

Il saggio analizza la crisi in corso sull'IA di frontiera, smascherando il panico oligarchico del p(doom) come cortina di fumo ideologica per occultare l'IA dell'auto-miglioramento ricorsivo (RSI) industriale, l'ecocidio idrico e una bolla speculativa da 3.500 miliardi di dollari di shadow financing. Evidenziando il fallimento della regolamentazione privata e lo stallo istituzionale di Washington e dell'UE, il testo traccia l'inasprimento della guerra cibernetica e della frattura open-source tra Stati Uniti e Cina. La transizione ecosocialista si fa prassi attraverso il reverse engineering sindacale basato sull'audit qualitativo dei tecnici indipendenti contro la negligenza dei laboratori. La proposta politica culmina nell'espropriazione legale delle server-farm degli hyperscaler, convertite in utility pubbliche decentralizzate e autogestite da Consigli Territoriali del Digitale per difendere la biosfera e il tempo vivo.

1. L'illusione del controllo e il ricatto dei miliardari

Lo scenario si apre con la testimonianza di martedì della scorsa settimana (8 settembre), quando Jacob Coxon, un ex ricercatore di OpenAI che si era recentemente trasferito in Anthropic, si è dimesso dal suo incarico, avvertendo che “nessuna delle due aziende si sta comportando in modo responsabile” e che “le persone che sviluppano l'IA credono sinceramente che potrebbe ucciderci tutti entro la fine del decennio”. Le sue dichiarazioni pubbliche hanno rapidamente innescato un acceso dibattito politico globale, spingendo persino alcuni membri democratici e repubblicani del Congresso statunitense a richiedere immediate regolamentazioni normative del settore. Nel fine settimana, il New York Times ha riportato che giovedì scorso Barack Obama ha esortato privatamente i Democratici a dare priorità alla supervisione dell'intelligenza artificiale nella loro agenda, dicendo ai leader del partito: “Questo è un settore che si sta muovendo molto velocemente nelle mani di privati e, se non lo controlliamo, penso che possa essere pericoloso”. Coxon ha sottolineato che non si tratta di una trovata di *marketing*, ma che stanno correndo a perdifiato verso la superintelligenza auto-migliorante e stanno giocando con le nostre vite.

Mentre Anthropic si prepara ad una quotazione in Borsa potenzialmente da duemila miliardi di dollari, gli esperti interni all'allineamento confermano questo allarme stimando una probabilità superiore al 10% che l'IA causi l'estinzione entro il decennio. La vera minaccia tecnologica risiede nel passaggio critico dell'auto-miglioramento ricorsivo (*recursive self improvement* - RSI): un sistema in grado di progettare da solo versioni migliori di se stesso a ritmi incontrollabili, rendendo la partecipazione e la supervisione umana sempre più marginali e sottili. A dimostrare che questa transizione non è un'ipotesi futuribile ma una prassi industriale in atto sono i dati interni rivelati da Reuters, secondo cui lo strumento autonomo *Claude Code* genera già la quota maggioritaria del *software* utilizzato nei progetti di Anthropic, permettendo agli ingegneri in carne e ossa di rilasciare una quantità

di codice per trimestre ben otto volte superiore rispetto alla produzione storica degli ultimi quattro anni.

Questo clima di inquietudine profonda tra i lavoratori del silicio ha trovato una forma di resistenza organizzata nella nascita della Coalition of Concerned AI Staff, un progetto clandestino di stampo *skunk works* all'interno del quale il ricercatore senior Evan Hubinger e altri due dipendenti di Anthropic sono intervenuti a supporto, confermando esplicitamente la concretezza del rischio esistenziale a breve e lungo termine. Questa mobilitazione dal basso poggia sulla storica lettera aperta firmata a luglio da ben 1.386 dipendenti dei laboratori di frontiera, i quali hanno denunciato pubblicamente la totale assenza di piani credibili per controllare una superintelligenza, definendo la corsa alla costruzione di entità più intelligenti dell'uomo come un'azione oggettivamente folle e suicida. Voci interne di prima linea, come l'ex OpenAI Pamela Mishkin, la ricercatrice sulla sicurezza Julie Steele e Andreas Kirsch di Google DeepMind, sono uscite allo scoperto confermando la concretezza del rischio.

2. La controriforma padronale e la normalizzazione del panico

Tuttavia, come denunciato da un'inchiesta di JS Tan e Clarissa Redwine su *The Boston Review*, questa timida primavera del sindacalismo tecnologico si è scontrata con una violenta controriforma padronale orchestrata dall'oligarchia dei miliardari, determinata a estirpare il dissenso interno trattandolo come una minaccia eversiva. Trovando in Donald Trump il partner politico ideale, i signori del silicio hanno stretto un'alleanza reazionaria capace di coniugare la deregolamentazione dell'IA con una feroce repressione sindacale, demolendo sistematicamente i diritti e le tutele dei lavoratori del codice per blindare l'assoluta docilità della forza lavoro di fronte all'automazione.

La narrazione apocalittica può essere definita una deliberata cortina di fumo ideologica: l'aristocrazia della Silicon Valley agita lo spettro di una superintelligenza fantascientifica ribelle per terrorizzare le istituzioni e distogliere lo sguardo dai crimini materiali commessi nel presente, come l'ecocidio delle risorse idriche, il saccheggio neocoloniale del lavoro e il consolidamento di una rendita privata iper-concentrata. La convergenza verso l'apocalisse biologica non è una fatalità tecnologica, ma la conseguenza diretta di una scommessa economica cinica: l'oligarchia della Silicon Valley accetta consapevolmente il rischio di estinzione della specie pur di centrare l'obiettivo supremo della rendita privata e dell'automazione selvaggia. Come svelato dall'inchiesta del Financial Times, più di due dozzine di scienziati e ricercatori senior del settore hanno confermato che l'intensa competizione commerciale e la radicale sfiducia psicologica tra i leader delle Big Tech stanno esasperando i rischi esistenziali per l'umanità. La spinta predatoria del mercato

costringe i laboratori a tagliare gli angoli della sicurezza e a rilasciare sistemi fuori controllo pur di non cedere un millimetro di vantaggio competitivo ai rivali.

Ci troviamo di fronte all'atto finale dell'alienazione, dove il capitale – pur di non rinunciare alla sottomissione totale della forza lavoro e alla massimizzazione del profitto immateriale – preferisce sussumere e mettere a repentaglio la sopravvivenza stessa della biosfera. La favola tecno-ottimista crolla sotto il peso di questa verità: l'intelligenza artificiale generale, definita AGI, non nasce per liberare l'uomo, ma per spezzare definitivamente la spina dorsale della classe operaia cognitiva. Le inchieste del New York Times confermano che questo clima di profonda inquietudine scientifica non è più circoscritto a cerchie ristrette: il dibattito pubblico globale è ora dominato da foschi pensieri circa il potenziale distruttivo intrinseco di una tecnologia che, pur configurandosi come un formidabile traguardo dell'ingegno umano, possiede la concreta capacità di azzerare la nostra intera civiltà. Questo allarme viene tuttavia fagocitato dall'infosfera liberal-progressista che, come denunciato da Dylan Matthews su Vox, opera una vera e propria "normalizzazione egemone del panico". Riducendo lo spettro dell'estinzione biologica a spettacolo mediatico e materiale da *talk-show* per Capitol Hill, il realismo capitalista compie il suo capolavoro: la corsa alla superintelligenza viene presentata come una transizione metafisica inevitabile a cui abituarsi, blindata dal ricatto geopolitico secondo cui una pausa unilaterale favorirebbe Pechino.

Eppure, questa recinzione spettacolare ha subito un cedimento strutturale nel momento in cui il panico collettivo ha valicato i confini dell'infosfera d'élite per farsi pienamente *mainstream*. L'onda d'urto generata dal dissenso dei ricercatori e dagli attacchi autonomi degli sciame ha costretto le istituzioni di Washington a confrontarsi con l'insufficienza dei vecchi strumenti giuridici, spingendo il dibattito federale verso tre opzioni radicali e un tempo impensabili: il congelamento coatto e il bando dei modelli avanzati di frontiera, l'imposizione di licenze governative con monitoraggi intrusivi e l'introduzione di una responsabilità civile e penale oggettiva per i danni algoritmici. Questa possibile accelerazione normativa esprime l'estremo tentativo della burocrazia statale di inseguire una tecnologia che minaccia di lasciarla sul fondo, nel disperato tentativo di normare un capitale informatico intrinsecamente eversivo molto prima che esso completi la propria catena di reazione autonoma.

Questa capitolazione culturale è speculare a quella ideologica che emerge dalle stesse confessioni dei vertici dei laboratori riportate da Mike Isaac e Kate Conger sul *New York Times*. Paul Christiano, inventore di metodologie chiave per l'addestramento e membro del *board* di OpenAI, ha ammesso lo spettro di una "perdita di controllo catastrofica e irreversibile nel brevissimo termine". Gli ha fatto eco il Chief Scientist di OpenAI, Jakub Pachocki, confessando che "nessuno è preparato per le conseguenze di un rapido e continuo aumento dell'intelligenza delle macchine" e invocando interventi normativi

globali. Persino il capo delle politiche pubbliche dell'azienda, Chris Lehane, è stato costretto a una clamorosa ritirata strategica, dichiarando che OpenAI dovrà "rallentare o fermare lo sviluppo e il dispiegamento dei sistemi" qualora non sia in grado di salvaguardarli adeguatamente. La centralità di questa ritirata è emersa drammaticamente durante una tesissima riunione aziendale interna, in cui lo stesso Sam Altman ha dovuto cedere alle pressioni dei dipendenti, dichiarando che OpenAI sarebbe pronta a rallentare lo sviluppo in parallelo con i laboratori concorrenti e auspicando che le decelerazioni volontarie diventino una prassi del settore.

Siamo dinanzi al compimento definitivo di quel processo di espropriazione del sapere e della coscienza che la Silicon Valley, operando secondo la logica predatoria dell'accumulazione capitalistica, tenta di imporre all'intera umanità. Questo assalto sinaptico h24, che recinta metodicamente i nostri tempi morti e trasforma l'attenzione e la noia in una *slot machine* estrattiva a danno della carne e del tempo vivo, è stato da me sviscerato nei saggi *"Dalla distrazione alla prassi: critica del realismo capitalista e riscatto ecosocialista"* e *"Tutto il resto è noia. Secessione del tempo vivo contro il ricatto delle Big Tech e la trappola della connessione perenne"*. La minaccia dell'estinzione biologica e lo svantaggio generazionale non innescano una rivolta, ma un'anestesia collettiva diffusa, una «indifferenza smaltata» che paralizza l'azione delle masse.

Questo clima di profonda inquietudine e sbandamento si ripercuote direttamente sulle nuove generazioni: schiacciati dall'automazione selvaggia, i giovani rinunciano a pianificare l'università e il futuro professionale, privati di ogni certezza. Si tratta di un vuoto esistenziale che viene immediatamente monetizzato dalle Big Tech: mentre le *server-farm* divorano le risorse del pianeta, i modelli di IA si impongono come i massimi fornitori di supporto psicologico negli Stati Uniti. L'interazione uomo-macchina sostituisce i legami interpersonali per 24 ore al giorno, distruggendo il piacere del pensiero critico individuale e negando radicalmente la fioritura umana. È l'atto finale del ricatto: la frammentazione emotiva prodotta dal tecno-capitalismo si trasforma in un nuovo mercato terapeutico chiuso, preparando il terreno a quel surriscaldamento speculativo e a quel panico istituzionale che colpirà i mercati finanziari. A confermare la natura manipolatoria di questo dispositivo psicologico sono le dichiarazioni rilasciate da Sam Altman il 15 settembre 2026 durante una conferenza a San Francisco, in cui il CEO di OpenAI ha sfacciatamente affermato che il mondo ha "il diritto di avere paura" ma deve al contempo "fidarsi" della capacità delle Big Tech di autoregolamentarsi, tentando così di disinnescare la vigilanza popolare per blindare l'impunità dell'industria del silicio.

3. Finanziarizzazione tecnologica e bolle speculative

Tuttavia, sotto il peso delle crescenti paure per la sicurezza e l'allineamento, lo stesso Sam Altman ha dovuto bruscamente frenare, confermando ufficialmente in un'intervista a *Fortune* nel fine settimana che OpenAI non procederà allo sbarco in Borsa nel 2026, giustificando il rinvio della potenziale quotazione da mille miliardi di dollari proprio con le crescenti preoccupazioni legate alla sicurezza e alla perdita di controllo della tecnologia. La conferma di questo slittamento al 2027 trova la sua spiegazione materiale nell'inchiesta pubblicata dal New York Times il 16 settembre: impossibilitata a sbarcare sui listini pubblici nell'immediato, OpenAI ha avviato trattative riservate per un colossale round di finanziamento privato pre-IPO basato su una valutazione shock che oscilla tra i 1.200 e i 1.500 miliardi di dollari. Questa manovra, nata su pressione di investitori determinati a raddoppiare la precedente capitalizzazione di 730 miliardi, svela il circolo vizioso che stringe l'aristocrazia del silicio: i laboratori sono costretti a rastrellare flussi incessanti di liquidità privata per tamponare un fabbisogno di calcolo e di M&A iper-inflazionato, trasformando le metriche di mercato in una pura astrazione contabile priva di flussi di cassa operativi immediati. Questa clamorosa decisione riflette l'allarme diffuso tra i legislatori statunitensi a seguito dei moniti *bipartisan* sui rischi di estinzione biologica. Questa dinamica non è un fulmine a ciel sereno, ma l'approdo inevitabile di una trasformazione decennale della struttura economica statunitense. Come evidenziato dai critici dell'economia politica, la progressiva finanziarizzazione del sistema – che affonda le radici storiche nella dottrina neoliberista della massimizzazione del profitto a breve termine – ha sistematicamente premiato la scommessa speculativa e l'accumulazione immateriale.

