

Per un'IA non produttivista

Tiziano Bonini

10 Agosto 2026

Sono le 7:14 di un mattino di marzo. Valeria, 52 anni, product manager, legge sul telefono la notifica del suo assistente medico automatizzato. Il tono è rassicurante, la prosa quasi affettuosa: nelle ultime tre settimane la variabilità della sua frequenza cardiaca ha mostrato un pattern leggermente fuori range, compatibile con diverse condizioni che è opportuno escludere. Una visita cardiologica è già prenotata per giovedì, alle 10:30, al centro diagnostico a quattrocento metri da casa — il più vicino con disponibilità compatibile con la sua agenda. Il ticket è già pagato, scalato dalla sua assicurazione integrativa. Valeria chiude la notifica, mette su il caffè. Non ha mai parlato con un medico di questa cosa. Non sa se il cuore le batte così perché è stanca, o perché da pochi mesi frequenta un nuovo compagno, o perché ha paura di qualcosa che non riesce ancora a nominare.

A qualcuno lo scenario sembrerà perfetto: salute monitorata, visita prenotata dietro casa. Eppure c'è qualcosa di sinistro — un rumore di fondo che richiama subito *Black Mirror*. Forse l'IA medica ha ragione, forse la visita è necessaria. O forse è un errore, un bias. Ma quello che mi spaventa di più di questo scenario, è la presunta, levigata efficienza di questo sistema.

"Efficienza" non è una categoria neutrale. È un'ontologia — un modo di organizzare il mondo che stabilisce cosa conta e cosa no, cosa si misura e cosa si lascia nell'ombra. Quando diciamo che un sistema è efficiente, lo diciamo sempre rispetto a qualcosa: una funzione obiettivo, un criterio di valore che qualcuno ha scelto al posto nostro. Nel caso dell'IA applicata alla medicina, all'istruzione, alla psicoterapia, alla consulenza finanziaria, la funzione obiettivo è quasi sempre la stessa: fare la stessa cosa con meno forza lavoro umana, o più cose con la stessa, nel minor tempo possibile. È il mandato implicito con cui l'IA viene venduta a ospedali, università, aziende. Non è una conseguenza tecnica inevitabile ma una scelta politica nascosta dietro la retorica dell'ottimizzazione dei costi e della forza lavoro.

Immaginate allora un futuro non lontano in cui il percorso si è compiuto: non solo Valeria non parla con un medico, ma non ne esiste più uno da incontrare. Il suo assistente virtuale le interpreta i parametri vitali rilevati 24/7 dal dispositivo al polso, le comunica le diagnosi, le prescrive i farmaci, prenota gli esami. Un algoritmo ha già deciso che i suoi dati rientrano nella categoria C, protocollo 7. Altrove, uno studente studia sul divano: un assistente didattico gli prescrive le letture, verifica la comprensione con quiz adattivi, valuta le esercitazioni, e alla fine gli conferisce la laurea. Personalizzato, ottimizzato, calibrato sulla sua velocità di apprendimento. Non ha mai discusso un'idea con un altro essere umano.

So bene che questi scenari speculativi sono semplificazioni. Il futuro reale sarà più caotico, pieno di frizioni, di usi imprevisti — e gli esseri umani non verranno semplicemente sostituiti: accadrà (accade già) qualcosa di più sottile. Verranno *deskillati*, svuotati delle competenze che non esercitano più, fino al punto in cui la sostituzione diventa inevitabile perché nessuno sa più fare quella cosa senza la macchina. Non serve che l'IA diventi davvero migliore di un medico o di un professore in quasi tutto ciò che facevano, basta che diventi abbastanza economica da rimpiazzarli, e abbastanza "efficiente" da poterlo spacciare per progresso. La sostituzione non ha bisogno di una superiorità reale — ha bisogno di un costo più basso e di un mandato: fare la stessa cosa con meno. È così che medici e professori — e le altre categorie di "esperti" umani — diventano, sulla carta, ridondanti.

Ivan Illich l'aveva visto cinquant'anni fa, senza l'IA. In *Nemesi medica* (1975) e *Descolarizzare la società* (1971) aveva descritto il momento in cui un'istituzione produce l'effetto opposto al suo mandato: la medicina che rende malati, la scuola che distrugge il desiderio di apprendere. Chiamava questo fenomeno controproduttività — la soglia oltre la quale il potenziamento tecnico di un sistema comincia a danneggiare la cosa stessa che dovrebbe proteggere. L'IA produttivista è la forma contemporanea, e più acuta, della controproduttività illichiana: uno strumento che promette di migliorare salute e istruzione eliminando le relazioni umane in cui salute e apprendimento si producono davvero.

