia o le telecomunicazioni.

Il filo comune

Come per gli altri titoli di questa serie, nessuna di queste proposte, presa da sola, risolve il problema descritto in questo libro. Ma condividono lo stesso principio di fondo: che un potere di questa portata, concentrato in così poche mani, non può restare governato solo dalle regole ordinarie pensate per aziende comuni — servono strumenti proporzionati alla scala reale delle conseguenze in gioco.

Capitolo 13 — Cosa puoi fare tu, adesso

Distinguere la persona pubblica dal potere strutturale

Il primo passo, quando leggi una notizia su Altman, Amodei o Musk, è provare a distinguere due cose spesso confuse nel dibattito pubblico: le opinioni o la personalità di ciascuno di loro, e il potere strutturale che deriva dal ruolo che occupano. Il Capitolo 10 di questo libro ha mostrato quanto sia facile, online, scivolare verso accuse personali non dimostrate — un errore che indebolisce, invece di rafforzare, la critica

reale a uno squilibrio di potere che esiste indipendentemente da chi occupa oggi quelle tre poltrone.

Informarsi alla fonte sulla governance, non solo sui prodotti

I documenti citati in questo libro — la struttura della OpenAI Foundation, quella del Long-Term Benefit Trust di Anthropic, i dati sul lobbying dichiarato — sono in gran parte pubblici, anche se dispersi tra siti aziendali, archivi normativi e inchieste giornalistiche. Cercare attivamente questo tipo di informazione, invece di limitarsi alle notizie sui nuovi prodotti lanciati da queste aziende, è un modo concreto per formarsi un giudizio basato sui fatti invece che sulla sola percezione pubblica costruita dalla comunicazione aziendale.

Sostenere chi lascia per raccontare quello che ha visto

Le persone raccontate nel Capitolo 8 — da Daniel Kotajlo, protagonista di *[[Libro_Ultima_Invenzione]]*, ai membri del team Superalignment di OpenAI — hanno pagato un prezzo personale, spesso economico, per parlare apertamente. Seguire il loro lavoro successivo, dividerlo, sostenere le organizzazioni indipendenti che fondano dopo l'uscita, è un modo concreto per rendere quella scelta meno costosa per chi, in futuro, si troverà davanti allo stesso dilemma.

Fare pressione politica sulle proposte del capitolo precedente

Le misure discusse nel Capitolo 12 — trasparenza sulla governance, regole sui conflitti di interesse, protezioni per chi segnala rischi, antitrust ripensato per il compute — non si scrivono da sole. Richiedono che diventino argomento di dibattito politico reale, non solo materia per esperti di settore. Chiedere esplicitamente a chi ci rappresenta quale posizione ha su questi temi specifici, non genericamente «sull'intelligenza artificiale», è un modo concreto per portarli fuori dalla nicchia degli addetti ai lavori.

Condividere questo libro

Se questo libro ti ha aiutato a vedere con più chiarezza chi decide davvero la traiettoria di una tecnologia che riguarda tutti, condividilo con chi, nella tua vita, ne parla senza aver mai riflettuto su chi, concretamente, prende quelle decisioni — non per sostituire il tuo giudizio al suo, ma per dargli gli stessi fatti verificabili su cui costruire il proprio.

Nota sulla trasparenza nell'uso dell'intelligenza artificiale

Questo libro è stato scritto con il supporto di strumenti di intelligenza artificiale, usati per la ricerca preliminare delle fonti, la prima stesura dei testi a partire

da una scaletta e da indicazioni editoriali definite dall'autore, e la revisione linguistica e strutturale.

Restano interamente umane: la scelta del tema e la tesi del libro, la verifica di ogni fatto, cifra e citazione contro una fonte primaria o istituzionale prima della pubblicazione, la decisione editoriale finale su cosa viene pubblicato o corretto, e la responsabilità editoriale dei contenuti — che resta in capo all'autore, mai allo strumento di intelligenza artificiale utilizzato.

Questa nota risponde agli obblighi di trasparenza previsti dall'art. 50 del Regolamento (UE) 2024/1689 (AI Act), applicabile dal 2 agosto 2026.

Se trovi un errore, una fonte imprecisa o un'affermazione che ritieni non verificata correttamente, scrivi a redazione@contropotere.com — ogni segnalazione fondata porta a una correzione pubblica del testo. Questo libro fa parte della collana **Contropotere**: gli altri titoli sono disponibili gratuitamente su **contropotere.com**.