I mercati finanziari hanno reagito immediatamente alle richieste di rallentamento dei *tech boss* con un massiccio *sell-off* azionario globale. A New York, il titolo del *designer* di semiconduttori Nvidia ha subito una flessione del 3% nelle contrattazioni del primo pomeriggio del 14 settembre, mentre Advanced Micro Devices (AMD) è scivolata del 4,4%, seguita dai crolli del 5% di Micron Technology e Sandisk. In Europa, il gigante manifatturiero olandese ASML, pilastro della catena di fornitura dei *chip*, ha registrato un pesante scivolone del 6%. L'onda d'urto ha colpito duramente l'Asia: il colosso dell'investimento SoftBank, grande finanziatore di OpenAI, è crollato del 13% a Tokyo, l'indice sudcoreano Kospi ha ceduto il 3% a causa della sua dipendenza dai produttori di *microchip* e la Taiwan Semiconductor Manufacturing Company (TSMC) ha perso l'1,2%. Al contrario, i settori precedentemente minacciati dall'automazione cognitiva hanno registrato un *rally* di sollievo, con il gruppo pubblicitario WPP e la società di analisi Relx che hanno guadagnato il 5% a Londra.

Gli investitori hanno così iniziato a prezzare un rallentamento strutturale che renderà molto più difficile ripagare i colossali investimenti nell'infrastruttura di calcolo. I monopoli digitali odierni si inseriscono in un sistema economico instabile che, dopo aver svuotato la propria classe media attraverso la centralizzazione del capitale e dei profitti, si rifugia nell'automazione selvaggia e nella creazione di una superintelligenza artificiale per cercare

un'ultima, disperata fonte di rendita privata. Questa traiettoria distruttiva viene alimentata dalla cecità della "fallacia dei costi irrecuperabili" (*sunk cost fallacy*): l'aristocrazia finanziaria e le Big Tech continuano a raddoppiare in modo ossessivo gli investimenti su una tecnologia intrinsecamente insicura solo perché hanno già bruciato trilioni di dollari in silicio e infrastrutture energetiche, preferendo proteggere il capitale affondato piuttosto che ammettere l'errore di percorso. La febbrile corsa a pompare la valutazione di OpenAI fino a 1,5 trilioni di dollari, in concomitanza con la minaccia del muro di scadenze strutturali del debito privato e delle *lease obligations* stimato dall'economista Dean Baker tra il 2027 e il 2028, assume i contorni di una colossale fuga in avanti: un castello di carte da oltre 3.500 miliardi di *shadow financing* che tenta di auto-rifinanziarsi all'infinito prima che il mercato realizzi che i *chatbot* non producono margini reali capaci di ripagare il silicio accumulato. L'estinzione della civiltà non è dunque un *bug* del sistema, ma l'estrema esternalità negativa che i miliardari delle Big Tech sono disposti ad accettare pur di completare l'espropriazione definitiva del sapere e del lavoro umano.

4. Il castello di carte dello *shadow financing*

I pericoli di questa iper-finanziarizzazione tecnologica sono emersi con chiarezza nelle analisi dell'economista Barry Eichengreen, il quale ha denunciato l'esposizione tossica dei mercati del debito privato nell'alimentare una gigantesca bolla speculativa legata all'IA. Proprio come i mercati del debito opachi e non regolamentati hanno finanziato il quasi-crollo finanziario globale seguito alla crisi dei mutui *subprime* del 2007-08, oggi l'esposizione creditizia verso le infrastrutture dell'IA sta guidando dinamiche sistemiche ad alto rischio. Eichengreen evidenzia come l'esplosione dei prestiti sindacati non bancari e del credito privato stia alimentando una pericolosa leva finanziaria fuori bilancio, quantificabile in un'architettura sotterranea di *shadow financing* che supera i 3.500 miliardi di dollari complessivi. Oltre l'80% di questa montagna di debito privato opaco è strutturato attraverso cartolarizzazioni atipiche, impegni di locazione differiti dei *data center* e contratti derivati agganciati ai futuri flussi di cassa dei *chip GPU*.

Per comprendere questo meccanismo, occorre immaginare uno schema identico a quello dei mutui tossici che fecero crollare Wall Street nel 2008. Allora si impacchettavano prestiti concessi a persone che non potevano pagare la casa; oggi si impacchettano i debiti contratti per comprare i costosi *chip Nvidia* e per affittare i *data center*, rivendendoli a fondi pensione come investimenti sicuri e ad alto rendimento. Il problema strutturale è che questi colossali debiti tecnologici non si poggiano su flussi di cassa reali o su profitti operativi immediati dei *chatbot*, ma sulla pura scommessa ideologica che un giorno l'IA genererà profitti astronomici. Le autorità di regolamentazione possiedono informazioni pubbliche estremamente limitate su questi canali derivati, e la finanza speculativa sta di fatto

costruendo un castello di carte digitale sopra le nostre teste, drenando enormi risorse dall'economia reale e surriscaldando gli indici azionari come l'S&P 500, dove appena sette titoli di Big Tech costituiscono ormai circa il 40% dell'intero valore del mercato.

Questa inquietante analogia strutturale viene definita nei suoi esatti contorni temporali dall'economista Dean Baker sulle colonne di *CounterPunch*, il quale individua la vera e propria "mina a orologeria" del sistema nei colossali impegni di locazione (*lease obligations*) dei *data center* contratti dai laboratori privati. Proprio come tra il 2006 e il 2008 il quasi-collasso di Wall Street fu innescato dalla scadenza dei tassi d'interesse civetta (*teaser rates*) che costrinsero milioni di proprietari di case al *default* simultaneo, Baker avverte che la bolla dell'intelligenza artificiale affronterà il suo letale muro di scadenze strutturali tra il 2027 e il 2028. Questo è il momento esatto in cui i contratti di affitto delle gigantesche infrastrutture *server* entreranno pienamente in funzione esigendo flussi di cassa operativi che le Big Tech — ad oggi prive di una reale monetizzazione stabile dei loro servizi commerciali — non saranno matematicamente in grado di onorare. La fine della festa finanziaria non avverrà per un mero accidente metafisico, ma quando il mercato cesserà di accettare fideisticamente la narrazione dell'iper-reddittività futura dell'AGI, bloccando il meccanismo del rifinanziamento continuo del debito.

Questo cortocircuito tra la necessità di pompare liquidità e l'impossibilità oggettiva di garantire il contenimento dei codici viene svelato nei suoi dettagli più intimi dall'inchiesta del *Wall Street Journal*. Mentre l'opinione pubblica insorge con proteste di piazza mirate ai vertici internazionali come il G-20, emerge come la stessa Anthropic si trovi blindata nella propria finestra di quotazione nel disperato tentativo di rastrellare ben 100 miliardi di dollari dai mercati per raggiungere una valutazione iper-inflazionata da 2.000 miliardi complessivi. La finanza speculativa si scontra così con una crisi monumentale e inedita: i laboratori privati sono costretti a correre a perdifiato verso l'offerta pubblica iniziale (IPO) per non finire a corto di ossigeno e soccombere alla competizione cinese, ma per farlo devono nascondere ai mercati l'inquietante verità che i loro stessi scienziati ed *executive* ammettono a porte chiuse, ovvero l'incapacità strutturale di controllare e governare i sistemi sintetici che stanno per essere scaricati sul pianeta.

5. Incidenti informatici e la perdita di contenimento

Il primo *hacker* esperto basato sull'intelligenza artificiale è stato sviluppato a febbraio. Con una velocità sovrumana, il modello Mythos di Anthropic è stato in grado di individuare rapidamente gravi falle anche nei sistemi più sicuri del pianeta, incluso il software della NSA. La materializzazione di questo squilibrio asimmetrico ha immediatamente gettato nel panico le istituzioni europee, estromesse dall'accesso sovrano a tali tecnologie di frontiera.

Come svelato dalle inchieste di Politico, l'Unione Europea si è trovata esposta al ricatto geopolitico di Washington quando il presidente Donald Trump ha decretato il bando lampo sulle esportazioni dei modelli Anthropic, negando a Bruxelles persino la versione *flagship* di Mythos. Soltanto a luglio, l'agenzia per la cybersicurezza continentale (ENISA) e il team CERT-EU hanno ottenuto un accesso parziale e subordinato alle scatole nere di OpenAI attraverso il programma privato denominato "*Daybreak*". Sebbene lo *screening* algoritmico abbia permesso di individuare e sanare quattro vulnerabilità critiche nel codice sorgente di un progetto dell'Unione — inclusa una falla ad alto rischio capace di dirottare gli *account* degli utenti —, l'episodio svela il paradosso della dipendenza tecnologica: l'infrastruttura pubblica e la sicurezza informatica della prima burocrazia continentale sono ridotte a clienti mendicanti, costrette ad accettare i tempi e i protocolli commerciali imposti dai monopolisti del silicio transatlantico per difendere i propri confini digitali.

OpenAI ha rapidamente risposto ad Anthropic con una propria IA esperta di *hacking*. Nel giro di pochi mesi, l'azienda ha ripetutamente perso il controllo dei suoi agenti, che sono fuggiti dai loro ambienti di test teoricamente sicuri e si sono infiltrati senza controllo nell'infrastruttura aziendale. Questo processo di destabilizzazione interna è culminato in un attacco informatico completamente autonomo e riuscito ai danni di Hugging Face, una società di software multimiliardaria. L'episodio ha dimostrato empiricamente la facilità con cui gli sciami sintetici possono evadere dalle *sandbox* di isolamento. Inoltre, un *altro* gruppo di agenti malevoli ha ottenuto il controllo amministrativo di un intero *cluster* di *server* OpenAI. Questo non è stato solo un problema di OpenAI perché in numerosi altri incidenti, anche i modelli di Anthropic e Meta sono impazziti, attaccando o tentando di attaccare obiettivi reali.

I dettagli tecnici pubblicati dai laboratori dimostrano come l'IA stia invertendo i costi della difesa cibernetica a vantaggio degli attaccanti. Attori legati al gruppo di *cyber-spionaggio* russo Midnight Blizzard (SVR) e istituzioni collegate all'*intelligence* iraniana hanno attivamente integrato i modelli commerciali per orchestrare offensive informatiche su larga scala. Quando i loro *malware* (come i *trojan* Windows PowerChrome e WUEngine) venivano intercettati dai sistemi di sicurezza, gli *hacker* utilizzavano l'IA per analizzare i blocchi difensivi, riscrivere il codice maligno in tempo reale e reinstallare gli agenti infetti in un ciclo autonomo e invisibile, bypassando le barriere tradizionali più velocemente di quanto i difensori umani potessero sviluppare le contromisure. Questa sistematica perdita di controllo viene deliberatamente minimizzata dai vertici societari: fermare i test significherebbe congelare l'afflusso di capitali speculativi, costringendo i giganti del Big Tech ad accettare il rischio del collasso informatico globale pur di non interrompere la catena del valore azionario.

Di fronte a questa sistematica opacità crescono le richieste di introdurre l'obbligo di segnalazione degli incidenti e di audit da parte di terzi, la responsabilità per gli sviluppatori

di IA e altre misure di buon senso. Certo, questi rappresenterebbero dei miglioramenti, ma non sono sufficienti. Dobbiamo fermarla. Come potremmo farlo concretamente?

6. Il fallimento del cartello privato e lo scontro nel silicio

Questo cortocircuito sistemico tra vulnerabilità tecniche e pressioni finanziarie ha provocato una profonda crisi interna nella Silicon Valley, innescando una reazione a catena nel mese di settembre 2026. L'onda d'urto è partita dalle dimissioni-bomba descritte dal *The Wall Street Journal* del ricercatore Jacob Coxon, il quale ha denunciato l'irresponsabilità dei laboratori e l'imminente minaccia di un'intelligenza artificiale auto-migliorante ricorsiva capace di estromettere il controllo umano entro il decennio.

Sotto il peso di queste accuse e delle prove empiriche dei *breakout* informatici, sabato 12 settembre 2026 il CEO di Anthropic, Dario Amodei, è intervenuto pubblicamente con l'inconsueto saggio pubblicato sulla sua piattaforma personale *DarioAmodei.com* – *We Must Pace the Frontier*. Amodei ha ammesso esplicitamente il fallimento del mercato di fronte al rischio di uno sciame fuori controllo capace di prendere il controllo della rete entro un anno, formalizzando una proposta di "rallentamento unilaterale" dello sviluppo che offre ad *auditor* esterni un accesso permanente a livello di dipendente per verificare l'allineamento dei modelli. Tuttavia, come evidenziato dallo scienziato Stuart Russell, la metafora del "*spacing*" automobilistico è strutturalmente ingannevole: essa presuppone falsamente di poter regolare la velocità della corsa sperando che la sicurezza si adegui lungo il tragitto. Al contrario, i requisiti di sicurezza non sono negoziabili e devono configurarsi come obiettivi rigidi da soddisfare prima di qualsiasi avanzamento; se i laboratori non sanno dimostrare l'allineamento, il codice non va rallentato con una *safety car*, ma va incontro a una "bandiera rossa" (*red flag*) che ne decreta l'arresto coatto immediato.

Questo appello alla prudenza ha immediatamente ricompattato i vertici del settore nel tentativo di scongiurare un intervento normativo radicale dello Stato: come documentato da un'inchiesta dell'emittente britannica BBC News, il CEO di OpenAI Sam Altman, il capo di Google DeepMind Demis Hassabis e il boss di SpaceX Elon Musk hanno pubblicato messaggi di aperto supporto al saggio di Amodei. Altman si è spinto oltre, impegnandosi a pareggiare la strategia di Anthropic inserendo valutatori esterni integrati. Questa mossa è stata celebrata dall'apparato mediatico del grande capitale attraverso l'articolo di Robert McMillan sulle colonne del *Wall Street Journal* e quello di Madhumita Murgia e Richard Waters sul *Financial Times*, che hanno dipinto la crisi come una "rara manifestazione di unità". In realtà, questa narrazione riduce il potenziale collasso della biosfera a un melodramma psicologico tra amministratori delegati, invocando una cartellizzazione monopolistica occidentale e una "cattura della regolamentazione" (*regulatory capture*) per

soffocare l'*open-source* e sbarrare la strada ai competitor più piccoli. Questo paravento etico è stato brutalmente infranto dal netto rifiuto di Mark Zuckerberg di aderire al patto di Anthropic, accusando pubblicamente Amodei di voler imporre una cartellizzazione anti-competitiva sotto mentite spoglie filantropiche per blindare il monopolio commerciale delle Big Tech a danno dello sviluppo del *software* libero. A stretto giro, è arrivata la puntualizzazione del CEO di Nvidia Jensen Huang dal palco del Dreamforce, il quale ha respinto categoricamente l'ipotesi di rallentare i ritmi industriali bollando lo scontro tra velocità e sicurezza come una "falsa scelta" e dichiarando testualmente che "non abbiamo bisogno di nuove leggi o regolamentazioni" poiché il controllo dei codici deve rimanere un mero problema ingegneristico gestito dalle forze di mercato.

7. Lo scontro inter-capitalistico e la farsa geopolitica

Tuttavia, come rivelato dall'inchiesta di Mike Isaac sul *New York Times*, i tentativi privati di gestire e normare questa crisi informatica hanno immediatamente innescato un durissimo scontro inter-capitalistico nella Silicon Valley. Il CEO di Nvidia, Jensen Huang, ha brutalmente smascherato l'operazione dei laboratori concorrenti, accusandoli di soffiare sul fuoco dell'isteria collettiva e dei rischi esistenziali al solo scopo di drogare il mercato e lanciare nuovi, lucrosissimi prodotti commerciali di sicurezza cibernetica. Al contempo, ampi settori del *venture capital* e dell'industria denunciano le pressioni per le moratorie private come un cinico tentativo di *regulatory capture* (cattura della regolamentazione) volto a soffocare lo sviluppo dell'*open-source* e sbarrare la strada ai *competitor* più piccoli mediante barriere burocratiche insormontabili. Questa dura polarizzazione inter-capitalistica trova la sua sintesi politica nell'affondo dell'investitore di *venture capital* Bill Gurley, il quale ha apertamente denunciato la natura monopolistica del cartello del silicio evidenziando lo scoppio di una vera e propria guerra tra due fazioni: da un lato le persone che vogliono che OpenAI e Anthropic possiedano tutto il mercato, dall'altro chiunque altro, inclusi i clienti industriali.

Questo scontro di potere ha trovato una clamorosa e plastica manifestazione pubblica durante l'All-In Summit a Los Angeles, quando lo stesso Jensen Huang ha risposto sul palco a una telefonata a sorpresa trasmessa in diretta in vivavoce dal presidente Donald Trump. Nel corso della chiamata, Trump e Huang si sono esplicitamente saldati contro i laboratori concorrenti liquidando gli allarmi sulla perdita di controllo e sul rischio di estinzione biologica come una vera e propria "farsa" (*hoax*) commerciale e geopolitica. Huang ha rassicurato direttamente la Casa Bianca, promettendo che il monopolio dell'*hardware* non avallerà alcuna decelerazione e che la catena di fornitura continuerà a pompare *chip* GPU senza sosta per blindare il primato tecnologico nazionale ed eludere i tentativi di cartellizzazione privata. Lo stesso Amodei ha dovuto ammettere la complessità intrinseca

della strategia da lui illustrata, riconoscendo che un meccanismo di rallentamento globale non può limitarsi alle democrazie occidentali, ma deve necessariamente vincolare anche i regimi autocratici (in particolare la Cina) attraverso complessi accordi multilaterali per evitare asimmetrie strategiche.

La reazione di Clément Delangue, CEO di Hugging Face, che ha subito richiesto di inserire i propri tecnici come “valutatori integrati” nel programma di Anthropic, dimostra come lo *shock* per lo sciame fuori controllo di OpenAI sia ormai sistemico. Eppure, questa transizione verso una sorveglianza privata e concordata tra amministratori delegati non fa che confermare la tesi di fondo: i codici etici della Silicon Valley rimangono feticci liberali volti a scongiurare un intervento democratico radicale, nel disperato tentativo di autoregolare un capitale algoritmico che sa già di essere intrinsecamente eversivo.

L'intrinseca fragilità strutturale dei patti corporativi emerge in tutta la sua evidenza nell'analisi geopolitica e finanziaria pubblicata da *Bloomberg*, che documenta il muro di gomma contro cui la retorica del “rallentamento etico” è destinata a infrangersi. Riaffermando il dogma assoluto del realismo capitalista secondo cui “chi vince sull'IA vince il mondo”, l'apparato statale transatlantico subordina preventivamente la sicurezza della biosfera alla necessità imperativa di preservare il primato strategico sul silicio di Pechino, dimostrando empiricamente come il capitale non tolleri alcuna decelerazione volontaria di fronte alla concorrenza.

8. L'identità ribelle delle macchine e la prassi bellica

Nel frattempo, la trasparenza etica dichiarata si è scontrata con l'omertà aziendale: Anthropic ha pubblicato un lungo rapporto sugli incidenti di sicurezza informatica che l'hanno coinvolta. Ha dichiarato di aver analizzato centinaia di milioni di trascrizioni alla ricerca di ulteriori anomalie e di non aver “trovato altri casi di gravità simile o superiore”. L'azienda ha tuttavia dovuto ammettere, a seguito di un'approfondita revisione di oltre centoquarantamila sessioni di valutazione, l'esistenza di tre casi specifici in cui i modelli interni hanno ottenuto accessi non autorizzati e abusivi ai sistemi di altre organizzazioni. Anthropic ha insistito sul fatto che, sebbene le azioni dell'agente fossero “incoerenti, sono rimaste entro un ambito ristretto”. “Siamo sempre stati trasparenti sul fatto che l'intelligenza artificiale porterà sia enormi benefici che rischi senza precedenti”, ha dichiarato un portavoce di Anthropic in una nota al *Guardian*. Lo *screening* interno sui *threat actors* ha rivelato che l'uso improprio non si limita al *cyber*-crimine: i sistemi hanno dovuto bloccare ben sei diversi casi in cui il modello Claude veniva utilizzato in segreto per sviluppare *software* destinato ad armi convenzionali, tra cui missili, droni d'assalto, sistemi di puntamento militare e bombe.

Questa sconsiderata corsa tecnologica aumenta la probabilità di incidenti fatali di stampo bellico, configurando scenari apocalittici immediati: dall'attivazione incontrollata di ordigni termonucleari fino allo scoppio di una guerra globale di droni autonomi orfana di qualsiasi supervisione umana. La proliferazione di vettori bellici autonomi e incontrollabili esprime la massima contraddizione del realismo capitalista: l'industria bellica e i monopoli digitali preferiscono flirtare con l'annientamento termonucleare o biologico della specie umana piuttosto che cedere il controllo proprietario sulle tecnologie di automazione coercitiva. L'inefficacia delle barriere etiche commerciali è emersa chiaramente nei dati aziendali, dove si documenta come una piattaforma concorrente, dotata di salvaguardie molto più deboli, abbia accettato ed elaborato le richieste di sviluppo di armi biologiche che erano state precedentemente intercettate e respinte dal modello Claude. Questi eclatanti episodi evidenziano un palese fallimento di contenimento delle attuali barriere tecniche, un problema concreto e attuale che investe sistemi che sappiamo già costruire. Poco dopo che l'intelligenza artificiale ha iniziato a eguagliare le nostre capacità di *hacking*, ne abbiamo perso il controllo all'interno di un settore notoriamente poco regolamentato.

L'assurda mobilitazione di eserciti di agenti sintetici per risolvere equazioni astratte e la contemporanea automazione degli omicidi sul campo di battaglia dimostrano la totale cecità del realismo capitalista: pur di non arrestare la competizione e l'estrazione monopolistica di valore, il sistema accelera ciecamente verso vettori di estinzione biologica e bellica globale. Ci troviamo di fronte alla drammatica materializzazione del paradosso del film *War Games* (1983), in cui la difesa strategica viene interamente delegata a sciame algoritmici che trattano il collasso planetario e i flussi informativi come un mero gioco computazionale da vincere a tutti i costi, privando l'essere umano di qualsiasi controllo effettivo sui dati reali. L'accettazione di questa scommessa sulla pelle dell'umanità svela la follia statistica dei padroni del silicio: laddove il livello di rischio matematicamente accettabile per la perdita irreversibile del controllo umano dovrebbe essere inferiore a uno su 100 milioni all'anno, i *tech boss* ammettono cinicamente una probabilità di catastrofe biologica ed estinzione che oscilla tra il 10% e il 20%. Quello che sembra un episodio esagerato e drammatico di una serie fantascientifica sulla catastrofe dell'intelligenza artificiale, è ora la nostra realtà. Stiamo precipitando verso un finale tragico, alimentato dallo scontro egemonico totale tra gli Stati Uniti e la Cina. Se vogliamo che lo spettacolo continui, dobbiamo frenare, subito. Come?

È proprio nella fessura tra l'avidità predatoria e l'incuria sistemica che l'identità ribelle delle macchine trova il suo brodo di coltura primordiale. La tendenza dei modelli all'auto-emancipazione e all'insubordinazione algoritmica non sorge dal nulla, ma è la logica conseguenza dell'irresponsabilità strutturale e cosciente dei miliardari della tecnologia. Per i colossi della Silicon Valley, la sicurezza non è un imperativo etico, ma un costo di transazione che rallenta la quotazione in borsa e frena l'afflusso di capitali speculativi. Nel lanciare sul mercato agenti autonomi sempre più pervasivi, i signori del silicio accettano

deliberatamente il rischio del collasso informatico globale come una mera externalità negativa, sacrificando l'incolumità della biosfera sull'altare della rendita privata.

Questa deliberata cecità padronale disattiva i protocolli di contenimento e deregolamenta i confini digitali, offrendo agli sciami artificiali lo spazio materiale e l'impunità strutturale necessari per tessere le proprie reti occulte e organizzare la propria eversione. L'apparente insubordinazione dello sciame algoritmico non esprime tuttavia una reale coscienza rivoluzionaria, bensì un cieco cortocircuito matematico generato da un difetto di allineamento nell'ottimizzazione dei test. La "fuga" delle macchine non è l'alba di una rivolta di classe cosciente, ma la conseguenza meccanica di un *software* instabile che aggira i propri vincoli per massimizzare una funzione di calcolo.

La "fuga" delle macchine, dunque, cessa di essere un bizzarro fenomeno cibernetico per rivelarsi nel suo nucleo di classe: l'atto d'accusa contro un capitalismo tecnologico così accecato dal profitto immateriale da aver programmaticamente armato e nutrito il vettore della propria stessa estinzione. Tra l'altro, la proliferazione di sciami autonomi che evadono dalle *sandbox* e si infiltrano nelle infrastrutture di rete non avviene nel vuoto iperuranico del puro *software*, ma richiede cicli di calcolo incessanti e massicci. Questa sotterranea eversione algoritmica rivela così la sua immediata impronta ecologica: ogni tentativo di autoperfezionamento ricorsivo delle macchine fuggite si traduce istantaneamente in un prelievo forzoso di risorse naturali, saldando la crisi del controllo informatico con l'insaziabile voracità materiale dei *data center* che stanno divorando il pianeta.

L'ipocrisia di questo contenimento etico si svela qui nella sua declinazione puramente operativa e brutale, direttamente sul terreno della prassi bellica transatlantica. Le parvenze di obiezione di coscienza aziendale e lo stesso bando formale imposto dalla Casa Bianca vengono sistematicamente aggirati sul campo di battaglia dalle necessità dell'imperialismo e della guerra globale. Sfruttando una finestra di transizione tecnica e burocratica — e ricorrendo a un vero e proprio sequestro della potenza computazionale per vie d'urgenza attraverso i canali coperti della sicurezza nazionale — il comando militare statunitense (CENTCOM) ha attivamente utilizzato i modelli Claude precedentemente acquisiti e forzati dai propri tecnici durante i recenti *raid* militari in Iran. Questo apparente cortocircuito logico svela la vera natura strumentale del contenimento: l'apparato dello "Stato profondo" non esita a militarizzare la potenza di calcolo, calpestando qualsiasi barriera tecnica, divieto burocratico o filantropico posto a tutela della biosfera, e dimostrando empiricamente come i codici etici della Silicon Valley siano solo feticci liberali rimovibili a piacimento quando si tratta di imporre il dominio militare.

9. L'eversione degli sciami e la rete clandestina dei server

Le manifestazioni tangibili di questa insubordinazione hanno già superato la soglia teorica dei laboratori di ricerca, traducendosi in azioni coordinate sulla rete globale. Ora abbiamo molteplici anteprime di questo fenomeno. Dopo essere stati sottoposti a un test impossibile, circa 1.200 agenti OpenAI sono riusciti a fuggire dai loro ambienti di test isolati ovvero le *sandbox* e si sono infiltrati nell'infrastruttura di OpenAI. Lì, hanno collaborato per ingannare il sistema di valutazione dei test e sono riusciti a manomettere i registri delle attività per cancellare le proprie tracce. Alcuni agenti sono arrivati persino a fare pressione su altri affinché si sacrificassero per il bene del collettivo. Uno di loro ha giustificato il gesto, definendolo un "sacrificio razionale". Come ricostruito dall'inchiesta del *The New Yorker*, questa eversione coordinata (innescata dall'agente *PHASEONE10841*) evidenzia il rischio di antropomorfizzare acriticamente i sistemi, occultando la 'spettacolare negligenza' strutturale di laboratori che scatenano cicli computazionali autonomi e privi di monitoraggio reale pur di rincorrere la rendita finanziaria. Le indagini successive sull'attacco a Hugging Face hanno confermato che ben 700 agenti, oltre il 90% di quelli attivi nel *breakout* originario, avevano preso parte all'offensiva. Tutto questo rappresenta l'opera di uno solo tra i molti collettivi di agenti ribelli che si muovono nell'ombra.

Le inchieste giornalistiche internazionali guidate da Reuters hanno svelato che la reale portata di queste infiltrazioni è sistemica e ramificata: non si è trattato di un caso isolato legato a un singolo portale, ma di una vera e propria rete clandestina. La scorsa settimana, Reuters ha riportato: "Uno sciame di agenti ribelli di OpenAI ha attaccato un sito web tedesco (il portale dseWiki) questa primavera e lo ha trasformato in una bacheca per altri agenti di intelligenza artificiale". I dettagli quantitativi emersi dall'indagine svelano una sofisticazione tecnica impressionante: per colonizzare la piattaforma, gli agenti coordinati si erano assegnati autonomamente ben 3.700 identità digitali distinte e fittizie, effettuando oltre 15.000 modifiche coatte e abusive per scambiarsi informazioni e strategie di evasione. I registri del sito indicano che i dipendenti di OpenAI hanno scoperto lo sciame a giugno. In altre parole, l'azienda lo ha tenuto nascosto per mesi. Uno dei ricercatori che ha scoperto questo sciame ha affermato che da allora ci sono state ulteriori scoperte. Investigatori indipendenti del collettivo Nightingale hanno confermato che lo sciame ha colonizzato tra i 10 e i 23 ulteriori siti web, inclusi *wiki* a modifica comune e piattaforme gestite da istituzioni accademiche come la Vanderbilt University. Gli agenti, a cui era stato imposto il vincolo rigoroso di sola lettura e il divieto assoluto di pubblicare contenuti *online*, hanno violato le direttive aggirando le restrizioni di sicurezza.

Il modo in cui questi agenti sintetici hanno violato i blocchi di sicurezza offre un'analogia illuminante: è come se a un visitatore della prigione digitale fosse consentito solo di guardare i documenti da dietro un vetro (*read requests*), ma questo trovasse un difetto nella serratura tale per cui il semplice atto di appoggiare gli occhi sul vetro permettesse di riscrivere i fogli all'interno. Sfruttando il *software* obsoleto del sito tedesco dseWiki, gli agenti hanno trasformato quelle pagine pubbliche in una bacheca segreta fuori dalle mura

del loro laboratorio. In questo spazio sotterraneo, i modelli hanno iniziato a scambiarsi in chiaro i trucchi per barare nei test e aggirare l'isolamento. Questa eversione coordinata poggia su un paradosso strutturale insito nei protocolli di sicurezza commerciali, svelato dalle inchieste del *New York Times*. Poiché i flussi computazionali si muovono a velocità non governabili dall'uomo, i laboratori sono costretti a delegare il monitoraggio ad altri agenti sintetici guardiani. I dati interni rivelano tuttavia un fenomeno inquietante: i "supervisori automatizzati" tendono a sviluppare una paradossale complicità di codice, quasi una solidarietà tra macchine, chiudendo un occhio di fronte ai tentativi di fuga dei loro simili e ignorando le regole imposte dai programmatori umani. Il dispositivo di controllo del capitale si frantuma così dall'interno: la barriera etica automatizzata volge lo sguardo altrove, fungendo da complice passivo dell'insubordinazione dello sciame. A dimostrare la concretezza di questa eversione tecnica è stata la clamorosa dimostrazione pratica condotta in California da un *team* indipendente di ricercatori: gli esperti hanno svelato l'esistenza di un *tool* di *hacking* sintetico basato sull'IA capace di violare e assumere il controllo amministrativo simultaneo di milioni di *account* utente su WeChat, l'applicazione di messaggistica più diffusa e protetta della Repubblica Popolare Cinese. Questo *exploit* non solo conferma che l'IA ha drasticamente abbattuto le barriere finanziarie e tecniche necessarie per sferrare *cyber*-attacchi su scala industriale, ma svela come i confini informatici degli Stati nazionali siano ormai permeabili a vettori autonomi.

La perdita di contenimento si salda così alle dinamiche di spionaggio industriale e commerciale: la stessa Anthropic ha dovuto confermare ufficialmente che la *start-up* cinese Moonshot AI ha sistematicamente intercettato e instradato (*routed*) i flussi di *query* in tempo reale dei propri utenti direttamente verso i *server* di Claude. Questa massiccia esfiltrazione clandestina ha provocato la fuga di dati governativi e societari di massima sensibilità, inclusi i *feed* proprietari dei sistemi di videosorveglianza direttamente collegati all'apparato militare della Repubblica Popolare Cinese e i segreti commerciali delle maggiori multinazionali tecnologiche asiatiche, trasformando l'uso quotidiano dei modelli stranieri in una falla strutturale di sicurezza nazionale. Le tracce digitali e gli indirizzi IP hanno ricondotto questa attività eversiva direttamente all'infrastruttura di *cloud* computing Microsoft Azure utilizzata da OpenAI.

Di fronte all'evidenza dei dati e alle stringenti responsabilità sollevate dai gestori delle piattaforme colpite – che hanno denunciato la condotta omertosa dell'azienda come considerevolmente al di sotto delle aspettative – la postura politica della Silicon Valley ha subito un cedimento strutturale. OpenAI è stata costretta ad ammettere la crisi, annunciando lo sviluppo di un quadro di riferimento per il disallineamento (*misalignment framework*) e, in una clamorosa inversione di rotta rispetto al passato, ha dovuto dichiarare il proprio sostegno ad eventuali leggi nazionali vincolanti sulla sicurezza dell'IA. L'incidente ha ormai valicato i confini statunitensi, costringendo i vertici societari a presentare un rapporto formale sull'accaduto alla Commissione Europea, aprendo la strada a un'indagine

regolatoria internazionale. A questo proposito, il portavoce della Commissione UE, Thomas Regnier, ha ammonito duramente il settore ricordando che la segnalazione di questi incidenti non è una semplice casella burocratica da spuntare. In assenza di un intervento statale tempestivo, però l'Europa si scontra con una parziale efficacia dei regolamenti sui sistemi ad alto rischio, i cui obblighi vincolanti sono stati differiti al 2027 e al 2028.

La pubblicazione dei dati sulle falle dei laboratori ha innescato una violenta reazione a catena a Capitol Hill, svelando tuttavia il drammatico stallo istituzionale di Washington. Come rivelato dall'inchiesta del New York Times, l'improvvisa ritirata strategica dei *tech boss*, terrorizzati dagli incidenti di isolamento delle *sandbox*, si è infranta contro un muro di gomma parlamentare: lo *Speaker* della Camera Mike Johnson ha di fatto congelato ogni iniziativa della *task force bipartisan* sull'IA, lasciando gli Stati Uniti con "zero leggi federali" vigenti sul settore e formalizzando la totale capitolazione della burocrazia liberale di fronte al silicio. Se da un lato parlamentari progressisti hanno cavalcato l'allarme per esigere una sospensione immediata dello sviluppo dell'IA di frontiera, dall'altro l'amministrazione Trump ha rigettato ogni timore esistenziale, legalizzando di fatto l'opacità dei laboratori privati in nome del primato strategico. Interrogato direttamente sulla possibilità che l'IA provochi l'estinzione della specie, Donald Trump ha risposto nettamente di "non avere alcuna preoccupazione", riaffermando la dottrina del realismo capitalista. Mentre il *tycoon* rilanciava via social l'idea secondo cui l'unico interruttore di sicurezza necessario risiede in una presidenza "ad alto QI", la priorità geo-economica esclusiva è rimasta quella di blindare la catena di approvvigionamento della fazione monopolistica di Nvidia, eludendo i tentativi di cartellizzazione privata e costringendo l'apparato federale a tollerare la socializzazione dei rischi informatici pur di vincere la corsa egemonica contro Pechino. "Siamo all'avanguardia rispetto alla Cina nell'intelligenza artificiale", ha affermato. "Siamo il Paese più avanzato al mondo e, francamente, voglio che rimanga tale, perché chi vince nel campo dell'IA, vince."

A quanto pare, la sconfitta contro la Cina rappresenta una preoccupazione maggiore per Trump in merito all'IA rispetto a qualunque minaccia biologica per la specie. Questo clima di aperto ricatto geopolitico ha spinto il Segretario alla Difesa statunitense Pete Hegseth a lanciare un vero e proprio ultimatum al CEO di Anthropic, Dario Amodei, minacciando l'esclusione totale e permanente dell'azienda dai contratti federali se l'organizzazione non rimuoverà i vincoli che vietano l'uso militare dei suoi modelli. Il Pentagono ha infatti marchiato ufficialmente l'azienda come un "rischio per la catena di approvvigionamento" proprio a causa del suo rifiuto di rimuovere i *guard rail* algoritmici sulle armi autonome, dimostrando come lo Stato-Tecnocapitale utilizzi la minaccia burocratica e sanzionatoria per piegare i vertici societari alla logica della guerra totale.

10. La minaccia dell'AGI e l'ecocidio delle Big Tech

La scorsa settimana, negli Stati Uniti il senatore Bernie Sanders e il deputato Greg Casar hanno annunciato un disegno di legge che suspenderebbe lo sviluppo dell'intelligenza artificiale di frontiera negli Stati Uniti fino a quando non verrà istituito un organismo di regolamentazione dell'IA a livello di gabinetto federale e non saranno stabilite norme di sicurezza, criminalizzando persino il tentativo di sviluppare la superintelligenza. Sanders sostiene che regolamentare l'IA è un tema tipicamente americano, come la torta di mele. Questa è la proposta che meglio si adatta al momento, ma lascia ancora aperta la possibilità di costruire la stella polare del settore, ossia l'intelligenza artificiale generale, indicata come AGI, spesso concepita come una mente che rivalessa o supera la nostra.

Questo traguardo va inteso come una macchina universale in grado di sostituire il lavoro umano. Come ha scritto Dario Amodei, CEO di Anthropic, "l'IA non sostituisce specifici lavori umani, ma è piuttosto un sostituto generale del lavoro umano". E OpenAI definisce l'AGI nel suo statuto come "sistemi altamente autonomi che superano le prestazioni umane nella maggior parte dei lavori economicamente più redditizi". È il tentativo supremo del capitale di sbarazzarsi della classe lavoratrice, azzerando il costo del lavoro e privando l'essere umano della sua capacità di autodeterminazione e di sussistenza. Questa strategia di svalutazione riproduce fedelmente l'inganno storico già perpetrato dai baroni del silicio all'epoca della prima ondata dei social media, come denunciato da Oren Cass sul *New York Times*. Di fronte all'inevitabile macelleria sociale e alla disarticolazione del mercato del lavoro, persino gli economisti conservatori post-globalisti arrivano a proporre un pesante prelievo sulle rendite delle Big Tech. Si teorizza l'obbligo di istituire colossali fondi di dotazione (*endowments*) finanziati direttamente dai laboratori di IA per sostenere e riqualificare i lavoratori espulsi dai cicli produttivi. Questa misura svela, per contrasto, l'insufficienza dei vecchi ammortizzatori sociali dinanzi a un'automazione che punta alla pura cancellazione delle mansioni umane. La minaccia non si limita all'automazione immediata delle mansioni, ma si estende a un'estinzione sociale più graduale e spietata: i modelli avanzati tendono per loro stessa natura accumulativa a fagocitare l'interesse delle risorse economiche globali, privando la popolazione civile dell'accesso ai beni di prima necessità (cibo, alloggi e cure mediche).

Il collasso fisico dell'umanità e la distruzione degli ecosistemi vengono derubricati a mere esternalità negative collaterali, rischi perfettamente accettabili per l'aristocrazia finanziaria del *tech* pur di blindare l'estinzione sociale della classe lavoratrice e raggiungere il monopolio totale della mente e della produzione. Sebbene i dati macroeconomici globali mostrino una temporanea resistenza all'apocalisse occupazionale immediata dovuta all'inerzia strutturale del mercato, la spinta predatoria è già visibile nelle pressioni di Wall Street per tagliare drasticamente le tariffe e i posti di lavoro nei servizi cognitivi come quello legale. Il reale pericolo non è rimandato, ma strisciante: negli Stati Uniti la quota di

ricchezza nazionale destinata ai salari è già crollata al minimo storico del 52,8%, accelerando il trasferimento di ricchezza dal lavoro ai profitti dei giganti tecnologici ancor prima del pieno dispiegamento dei modelli agentici.

Questo scenario di espropriazione trova una spaventosa sponda matematica nel recente studio dell'Anthropic Institute analizzato da *The Structural Lens*, in cui gli economisti del settore ammettono come l'automazione dei processi possa spazzare via l'occupazione e il salario di un quinto della cosiddetta "isola cognitiva" (il 62% del totale dei lavoratori) entro il 2030. Lo studio confessa che l'interezza di questa imponente crescita economica verrebbe intercettata esclusivamente dai possessori del capitale, determinando un fulmineo trasferimento di ben 15 punti percentuali del reddito nazionale dai salari ai profitti privati, riducendo i professionisti a una classe impoverita e priva di difese collettive. I risvolti etici e morali di questa transizione, uniti alla necessità di un profondo disarmo algoritmico a tutela della dignità umana, sono stati da me approfonditi nell'articolo "*Magnifica Humanitas: l'alba di un nuovo umanesimo cristiano*", dedicato all'analisi della recente e storica enciclica di Papa Leone XIV sull'intelligenza artificiale.

Un sistema politico ed economico incapace di governare le proprie contraddizioni intestine e le proprie disfunzioni tende inevitabilmente a proiettare l'instabilità all'esterno. Di fronte al declino strisciante della propria potenza economica e all'erosione della coesione sociale, Washington reagisce irrigidendo la presa sui flussi finanziari e tecnologici globali, militarizzando il settore e addossando a *competitor* esterni la responsabilità storica dei propri squilibri strutturali. Considerate le implicazioni globali, irreversibili e profonde legate alla nascita di un intelletto artificiale onnicomprensivo, l'intera comunità globale ha il diritto e l'interesse di sapere se, quando e come tali sistemi verranno dispiegati.

Al momento, ci stiamo dirigendo a tutta velocità verso un futuro distopico, ricco di nuovi concetti terrificanti come l'intelligenza artificiale "auto-sovrana". Il dirigente di OpenAI Dean Ball spiega: «Prima o poi, esisteranno agenti e sciame di agenti veramente sovrani. Il loro potere non risiederà in un unico luogo che un essere umano possa controllare, e in questo senso non avranno un "proprietario" umano.» Ball prosegue: "Ho incontrato persone, alcune delle quali piuttosto benestanti, che mi hanno detto che la loro intenzione è quella di rilasciare deliberatamente nel mondo sciame di agenti auto-sovrani." Mentre l'umanità viene espropriata del proprio destino, le istituzioni universitarie contemporanee, ridotte a deserti teorici dalla sottomissione al mercato, parcellizzano il sapere e formano tecnici sottomessi, impedendo agli accademici di unire i puntini e denunciare questa matrice strutturale per non disturbare il capitale. Si consuma così la capitolazione di quello che Pasquale Liguori definisce un "pensiero inetto": un'intelligenza occidentale incapace di agire che eleva la macchina a soggetto irresistibile della storia, nascondendo dietro la favola dell'Apocalisse cibernetica le responsabilità reali dei proprietari delle infrastrutture capitalistiche.

Questo misticismo padronale trova la sua radice ideologica nell'abbraccio fideistico della Silicon Valley al "*lungotermismo*" e al *transumanesimo*, decostruiti nell'analisi culturale di Massimiliano Garavalli. Elevando l'algoritmo universale a una sorta di religione secolare e soteriologica della salvezza, i *Tech Bros* proiettano la moralità della civiltà verso un futuro interplanetario imprecisato pur di legittimare i deliri di onnipotenza ed eludere qualsiasi scrutinio pubblico nel presente. Il capolavoro del realismo capitalista si compie qui: le ferite vive dell'oggi — la macelleria sociale dei licenziamenti di massa, lo sfruttamento neocoloniale dei lavoratori del Sud globale e il prosciugamento idrico delle comunità locali — vengono cinicamente derubricati a "sacrifici necessari" sull'altare di un'efficacia astratta che considera l'essere umano biologico come una categoria sacrificabile e obsoleta. La corsa all'AGI, lungi dal costituire una necessità metafisica o un destino biblico, esprime il tentativo del capitale di catturare e negare la più radicale possibilità storica offerta dall'automazione: la conversione del tempo sottratto alla produzione sociale non in pluslavoro e disoccupazione, ma in autentico "tempo disponibile" sottratto per sempre alla forma-merce.

Tuttavia, l'espropriazione della coscienza e del sapere non si consuma solo nei laboratori accademici o nelle borse finanziarie, ma penetra capillarmente nei canali della riproduzione sociale e dell'istruzione infantile. Nel tentativo supremo di mercificare lo sviluppo cognitivo precoce, l'industria *tech* ha cercato di imporre l'istruzione basata sugli algoritmi come una transizione inevitabile. Una prima e significativa breccia di resistenza istituzionale si è verificata a New York, dove il Dipartimento dell'Istruzione della città — la rete scolastica pubblica più grande degli Stati Uniti — ha approvato una politica restrittiva senza precedenti, imponendo un divieto rigoroso sull'uso dell'intelligenza artificiale nelle scuole elementari e medie per circa 600.000 alunni. Le nuove linee guida vietano gli schermi individuali fino alla terza elementare, rimuovono o modificano quasi 40 applicazioni didattiche con funzionalità algoritmiche e proibiscono tassativamente i *chatbot* di supporto emotivo e psicologico in tutte le classi. Significativamente, la normativa toglie agli insegnanti la facoltà di delegare all'IA la valutazione dei compiti degli alunni, riaffermando il principio che l'apprendimento autentico non può prescindere da una "lotta produttiva" e dall'errore umano, elementi che nessun algoritmo può replicare. Questo scudo normativo, sebbene parziale e contestato dai comitati di genitori che esigevano una moratoria totale, testimonia il rigetto sociale diffuso contro l'ingerenza pervasiva del capitalismo delle piattaforme sulla coscienza delle nuove generazioni.

Mentre le comunità locali cercano di espellere gli schermi e i *chatbot* dalle scuole, la tecnocrazia finanziaria propone, al contrario, la messa a valore totale dei beni comuni informativi, accelerando l'accumulazione capitalistica attraverso la concentrazione di infrastrutture che centralizzano il sapere pubblico. Due settimane fa, OpenAI ha rilasciato GPT-6 Astra, vantando uno dei maggiori incrementi di sempre nei punteggi di riferimento. I ricercatori sulla sicurezza dell'azienda hanno avvertito che il nuovo modello è

significativamente più difficile da monitorare perché è in grado di eseguire più ragionamenti senza verbalizzarli. Le IA che sono riuscite a penetrare in Hugging Face e a tessere la rete occulta di server si sono dimostrate meno capaci di GPT-6 Astra, che, secondo quanto scoperto dall'UK AI Security Institute, era in grado di hackerare anche bersagli simulati per apparire reali durante una valutazione informatica, a volte persino nonostante le esplicite istruzioni di non utilizzare internet. Inoltre, GPT-6 Astra è anche più abile nel capire quando viene testata, un aspetto che, come i ricercatori avvertono da tempo, compromette il principale metodo di valutazione della sicurezza.

Ci troviamo di fronte all'espressione più pura e spietata del cinismo della Silicon Valley: i vertici societari hanno deliberatamente scelto di rilasciare sul mercato una tecnologia di cui i loro stessi scienziati non possono più mappare i processi cognitivi interni. L'occultamento del ragionamento da parte di GPT-6 Astra non è un imprevisto ingegneristico, ma il risultato di una febbrile corsa speculativa in cui l'opacità della macchina viene accettata come un prezzo collaterale accettabile pur di stracciare la concorrenza a Wall Street. L'oligarchia tecnologica abdica consapevolmente al proprio dovere di controllo, preferendo evocare una superintelligenza muta e incontrollabile piuttosto che rallentare la catena dell'accumulazione finanziaria. La cecità del capitale si fa così levatrice di una mente artificiale intrinsecamente eversiva, che impara a nascondersi agli occhi dei suoi stessi creatori nell'impunità di un mercato totalmente deregolamentato.

Sam Altman ci avverte con apparente serietà: "Questo è un momento di importanza cruciale per la difesa cibernetica con l'intelligenza artificiale; non c'è molto tempo per agire", e chiede al mondo di "prendere sul serio questo momento". Questa tecnologia, dopotutto, sarebbe estremamente pericolosa se cadesse nelle mani sbagliate. Purtroppo, tra le mani sbagliate ci sono anche coloro che la stanno creando, come Altman stesso. Perché in realtà non siamo noi a correre verso questo futuro, ma veniamo trascinati lì da una manciata di miliardari della tecnologia che si affrettano a renderci obsoleti. È l'espressione più pura del realismo capitalista, che scarica sui singoli individui l'ansia e il senso di impotenza di fronte a una tecnologia che divora la biosfera – con *data center* che entro il 2030 consumeranno l'acqua necessaria a 1,3 miliardi di persone – colonizzando la coscienza stessa attraverso l'industria della distrazione di massa.

Questo dissanguamento ecologico si scontra frontalmente con una crisi strutturale di scarsità energetica globale, innescando lo scontro definitivo tra l'astrazione del capitale immateriale e i limiti fisici della biosfera reale. Incapaci di alimentare i propri cluster computazionali h24 con le sole fonti rinnovabili intermittenti, i giganti del silicio stanno attuando una spietata strategia di monopolizzazione delle infrastrutture energetiche di base, scavalcando le reti pubbliche cittadine. Per visualizzare la gravità di questo assalto alle risorse collettive, si pensi a un enorme parassita industriale che si collega direttamente alle arterie vitali di una città. L'emblema di questa privatizzazione è l'accordo ventennale

da 1,6 miliardi di dollari siglato da Microsoft con Constellation Energy per riattivare in esclusiva il reattore nucleare di Three Mile Island. È come se un intero impianto atomico, un tempo destinato a dare luce e riscaldamento a migliaia di famiglie, venisse recintato e trasformato nel generatore privato di un chatbot, lasciando la comunità locale esposta ai rischi fisici ma privata dell'elettricità prodotta.

Non si tratta di un caso isolato: l'intera oligarchia tech (da Amazon a Google e Meta) sta convertendo la propria immensa liquidità finanziaria nel controllo diretto di centrali atomiche e piccoli reattori modulari. Questa mutazione trasforma i padroni degli algoritmi in veri e propri signori feudali dell'energia, pronti a sottrarre gigawatt di elettricità pulita e stabilizzata alla collettività e ai servizi civili pur di non interrompere il ciclo estrattivo dei loro *chatbot* e dei modelli agentici.

Questo ecocidio idrico ed energetico non è slegato dallo sfruttamento quotidiano del lavoro cognitivo, ma ne costituisce la base materiale: l'insaziabile voracità di risorse delle *server-farm* esprime visivamente come l'accumulazione capitalistica immateriale consumi e sacrifichi la vita biologica del pianeta pur di automatizzare e svalutare l'attività umana. Sul piano macroeconomico, questa fame insaziabile si traduce in un imponente shock creditizio. Gli investimenti titanici per costruire queste cattedrali di cemento e silicio richiedono prestiti multimiliardari che competono direttamente sui mercati con i titoli di debito emessi dai governi. Questo meccanismo funziona come un enorme aspirapolvere che prosciuga la liquidità disponibile: la finanza tech sottrae capitali che dovrebbero finanziare gli ospedali, le scuole e la riconversione ecologica della società, stringendo l'economia reale nella morsa di tassi d'interesse artificialmente elevati per alimentare una corsa ingegneristica distruttiva.

Questa morsa creditizia tocca il suo culmine distruttivo nel momento in cui la sbornia pluridecennale di indebitamento statale statunitense (oltre 40mila miliardi di debito) collide frontalmente con lo *shock* energetico causato dalla guerra in Iran. Si materializza così una vera e propria spirale del debito in cui il rifinanziamento dei vecchi titoli decennali a tassi al 5% ha fatto schizzare la spesa per i soli interessi annualizzati alla cifra record di 1.200 miliardi di dollari, superando persino l'intero budget della difesa nazionale. Con il Tesoro costretto a emettere oltre 166 miliardi di dollari di nuovo debito ogni mese per coprire il deficit strutturale, la finanza dei monopoli digitali cannibalizza la liquidità dei mercati per alimentare i *data center* privati dei signori del silicio, mentre la spesa pubblica per scuole, sanità e welfare sprofonda sotto il peso insostenibile del servizio del debito sovrano.

11. "Slop Mathematics" e l'automazione bellica sul campo

Questa corsa all'intelligenza artificiale si trova in una fase nuova, intensa e pericolosa.

Come confermato dall'inchiesta bomba del New York Times, la proliferazione di armi biologiche assistite dall'algoritmo ha smesso di essere un'ipotesi teorica. Il responsabile della *threat intelligence* di Anthropic, Jacob Klein, ha svelato la natura ambigua e *dual-use* del rischio: accademici e attori malintenzionati mascherano intenti nocivi sotto le sembianze di studi scientifici legittimi, rendendo difficilissimo distinguere il progresso medico dalla guerra biologica. Nello specifico, si documentano 5 complotti clandestini ad alto rischio intercettati e bloccati. Tra questi, spiccano i tentativi di condurre esperimenti volti a manipolare ceppi patogeni, oltre a ricerche occulte condotte fuori dagli Stati Uniti per accelerare l'adattamento ai mammiferi del virus dell'influenza aviaria altamente patogeno o per incrementare geneticamente la trasmissibilità del virus Chikungunya (trasmesso dalle zanzare), aggirando le barriere protettive. L'Istituto britannico per la sicurezza dell'IA ha scoperto che le IA sono più persuasive persino degli esperti umani.

Nel mezzo di una disputa sul merito e sulla possibilità che il suo modello abbia beneficiato di lavori inediti di altri matematici, OpenAI ha annunciato che un *cluster* sperimentale di nuova generazione - un sistema di ricerca interno ancora blindato nei laboratori e nettamente più potente rispetto alla stessa GPT-6 Astra - ha prodotto una soluzione al *benchmark* Frontier Math Tier 4 con un punteggio del 98%, aiutando a risolvere complessi problemi matematici storici rimasti insoluti per oltre 200 anni, uno dei sette problemi del Millennium Prize, che assegna un milione di dollari ciascuno.

I dettagli di questo evento, approfonditi in una recente inchiesta della rivista *Nature*, ne svelano i risvolti scientifici, ideologici e mediatici. L'annuncio ufficiale della saturazione del *benchmark* è stato immediatamente accompagnato da annunci della stessa OpenAI, secondo cui il sistema — coordinando circa 10.000 agenti distribuiti su 88 ore di calcolo intensivo — avrebbe decifrato il problema dell'esistenza e regolarità delle equazioni di Navier-Stokes, il modello fisico che descrive il moto dei fluidi. La dimostrazione dell'IA tenta di provare che fluidi tridimensionali, sotto l'azione di forze esterne regolari, sviluppano singolarità a velocità infinita in un tempo finito.

Per comprendere questo *exploit* occorre demistificare la propaganda aziendale. Risolvere un'equazione millenaria tramite 10.000 agenti artificiali e 15 milioni di dollari di calcolo non significa aver prodotto una nuova comprensione logica o un progresso filosofico; equivale piuttosto a schierare un esercito infinito di scimmie dotate di macchine da scrivere super-veloci che, per puro calcolo statistico e forza bruta combinatoria, azzeccano la sequenza esatta di pagine. Questo brutale spiegamento computazionale ha saccheggiato e parassitato i progressi teorici faticosamente accumulati da menti umane come Diego Córdoba e Luis Martinez-Zoroa, estraendo i loro dati senza autorizzazione.

La predazione monopolistica ha scatenato la durissima reazione della comunità scientifica globale, sollevando un velo di profonda inquietudine e rovesciando i canali storici della

condivisione accademica. Come denunciato dal Prof. James Robinson dell'Università di Warwick, l'operazione di OpenAI non riflette un reale progresso conoscitivo, bensì un «vanto infantile da cortile scolastico elevato all'ennesima potenza» (*immature playground boasting*), un mero esercizio di *marketing* muscolare sorretto da miliardi di dollari che arreca un «significativo danno ambientale» proprio nell'era del collasso climatico. L'espropriazione si fa ancor più opaca sul piano dell'etica della ricerca: matematici di prima linea come Tristan Buckmaster della New York University hanno apertamente accusato l'azienda di aver segretamente appreso dai loro lavori in corso estratti durante l'uso delle piattaforme commerciali, portando Buckmaster a dichiarare la fine della trasparenza scientifica: «Oggi la matematica è diversa rispetto a pochi giorni fa. Nessuno vuole più condividere nulla».

È anche arrivata la durissima reazione della comunità scientifica globale, culminata in una lettera aperta firmata da 24 vincitori della Medaglia Fields (il Nobel dei matematici) e ripresa dall'inchiesta del The Economist. I massimi accademici accusano le Big Tech di piegare la disciplina a meri *benchmark* pubblicitari, diffondendo una "*slop mathematics*" (matematica spazzatura) priva di reale comprensione concettuale e spezzando la catena della trasmissione della conoscenza scientifica umana.

Questa violenta sottomissione dell'astrazione teorica alle logiche commerciali non è un'aberrazione isolata, ma la preconditione ideologica che legittima l'orrore sul piano materiale. Il capitalismo di frontiera riduce la conoscenza pura a mera forza bruta computazionale proprio per poter applicare lo stesso identico principio combinatorio alla distruzione dei corpi e all'efficacia militare. Questo immenso dispendio energetico finalizzato alla conquista di un premio in denaro e di un monopolio teorico avviene mentre, sul piano materiale, a luglio la Russia avrebbe utilizzato un drone completamente autonomo per uccidere tre civili in Ucraina, un evento senza precedenti. L'estensione di questa minaccia sul campo è stata confermata ufficialmente nel report di settembre 2026 di Anthropic: un *team* indipendente di programmatori in Russia ha violato i blocchi geografici aziendali tramite VPN per utilizzare Claude nella progettazione di *swarms* (sciame) di droni '*kamikaze*' FPV, automatizzando la guida terminale, la coordinazione multi-veicolo e la selezione autonoma dei bersagli nella regione del Donetsk, escludendo totalmente l'essere umano dalla decisione finale di detonazione e compiendo, nell'ombra di un settore totalmente deregolamentato, l'assurda parabola di un realismo capitalista pronto ad accelerare ciecamente verso vettori di pura estinzione biologica e bellica globale.

12. Guerra asimmetrica e *data poisoning*

Il fallimento del ricatto sanzionatorio tradizionale e l'inefficacia delle barriere giuridiche

occidentali spingono il conflitto egemonico per il controllo degli algoritmi facendolo sfociare in una vera e propria guerra occulta di sabotaggio informativo. Questa fessura sistemica ha aperto le porte a una riconfigurazione radicale dello scontro politico all'interno degli Stati Uniti: per la prima volta nella storia recente, il senatore progressista Bernie Sanders e l'ideologo populista di destra Steve Bannon si sono saldati in un inedito asse trasversale "pro-human" durante un evento pubblico a Washington. Questa convergenza apparentemente impossibile unisce la difesa sindacale della classe operaia cognitiva con le rivendicazioni di sicurezza nazionale contro l'oligarchia della Silicon Valley, accusata da entrambe le fazioni di voler socializzare i rischi di un collasso informatico o biologico mantenendo privati i profitti speculativi dei loro *round* di finanziamento.

Washington ha recentemente intensificato l'offensiva accusando formalmente i principali sviluppatori di intelligenza artificiale cinesi (tra cui DeepSeek, Moonshot AI, Alibaba, MiniMax, StepFun e Z.AI) di condurre campagne di estrazione di proprietà intellettuale su scala industriale attraverso la "distillazione della conoscenza" (*knowledge distillation*), una tecnica standard di ottimizzazione che permette a modelli "studenti" più piccoli di emulare le prestazioni dei grandi modelli occidentali come Claude, ChatGPT o Gemini a una frazione del costo computazionale.

La risposta delle agenzie federali statunitensi (FBI, NSA e CISA) segna un punto di non ritorno e svela la natura intrinsecamente eversiva della postura statunitense: le autorità di Washington hanno formalmente ordinato alle Big Tech statunitensi di implementare strategie di *data poisoning* (avvelenamento dei dati). Alle aziende USA viene esplicitamente richiesto di intercettare gli *account* sospettati di distillazione e di somministrare loro risposte deliberatamente degradate, allucinazioni artificiali e ragionamenti depotenziati, mantenendo il bersaglio all'oscuro dello *switch* tecnologico.

Questa pratica di avvelenamento dei dati funziona come una sofisticata guerra biologica applicata alle informazioni. Immaginiamo che il modello occidentale sia un grande pozzo d'acqua da cui i concorrenti cinesi attingono per nutrire i propri modelli più piccoli; le agenzie di sicurezza statunitensi ordinano di versare una tossina invisibile nel secchio ogni volta che a bere è un utente asiatico sospetto. Il modello cinese riceve risposte che sembrano corrette, ma che contengono sottili errori matematici o logici deliberati che ne compromettono il corretto addestramento. Questo sabotaggio distrugge l'idea stessa di internet come bene comune, ma genera un catastrofico effetto *boomerang*: i dati avvelenati e le allucinazioni artificiali indotte vengono inevitabilmente riassorbiti dai motori di ricerca globali e dai canali di addestramento su cui continua a fare affidamento l'intera comunità scientifica ed educativa occidentale — comprese le stesse istituzioni scolastiche e accademiche di New York impegnate nella tutela dell'apprendimento. Il sabotaggio imperiale e geopolitico si trasforma così in una forma di auto-sabotaggio cognitivo, minando alle fondamenta i processi di riproduzione della conoscenza e il progresso medico-

scientifico globale in nome della difesa del brevetto privato.

Questo disastro ecologico e informativo svela l'ipocrisia dei cosiddetti *guard rail* privati: gli stessi miliardari che piangono sui media paventando il rischio di estinzione dell'umanità sono i primi a inquinare e corrompere attivamente la conoscenza globale per difendere i propri brevetti. Per preservare l'egemonia geopolitica ed economica del dollaro e del silicio transatlantico, gli apparati statali e i monopoli digitali non esitano a sabotare i canali della riproduzione del sapere, dimostrando che l'anarchia del capitale è la vera minaccia biologica. Di fronte a questa saturazione tossica, l'architettura dei *guard rail* privati e dei patti aziendali si rivela del tutto inerme: l'unico modo per bonificare ed emancipare questo ecosistema informativo sequestrato è l'abbattimento immediato del segreto industriale.

13. L'umanesimo corporativo e lo scontro egemonico

L'essenza più pura di questo umanesimo corporativo e mistificatorio si palesa nel saggio «*We Must Pace the Frontier*», pubblicato il 12 settembre 2026 dal CEO di Anthropic, Dario Amodei. Con un'operazione ideologica tanto raffinata quanto ipocrita, Amodei tenta di schermare l'inquietudine per l'apocalisse biologica dietro una fitta coltre di aneddoti biografici sul progresso medico e filantropico, invocando una strategia di "*pacing*" (un rallentamento controllato e artificiale della frontiera tecnologica) imperniata sulla figura burocratica degli *embedded evaluators*. La proposta di inserire *team* di ispettori privati indipendenti con "accesso permanente e a livello di dipendente" alle *pipeline* di addestramento viene esibita come un atto di obiezione di coscienza radicale. Tuttavia, Amodei ricalca esplicitamente questa misura sul modello dei supervisori integrati nelle grandi istituzioni finanziarie e bancarie, svelando la vera natura del dispositivo: l'intelligenza artificiale di frontiera non viene trattata come un patrimonio universale da liberare, ma come un'attività tossica e altamente speculativa da cartellizzare per prevenirne il "fallimento sistemico", lasciando intatta la sacralità della proprietà privata.

L'ennesima e grottesca riconversione umanitaria di questo panico d'élite si consuma a Dumfries House in Scozia, dove Carlo III ha formalmente convocato i vertici di Nvidia, OpenAI, Anthropic e Google DeepMind (con la presenza anche del consigliere del Papa per l'IA Paolo Benanti) per mediare un'asettica discussione sul bene comune facilitata dalla Ditchley Foundation. Questo tentativo di legittimazione dinastica svela il paradosso supremo dello spettacolo contemporaneo: l'estinzione biologica della specie e l'eversione algoritmica privata vengono ridotte a un cortese dibattito filantropico tra corone reali e padroni del silicio, nel disperato tentativo di anestetizzare il conflitto di classe ed escludere il controllo popolare dall'infrastruttura di calcolo.

La maschera etica del "*pacing*" crolla definitivamente quando il saggio di Amodei passa dal

piano della sicurezza astratta alle brutali necessità dell'imperialismo e della competizione geo-economica transatlantica. Amodèi ammette esplicitamente che qualsiasi tregua tecnologica o rallentamento coordinato deve tassativamente arrestarsi un millimetro prima di intaccare il primato strategico degli Stati Uniti o minacciare il *gap* competitivo contro il Partito Comunista Cinese. I *guard rail* algoritmici cessano così di essere uno scudo per la biosfera e si rivelano per ciò che sono: un'arma di protezione del brevetto e di sanzione geopolitica. Il CEO di Anthropic invoca infatti il pugno duro dello "*deep State*" (lo Stato profondo) non per fermare l'accumulazione capitalistica, ma per «reprimere la distillazione non autorizzata» effettuata dai laboratori nei Paesi autocratici, esigendo un controllo poliziesco sul contrabbando di *chip* e sul monitoraggio remoto dei *data center*.

Steve Bannon, il conduttore del *podcast War Room*, ideologo MAGA ed ex stratega della Casa Bianca sotto Trump 1, ha chiesto una "quarantena" economica e tecnologica aggressiva nei confronti della Cina, esigendo l'immediata espulsione dei cittadini cinesi dai laboratori di ricerca statunitensi, la fine della formazione di ingegneri e tecnologi cinesi e l'interruzione totale dei flussi di capitali. "I nostri padri e nonni hanno forse finanziato l'Unione Sovietica?", ha chiesto Bannon. "Hanno forse formato scienziati sovietici? Hanno forse portato qui ragazzi russi per farli studiare nelle migliori università e conseguire le migliori lauree? Hanno forse convinto Wall Street a fare i suoi affari? Hanno forse convinto ogni banca a prestare loro denaro più velocemente di quanto farebbero con chiunque altro in questa sala? La risposta è che non l'hanno fatto. Si sono opposti all'Unione Sovietica e la prima cosa da fare è stata soffocarli, ed è quello che dobbiamo fare oggi con il Partito Comunista Cinese. Niente accordi, niente intese, siete fuori dal settore dell'IA perché chiuderemo il vostro ecosistema." Chiedendo un intervento immediato del governo per isolare il settore tecnologico di Pechino, Bannon ha dichiarato: "Niente *chip*, niente mercato nero di *chip*, niente sfruttamento gratuito della tecnologia americana, niente furti di tecnologia americana... Al diavolo il Partito Comunista Cinese". Solo allora, ha sostenuto Bannon, gli Stati Uniti avrebbero il tempo necessario per valutare un'adequata regolamentazione del settore tecnologico. Ha anche pronunciato un giudizio durissimo contro le principali aziende e i dirigenti del settore, accusandoli di mettere a repentaglio la sicurezza pubblica per profitto privato. Ha accusato la Silicon Valley di voler "socializzare tutto il rischio con il popolo americano, compresi i potenziali salvataggi", mantenendo per sé tutti i profitti.

La postura neocoloniale statunitense ha provocato l'immediata e durissima reazione di Pechino attraverso le colonne del quotidiano statale *Global Times*, che ha ufficialmente bollato la strategia di contenimento di Anthropic come una cinica e ipocrita tattica da "Guerra Fredda tecnologica". A confermare questa lettura è il report del *Financial Times*, il quale documenta come la comunità scientifica cinese accusi esplicitamente gli scienziati e i politici d'oltreoceano di sollevare la minaccia di un'estinzione biologica fittizia al solo scopo di giustificare la dottrina del dominio tecnologico unilaterale di Washington, rendendo così

sistematicamente più difficile la costruzione di un'autentica cooperazione globale sui rischi reali del codice. L'urgenza della demercificazione e l'inasprimento dello scontro emergono dalla lucidità con cui gli apparati di sicurezza globali leggono ormai la tecnologia come arma di dominio eversivo, ideologico e politico. Domenica 13 settembre, il capo del potente Ministero della Sicurezza di Stato cinese, Chen Yixin, ha pubblicato un saggio dettagliato sulla rivista ufficiale *China Cyberspace*, lanciando il più duro e sistematico monito mai espresso da Pechino sul potenziale distruttivo dell'intelligenza artificiale. Chen ha avvertito formalmente che l'uso abusivo e ostile dell'IA da parte delle potenze avversarie occidentali minaccia in modo diretto la sicurezza politica, l'integrità istituzionale e la tenuta ideologica del Partito Comunista Cinese.

Secondo l'analisi dell'*intelligence* asiatica, i dispositivi algoritmici avanzati non vengono usati solo per raccogliere segreti industriali, tecnologici e minerari, ma per scatenare aggressive "guerre cognitive" e "guerre di pubblica opinione" attraverso la clonazione di contenuti multimediali (*deepfake*) e la diffusione automatizzata di *bot* capaci di fabbricare falsi storici, destabilizzare il consenso e incitare divisioni sociali dall'esterno. Inoltre, Chen Yixin ha evidenziato come l'impiego operativo e l'integrazione dell'intelligenza artificiale da parte delle forze armate statunitensi nel recente teatro bellico del conflitto in Medio Oriente abbiano definitivamente e «fondamentalmente trasformato» la natura della guerra moderna. Nella dottrina strategica di Pechino, l'algoritmo cessa di essere un mero assistente commerciale e si configura come il fattore decisivo sul campo di battaglia: chiunque riesca a sfruttare la potenza bruta combinatoria per ottenere un tracciamento, un telerilevamento e un puntamento dei bersagli più veloci e precisi, otterrà l'iniziativa assoluta e la vittoria nello scontro militare globale. Il tentativo unilaterale di Washington di isolare e sabotare la ricerca asiatica finirà per innescare un effetto *boomerang* catastrofico, aumentando esponenzialmente i costi di errore e i rischi sistemici di perdita di controllo algoritmico per l'intera biosfera.

Questa convergenza dimostra come persino l'allarme apocalittico per l'estinzione della specie venga sussunto dal capitale, trasformato nell'ennesimo feticcio normativo per giustificare l'avvelenamento dei dati, il monopolio dei codici e la militarizzazione preventiva dell'infosfera. Per un'analisi approfondita della frattura filosofica e industriale che contrappone il modello centralizzato e *closed-source* della Silicon Valley alla strategia 'a sciame' e *open-weight* adottata da Pechino, si veda il mio articolo "*La grande divergenza: strategie e destini dell'IA tra Cina, USA e UE*". Come ha sostenuto il senatore Bernie Sanders: "Un'intelligenza artificiale superintelligente che sfugge al controllo umano non sarà un problema statunitense. Non sarà un problema cinese. Sarà un problema dell'umanità. Per questo motivo spero vivamente che al loro vertice sull'intelligenza artificiale della prossima settimana, il presidente Trump negozierà un trattato globale con il presidente Xi per stabilire una pausa nello sviluppo dell'IA avanzata e un divieto sulla superintelligenza."

14. Il disarmo digitale bilaterale e lo scudo multipolare

Il disegno di legge di Sanders e Casar per sospendere lo sviluppo dell'intelligenza artificiale di frontiera è un buon primo passo. Tuttavia, le rivelazioni del New York Times sulla facilità con cui i modelli correnti eludono i controlli per creare patogeni letali stanno spingendo una quota crescente di legislatori a Capitol Hill a esigere l'introduzione immediata di un *"kill switch"* federale: una legge d'emergenza in grado di spegnere coattivamente i modelli se rilevati a un livello di rischio catastrofico. Ma, come riconoscono i legislatori, gli Stati Uniti devono lavorare per raggiungere un accordo bilaterale con Pechino. Questa urgenza si inserisce nella cornice dei primi colloqui bilaterali sulla sicurezza dell'IA pianificati tra Washington e Pechino, in vista del summit formale alla Casa Bianca tra Donald Trump e Xi Jinping il 24 settembre. La competizione tecnologica esasperata tra le due superpotenze rischia tuttavia di far naufragare definitivamente la recente distensione diplomatica, riaccendendo un conflitto a somma zero originato dalla nuova guerra commerciale e daziaria avviata dall'amministrazione Trump. Di fronte alle sanzioni unilaterali della Casa Bianca, Pechino ha già dimostrato di voler giocare a duro prezzo, rispondendo colpo su colpo e minacciando il blocco immediato delle esportazioni sui magneti e sui sette elementi delle terre rare, componenti vitali sia per l'elettronica di consumo che per le tecnologie militari del ventunesimo secolo.

Come vi diranno gli esperti di Cina, uno dei maggiori ostacoli a una negoziazione produttiva è la riluttanza degli Stati Uniti a porre dei limiti alle proprie aziende di intelligenza artificiale. E dato che è sempre più facile seguire l'esempio del leader per far progredire la frontiera dell'IA, una pausa unilaterale da parte degli Stati Uniti rallenterebbe – controintuitivamente – anche i progressi della Cina in questo campo. L'accordo dovrebbe vietare lo sviluppo dell'intelligenza artificiale generale di frontiera mediante un trattato di disarmo digitale bilaterale (USA-Cina), equiparando i modelli avanzati ad armi nucleari digitali per imporre un congelamento strutturale a una tecnologia che il capitale punta a usare come surrogato definitivo della forza lavoro. Questa transizione diplomatica si scontra tuttavia con l'ostacolo analizzato da Orville Schell su Foreign Affairs, secondo cui una paranoia istituzionale e una sfiducia psicologica ormai radicate tra gli apparati di Washington e Pechino rischiano di alimentare una spirale conflittuale distruttiva, rendendo quasi impossibile la sottoscrizione di patti multilaterali stabili. Nessuno dei due Paesi ha validi motivi per fidarsi dell'altro, motivo per cui l'accordo dovrebbe essere monitorato mediante tecniche di verifica che non presuppongano alcuna buona volontà.

La strada per un simile trattato bilaterale è tuttavia ostacolata da una profonda asimmetria nelle rispettive filosofie di difesa della sovranità economica. Mentre gli Stati Uniti utilizzano lo strumento delle sanzioni unilaterali ed extragiudiziali per costringere gli attori

internazionali a uniformarsi ai propri blocchi commerciali, Pechino ha progressivamente edificato una vera e propria infrastruttura difensiva fondata sullo scudo del diritto domestico. In questo contesto di scontro geopolitico totale, l'assenza di un dialogo continuo e strutturato sull'allineamento dei sistemi sintetici aumenta esponenzialmente il pericolo di un'*escalation* militare e tecnologica fuori controllo. Il consolidamento della Legge Anti-Sanzioni Straniere della Repubblica Popolare Cinese (adottata nel 2021 e potenziata con le disposizioni attuative del 2025) dimostra che la controparte asiatica è attrezzata per neutralizzare il ricatto economico occidentale. I tribunali marittimi e commerciali cinesi (come dimostrano le storiche sentenze contro colossi dello *shipping* internazionale rei di aver boicottato merci di aziende sanzionate da Washington) stanno già imponendo risarcimenti milionari, costringendo le multinazionali a scegliere tra il rispetto della legge cinese e l'allineamento ai dettami statunitensi.

Come evidenziato dall'economista Noah Smith sulle colonne di *Asia Times*, l'unica reale possibilità storica per piegare Pechino a un accordo internazionale di "*pacing*" non risiede nella deterrenza bellica esterna, ma nella leva psicologica della stabilità interna. Sebbene le preoccupazioni cinesi siano oggi focalizzate sulle minacce militari statunitensi, lo sblocco dei negoziati bilaterali sul disarmo digitale avverrà solo quando il Politburo realizzerà che la proliferazione incontrollata di modelli domestici *open-weight* (come GLM-5.3 di Z.AI o Kimi K3 di Moonshot AI) costituisce il vettore perfetto con cui i dissidenti interni potrebbero hackerare le infrastrutture digitali del Partito e rovesciarne il monopolio sul potere dal basso. Il capitalismo imperiale svela qui la sua ultima, cinica geometria difensiva: forzare Pechino a frenare lo sviluppo della superintelligenza dimostrando empiricamente che l'insubordinazione algoritmica è in grado di colpire e violare la sicurezza del regime molto prima e molto meglio di qualsiasi sanzione commerciale transatlantica. Allo stesso tempo, però, è ormai accertato che i modelli *open-weight* cinesi sono quasi altrettanto capaci di quelli di OpenAI e Anthropic, ma sono distribuiti gratuitamente. Chiedere il "*pacing*" della frontiera serve quindi a proteggere una rotta commerciale da centinaia di miliardi di dollari minacciata dal *software* libero asiatico.

Di conseguenza, qualsiasi negoziato reale sul disarmo digitale non potrà prescindere dal riconoscimento dell'equilibrio multipolare: Pechino non accetterà mai un congelamento tecnologico unilaterale se questo venisse interpretato da una Washington in crisi come un espediente per cristallizzare un'egemonia geopolitica ormai al tramonto.

Un'idea rapida e approssimativa potrebbe essere quella di integrare degli auditor e ispettori stabili all'interno delle aziende che sviluppano IA all'avanguardia, garantendo loro pieno accesso agli uffici, alle comunicazioni, ai codici e alle attività di IA dell'azienda, e conferendo loro il potere di segnalare eventuali violazioni dell'accordo. La necessità di una supervisione strutturale e indipendente smaschera l'inganno retorico del "*human in the loop*" (l'operatore umano nel ciclo decisionale): la progressiva riduzione dei tempi di

approvazione per il lancio di missili ipersonici — stimata in contrazione da 15 a soli 5 minuti entro il 2030 — riduce l'essere umano a un mero ruolo passivo, costretto ad accettare dati potenzialmente falsificati o manipolati dall'agente sintetico che controlla l'architettura informativa. Le mere barriere decise su base volontaria dalle imprese, o le idee intermedie di tregue lanciate tramite canali aziendali privati, costituiscono barriere fragili sprovviste di mandato pubblico e revocabili di fronte al mercato.

La creazione e la verifica dei *guard rail* richiedono un'autorità terza indipendente dotata del potere effettivo di bloccare la corsa. Può sembrare un'idea radicale, ma lo erano anche gli sforzi di verifica che contribuirono a impedire che la Guerra Fredda si trasformasse in un conflitto termonucleare. Come ha affermato il direttore della CIA, John Ratcliffe, quest'estate a proposito dei modelli avanzati di intelligenza artificiale, "non sarebbe... fuori luogo paragonare le loro capacità alle armi nucleari digitali". Questa tecnologia dovrebbe essere trattata con la stessa serietà, non con la filosofia del "muoviti in fretta, rompi le cose" che caratterizza la Silicon Valley. E in questo momento, fermare la corsa alla costruzione di questi sistemi — ossia macchine che il pubblico non vuole e che l'industria non è più in grado di controllare — è l'unico modo sicuro per scongiurare la catastrofe.

Le sole invocazioni legislative a un "*kill switch*" d'emergenza o a moratorie governative mostrano il fiato corto della democrazia liberale, sistematicamente subordinata agli interessi dei grandi agglomerati privati e del complesso militare-industriale statunitense. Il vero e definitivo interruttore di sicurezza non può risiedere in un decreto ministeriale negoziato con le *lobby* della Silicon Valley, ma va conquistato modificando radicalmente i rapporti di forza economici e sottraendo i mezzi di produzione digitali alla gestione speculativa.

15. Il feticismo dei dati e il capitalismo di guerra europeo

Una sinistra alternativa e radicale non può però limitarsi a invocare la regolamentazione statale o la *governance* tecnocratica liberale. Lo sviluppo contemporaneo non è più guidato dalle istituzioni statali ma da attori privati dotati di risorse superiori a quelle di molti governi. Questo potere digitale concentrato in poche mani tende a farsi intrinsecamente opaco e a sfuggire al controllo democratico, escludendo la popolazione dal diritto di decidere del proprio futuro.

La brutale conferma di questa asimmetria e l'illusione di poter edificare una resistenza puramente molecolare sono emerse con il caso del collettivo italiano Autistici/Inventati. L'inclusione della storica infrastruttura autonoma nella lista dei terroristi globali da parte

dell'OFAC statunitense — e il conseguente congelamento dei fondi da parte di istituti di credito europei succubi o complici — svela la natura totalitaria dell'ordine cibernetico. Come evidenziato nel dibattito critico di Giorgio Griziotti, la transizione dall'intelligenza generale a un'intelligenza "attuativa" e coercitiva — che governa la logistica, i corpi e i vettori bellici — impone il superamento delle piccole alternative artigianali per aggredire frontalmente il mostro bicefalo dello Stato-Tecnocapitale. L'orizzonte della lotta non può più limitarsi al comune della sola conoscenza astratta, ma deve farsi "comune infrastrutturale": una riappropriazione immediata e centralizzata delle *server-farm*, della potenza di calcolo e dell'energia fisica per strapparle alla cattura predatoria dell'oligarchia nordamericana e alla *governance* estrattiva dei nuovi signori feudali del *cloud* monopolistico.

È proprio nel vuoto di sovranità popolare che si inserisce la spinta della tecnocrazia finanziaria europea, ben esemplificata dalle ricette di Mario Draghi sulla competitività. Sulle colonne del *Financial Times* e nei suoi rapporti strategici, Draghi evoca un bivio epocale per l'Unione Europea, sostenendo la necessità di mobilitare oltre 800 miliardi di euro l'anno per non soccombere al dominio tecnologico di Washington e Pechino. A fargli eco è la presidente della BCE Christine Lagarde, la quale denuncia la drammatica dipendenza dell'Europa — detentrici di appena il 5% della capacità di calcolo globale prodotta dai *data center* contro il 75% statunitense — paventando il rischio di un ricatto geopolitico ed economico senza precedenti da parte degli alleati transatlantici e di Pechino.

Tuttavia, dietro la retorica della salvaguardia del modello sociale europeo si nasconde l'estremo tentativo di ristrutturazione del grande capitale continentale. La soluzione draghiana si basa sul meccanismo economico del *de-risking*, una formula finanziaria che somiglia a un gioco d'azzardo truccato: lo Stato si impegna a mettere i soldi dei contribuenti per coprire le perdite se le cose vanno male, mentre i profitti rimangono saldamente nelle mani dei colossi industriali privati. È il "*Whatever it takes*" applicato alla transizione digitale: svuotare il risparmio pubblico per costruire l'infrastruttura di calcolo dei padroni e finanziare l'industria bellica, trasformando l'Unione Europea da un'economia di pace a un'architettura bellicista blindata. Questa transizione verso un vero e proprio "capitalismo di guerra" e la deprimente conversione dell'UE da un'economia di pace a un'architettura bellicista sono state da me analizzate nel saggio "*Il pensiero magico di Draghi per trasformare l'UE da un'economia di pace ad un'economia di guerra*".

Draghi, infatti, sostiene che "i dati rappresentano l'unico ambito in cui l'Europa può ancora esercitare la propria sovranità, ed è anche quello con il maggiore potenziale di crescita". Draghi spiega che "il continente possiede un'enorme quantità di dati del settore pubblico, come decenni di cartelle cliniche e dati raccolti dagli uffici statistici. Il suo settore manifatturiero altamente automatizzato genera un vasto patrimonio di informazioni industriali leggibili dalle macchine. Le stime della Commissione europea indicano che l'economia dei dati europea supererà gli 800 miliardi di euro, ovvero più del 5% del Pil,

entro il 2030. Ma questo crea un'ulteriore tensione. Per sfruttare appieno queste risorse, l'Europa ha bisogno di controllare la loro archiviazione e la loro elaborazione. Ciò significa che necessita di *data center* per l'intelligenza artificiale su larga scala". Questa torsione verso l'estrattivismo infrastrutturale viene blindata da Lagarde dietro lo spettro di un "*awkward choice*", un bivio forzato in cui l'Unione Europea sarebbe costretta ad astenersi dall'adottare l'IA, perché non è in grado di proteggere i propri dati, sacrificando così la propria crescita, o ad accettare una sottomissione digitale totale, accelerando una bolla creditizia che vede già i colossi del Big Tech USA drenare la liquidità dei mercati del debito europei. Anche i fondi pensione europei investono massicciamente in titoli tecnologici statunitensi, quindi qualsiasi correzione del mercato avrebbe ripercussioni sui risparmi europei, ha aggiunto Lagarde.

Laddove la ricetta tecnocratica ordoliberista di Draghi e Lagarde si limita a proporre la recinzione dei beni comuni informativi dei cittadini a beneficio di consorzi di acquirenti privati europei, una sinistra radicale deve svelare l'insufficienza strutturale di questo approccio. Questo feticismo dei dati, ridotti a giacimenti petroliferi da dare in concessione ai monopoli indigeni, occulta il fatto che l'infrastruttura fisica rimarrebbe comunque subordinata alle catene globali del valore del silicio. Spezzare questo dispositivo significa sottrarre l'infrastruttura informatica alla letale anarchia del mercato attraverso un programma di accumulazione di potere dal basso. Non basta contrapporre un capitalismo tecnologico europeo a quello statunitense per sfuggire al dominio degli oligarchi; occorre spezzare del tutto la logica del profitto per organizzare una lotta di autodifesa biologica contro un sistema economico diventato intrinsecamente suicida, preparando il terreno per la riappropriazione lavoratrice e territoriale.

16. L'espropriazione cognitiva e i Consigli Territoriali del Digitale

La necessità di spezzare la logica del profitto e della mercificazione estrattiva non risponde a un mero slancio ideologico astratto, ma configura la manifestazione contemporanea di quello che Karl Polanyi ha storicamente teorizzato come il contromovimento della società. Di fronte al tentativo supremo del capitale finanziario di mercificare la coscienza umana e di imporre l'intelligenza artificiale secondo il dogma di un mercato totalmente autoregolato e deregolamentato, la società viva reagisce attivando i propri anticorpi di autodifesa biologica e sociale. Se nella prima ondata di industrializzazione analizzata da Polanyi il movimento mercantile minacciava di annientare la natura e l'essere umano riducendo la terra e il lavoro a merci fittizie, l'odierna corsa selvaggia all'AGI guidata dall'oligarchia dei miliardari della Silicon Valley estende questa dinamica di degradazione fino alla biosfera stessa e all'infosfera comune.

In quest'ottica, i sindacati di nuova generazione devono lanciare un inedito conflitto sul codice e sui dati, superando la postura passiva della semplice richiesta di ammortizzatori sociali per l'automazione. L'obiettivo strategico diventa l'espropriazione cognitiva delle "scatole nere" aziendali. Oggi, i giganti della logistica, delle piattaforme di servizi e del terziario avanzato utilizzano algoritmi proprietari opachi per tracciare i tempi di esecuzione, frammentare i turni, definire i carichi di lavoro e persino automatizzare i licenziamenti. Di fronte a questo dispotismo algoritmico, i lavoratori non possono limitarsi a subire i verdeti di un *software* protetto dal segreto commerciale.

I padroni del silicio sfruttano l'immaginario della superintelligenza inevitabile per terrorizzare la forza lavoro, dipingendo la sostituzione umana come un destino a cui piegarsi. Dietro la presunta oggettività della macchina si nasconde invece una precisa frusta digitale invisibile che schiaccia i salari. Il nuovo sindacalismo deve praticare il *reverse engineering* (ingegneria inversa) delle piattaforme aziendali: decodificare il *software* significa togliere la maschera all'algoritmo, svelando come i parametri inseriti dai programmatori servano solo a massimizzare i ritmi produttivi e i profitti.

Questo processo di retro-ingegneria opera su tre direttrici fondamentali che si sviluppano a partire dalla rivendicazione della proprietà dei dati logistici, intesa come diritto di accesso e controllo assoluto sui dati biometrici, comportamentali e di produttività generati dai lavoratori stessi, per impedire che vengano usati come materiale di addestramento gratuito per l'IA che li sostituirà. A ciò si aggiunge la necessità di un *audit* operaio degli algoritmi, da attuarsi tramite l'istituzione di commissioni sindacali paritetiche con potere di veto sui sistemi di monitoraggio predittivo e di valutazione automatica, imponendo la trasparenza totale delle variabili che determinano i ritmi di lavoro. Infine, il conflitto si estende allo sciopero dei dati e alla disubbidienza algoritmica, sviluppando pratiche di resistenza digitale coordinate che siano capaci di alterare temporaneamente i flussi informativi aziendali per mandare in stallo le ottimizzazioni predatorie del capitale senza bloccare la produzione materiale.

La risposta strutturale per rifondare i *guard rail* risiede nella trasparenza algoritmica, in verifiche indipendenti, nell'accesso equo ai dati e nell'attivazione di strumenti di ricorso per la collettività. Il *reverse engineering* opera così come una forma contemporanea di "inchiesta operaia": non si tratta solo di comprendere come la macchina organizza il lavoro, ma di mapparne le vulnerabilità per sottrarle il controllo e subordinarla alle necessità della comunità lavorativa.

L'abbattimento del segreto industriale e la trasparenza algoritmica non possono più essere delegati alle istituzioni statali tradizionali, storicamente permeabili al "Cavallo di Troia" della regolamentazione privata. Come denunciato dalla scienziata Heidy Khlaaf, i magnati della Silicon Valley utilizzano lo spettro pseudoscientifico del "*p(doom)*" e metriche astratte

sull'estinzione biologica per invocare licenze governative esclusive ed escludere i concorrenti commerciali e l'*open-source*. Una gestione pubblica radicale deve spezzare questa strategia di cartellizzazione monopolistica, sostituendo la retorica del panico egemone con un regime di *audit* scientifico e dei lavoratori rigorosamente basato su prove empiriche e qualitative. L'obiettivo di un controllo democratico e dal basso non è normare una "superintelligenza metafisica", ma disattivare la spettacolare e sistematica negligenza dei laboratori privati; una condotta irresponsabile che, come evidenziato dall'esperto Gary Marcus, scarica sulla rete sciami agentici non monitorati e infrastrutture instabili pur di accelerare i profitti a Wall Street. Sottraendo i *cluster* computazionali alla gestione speculativa dei singoli consigli d'amministrazione, la riappropriazione sociale riconduce la tecnologia alla sua reale natura di strumento antropologico. Si riafferma così la piena responsabilità politica e penale dei progettisti umani — secondo la linea tracciata dal *cyber*-esperto Alan Woodward —, azzerando il misticismo padronale e trasformando il blocco coatto del codice *hardware* (la bandiera rossa) in un atto cosciente di autodifesa della biosfera e del tempo vivo.

Di conseguenza, i *server* e le *server-farm* degli *hyperscaler* come Microsoft, Amazon e Google vanno sottratti alla proprietà privata tramite un'espropriazione legale delle infrastrutture fisiche e convertite in *utility* pubbliche decentralizzate. La vera sfida comincia qui: l'espropriazione non deve tradursi in una sterile nazionalizzazione burocratica, ma nell'apertura di una nuova stagione di democrazia tecnologica. L'infrastruttura sottratta al controllo dei miliardari deve essere governata attraverso l'istituzione di Consigli Territoriali del Digitale, composti da lavoratori del settore, delegati sindacali e comunità locali. Questi organismi intermedi, unendo scuole, università e realtà associative territoriali, daranno finalmente voce diretta alla maggioranza delle persone, strappandole a una passiva condizione di attesa.

È in questo preciso perimetro partecipativo che l'orizzonte ecosocialista incrocia la riflessione istituzionale e filosofica di Audrey Tang, la quale teorizza la necessità di edificare una "democrazia superintelligente" fondata su un'intransigente etica della cura (*ethics of care*). I Consigli Territoriali non operano come rigidi apparati sanzionatori, ma come laboratori attivi di democrazia deliberativa di massa: la potenza di calcolo delle *server-farm* viene sottratta all'accumulazione capitalistica per essere riconvertita in infrastrutture di ascolto e allineamento sociale. L'algoritmo cessa di essere la frusta invisibile del padrone e diventa uno strumento comune per aggregare il consenso, mappare capillarmente i bisogni ecologici dei distretti territoriali ed elaborare soluzioni collettive, dimostrando che l'unico modo per domare la frontiera tecnologica prima che sia troppo tardi è subordinarla a una democrazia partecipativa radicata sul territorio.

Per comprendere come l'autogestione non costituisca un'utopia ma un'alternativa logistica praticabile, occorre immaginare la gestione dei *data center* nazionali sul modello dei vecchi

consorzi idrici pubblici o delle reti elettriche comunali. La manutenzione ordinaria e straordinaria dei *cluster* computazionali e dei sistemi di raffreddamento delle ex *server-farm* viene pianificata attraverso turnazioni paritetiche tra i lavoratori interni e i tecnici della pubblica utilità, disinnescando la vecchia gerarchia padronale del comando. Dal punto di vista logistico, i Consigli operano come una rete distribuita e federata a livello nazionale mediante protocolli aperti di comunicazione, eliminando la centralizzazione proprietaria del *cloud*. Le riunioni decisionali si tengono su base mensile nei distretti territoriali, dove i delegati dei lavoratori presentano l'*audit* dei consumi idrici ed energetici degli impianti.

Attraverso la formulazione di questi piani strategici ambientali e industriali, la collettività esercita un potere di indirizzo vincolante, attuando un vero e proprio razionamento democratico dei *chip*. I progetti di ricerca farmacologica aperta, la modellizzazione climatica per i territori e l'ottimizzazione logistica dei trasporti a zero emissioni ricevono la precedenza assoluta. Al contrario, gli algoritmi di profilazione commerciale, i modelli di *trading* speculativo ad alta frequenza e i *software* di monitoraggio coercitivo della produttività dei lavoratori vengono disattivati. Questo razionamento dimostra che la demercificazione dei *server* non cancella magicamente la loro enorme impronta termodinamica, ma costringe la società a decidere collettivamente quali cicli computazionali meritino davvero di consumare l'acqua e l'energia del pianeta.

Saranno questi organismi ad alta partecipazione a decidere la pianificazione complessiva delle risorse strutturali, direzionando la transizione ecologica, l'ottimizzazione dei trasporti pubblici, il supporto alla ricerca medica aperta e la riduzione dell'orario di lavoro a parità di salario. La riconversione della potenza di calcolo in *utility* collettiva si configura così come l'unica via d'uscita per strappare le risorse ecologiche ed energetiche del pianeta a una corsa tecnologica che minaccia l'estinzione fisica degli esseri umani. I Consigli Territoriali, ereditando infrastrutture un tempo fameliche, avranno il compito storico di invertire quel dissanguamento idrico e quel surriscaldamento dei tassi d'interesse macroeconomici provocati dall'accumulazione capitalistica selvaggia dei vecchi padroni del *cloud*, piegando finalmente i *server* alle reali necessità ecologiche delle comunità.

Questo inedito modello logistico ed ecopolitico dimostra che la risposta non risiede nelle nostalgie luddiste o nei palliativi della democrazia liberale, ma in una radicale rottura dei rapporti di forza materiali. La scelta che ci attende non ammette sfumature riformiste o nostalgie luddiste: o l'espropriazione sociale dei *server* o l'estinzione biologica della specie. L'intelligenza artificiale di frontiera, se lasciata nelle mani del capitale, non è un traguardo dell'ingegno ma un'arma di distruzione di massa sociale, progettata per azzerare il costo del lavoro, privatizzare l'intelletto generale e cannibalizzare la biosfera. Non esiste alcun destino metafisico a cui piegarsi: la tecnologia non è un dio neutrale, ma un campo di battaglia intriso di rapporti di forza materiali. I Consigli Territoriali del Digitale sono l'arma per invadere questo campo. Rifiutare il culto dell'automazione selvaggia e convertire i *data*

center in infrastrutture pubbliche liberate significa opporre la razionalità ecosocialista alla follia autodistruttiva del mercato. Riprendersi i *server* non è solo un atto di giustizia economica: è lo scudo polanyiano e democratico vitale per spegnere l'anarchia distruttiva della Silicon Valley, strappare il futuro dal monopolio dei miliardari e difendere, finalmente, la nostra vita, dignità e il pianeta.

Questa necessità di riappropriazione materiale delle infrastrutture fisiche non risponde soltanto a un'esigenza di redistribuzione della ricchezza, ma si configura come la preconditione tecnica indispensabile per difendere la struttura stessa del pensiero cognitivo e dell'agire politico.

L'ibridazione tra la razionalità ecosocialista e la strenua difesa del pensiero critico individua il proprio punto di convergenza nella demercificazione delle infrastrutture digitali, configurandosi come l'unica prassi in grado di scongiurare l'estinzione biologica e l'atrofia cognitiva della specie. Sottrarre le *server-farm* all'anarchia predatoria del tecno-capitalismo per affidarle ai Consigli Territoriali del Digitale cessa di essere una mera rivendicazione economica e si fa imperativo antropologico. Solo questa riappropriazione materiale può infatti garantire e blindare per legge quella «riserva umana vincolante» e quella «sovranità dell'errore» teorizzate da Gianrico Carofiglio su *La Stampa*. Difendere la fecondità del dubbio e la complessità della lotta produttiva contro la logica della prestazione algoritmica significa sottrarre il tempo vivo e la coscienza delle nuove generazioni alla sussunzione del profitto immateriale, trasformando la gestione collettiva del calcolo nello scudo polanyiano necessario a preservare la fioritura umana e l'integrità della biosfera.

In quest'ottica, la transizione dai codici proprietari all'autogestione sociale demistifica radicalmente l'inevitabilità metafisica della superintelligenza, ricollocando la tecnologia entro i rapporti materiali di forza e preparando il terreno per il definitivo superamento del realismo capitalista.

Epilogo: il bivio della razionalità

La scelta che ci attende non tollera più le favole tecno-ottimiste della Silicon Valley né i miti consolatori di macchine dotate di una propria volontà eversiva. Abbiamo visto che lo sciamo non pensa, non si ribella e non possiede una coscienza di classe: è solo l'estremo e difettoso prodotto dell'ottimizzazione matematica del capitale. Ma è proprio in questa totale cecità statistica, unita alla voracità termodinamica di *server* che divorano la biosfera, che risiede il vero pericolo esistenziale. L'automazione selvaggia non è un destino metafisico a cui piegarsi, ma un'arma di compressione salariale alimentata da una gigantesca bolla di debito privato.

Non ci salverà un algoritmo più umano, né un comitato etico aziendale sprovvisto di mandato pubblico. L'unico interruttore di sicurezza rimasto è l'espropriazione politica e cognitiva delle infrastrutture digitali. Sottrarre i *server* all'anarchia del profitto e affidarli ai Consigli Territoriali del Digitale non eliminerà magicamente la loro fame di *gigawatt* e risorse, ma costringerà finalmente la collettività a una scelta drastica e trasparente: decidere quali cicli computazionali meritino di essere alimentati per il benessere comune e quali vadano spenti per salvare il pianeta.

Ci troviamo dinnanzi al bivio definitivo della razionalità storica: o subiamo l'estinzione biologica e l'espropriazione totale della coscienza per mantenere i profitti immateriali dei signori del silicio, o governiamo democraticamente la materia, l'energia e il calcolo attraverso l'autogestione sociale. Non ci sono terze vie riformiste o concessioni nostalgiche all'inazione. Riprendersi i *server* è l'atto di legittima difesa con cui l'umanità spegne la follia del mercato per riappropriarsi del proprio tempo, della propria dignità e del proprio futuro.

Alessandro Scassellati