Daron Acemoglu
Simon Johnson

Potere e progresso

La nostra lotta millenaria
per la tecnologia e la prosperità

Traduzione di
Fabio Galimberti
e Paola Marangon

ilSaggiatore

L'efficienza promessa dall'IA è il vero pericolo a breve termine per la nostra società. È lo strumento più potente che il capitale abbia mai avuto per giustificare la riduzione della spesa pubblica e del costo del lavoro — e lo fa col vocabolario della modernizzazione, dell'innovazione, della personalizzazione, dell'efficienza. Presenta come progresso tecnico quella che è, nella sostanza, una scelta politica di redistribuzione verso il basso: meno risorse per la sanità e l'istruzione

pubbliche, meno potere contrattuale per i lavoratori della conoscenza. Il problema è l'economia politica in cui è fiorita questa IA, questa che ci stanno vendendo oggi.

Un passo di lato. Nel 1997 Deep Blue batté Kasparov e scoprimmo improvvisamente che la macchina era più brava di noi nel gioco più complesso che avessimo inventato. Vent'anni dopo, con AlphaZero, la riflessione si è fatta più interessante: il modo in cui le IA giocano rivela una capacità qualitativamente diversa dalla nostra, non semplicemente superiore. Esplorano lo spazio delle possibilità in modo sistematico, senza la pressione del tempo reale, senza l'effetto tunnel delle abitudini. Vedono mosse che noi non vediamo non perché siano più intelligenti, ma perché operano in un regime cognitivo diverso: possono esplorare sette varianti invece di tre, andare avanti di dieci mosse invece di quattro, senza il costo che questo ha per un cervello umano in competizione. Gli scacchi, certo, sono un mondo chiuso e a informazione perfetta, mentre la medicina e l'istruzione sono aperte e ambigue — ma è la capacità esplorativa, non la cornice, il punto. Perché quella capacità potrebbe essere usata in due modi opposti. Nel primo — quello che stiamo scegliendo, quasi senza accorgercene — la usiamo per comprimere il processo decisionale: l'IA produce la risposta migliore nel minor tempo possibile, sostituendo il percorso deliberativo umano. Nel secondo potremmo usarla per espanderlo: l'IA mappa lo spazio delle ipotesi, porta sul tavolo le possibilità che non avremmo visto, e poi ci lascia il tempo di ragionarci sopra, sia da soli, che in gruppo. Invece di accelerare le decisioni, le rallenta. Invece di automatizzare i task, arricchisce il campo in cui le decisioni si prendono.

Immaginate ora un dispositivo medico che per legge non può emettere una diagnosi, ma solo generare una mappa di ipotesi da portare al medico, o un assistente didattico che non valuta e non prescrive, ma elabora le domande che lo studente non avrebbe saputo formulare da solo, e che diventano il punto di partenza di una conversazione con il professore.

Torniamo a Valeria. Immaginiamo un'altra versione dello stesso mattino. Valeria riceve la stessa notifica — lo stesso pattern, le stesse tre settimane di dati. Ma la notifica non ha prenotato niente, non ha deciso niente. Le ha preparato una mappa: cinque ipotesi diagnostiche, ordinate non per probabilità statistica ma per pertinenza rispetto alla sua storia clinica. Le ha generato dieci domande — che lei non avrebbe saputo farsi — da portare al medico. E le ha fissato un appuntamento per il giorno dopo, non al centro diagnostico più vicino ma dal suo medico di base. Il medico ha ricevuto la stessa mappa. Arriva all'incontro senza dover ricostruire l'anamnesi in dieci minuti, senza l'ansia di sbrigare una cosa per

passare alla successiva. Hanno venti minuti di conversazione vera — non di burocrazia clinica, ma di ragionamento condiviso su un corpo specifico, su una persona di 52 anni che magari è stanca, o innamorata, o spaventata di qualcosa che non sa ancora nominare. Alla fine Valeria scopre qualcosa che l'algoritmo non le avrebbe mai detto: che il medico pensa sia probabilmente stress, ma che vale la pena escludere lo stesso la causa cardiaca — e le spiega perché.

Questo secondo scenario, però, presuppone qualcosa che il primo non presuppone: che il medico ci sia. Che abbia tempo. Che non abbia trenta pazienti in sala d'aspetto. Che il sistema sanitario sia finanziato per permettere visite di venti minuti invece di sette. Significa, con tutta la brutalità che la cosa richiede, più medici, non meno. Più professori, non meno. Più psicologi, più esperti umani con cui condividere il ragionamento che l'IA ha reso più ricco.

Dobbiamo disaccoppiare l'IA dall'ideologia dell'efficienza. Dobbiamo immaginare una ontologia alternativa a quella neoliberale, dove l'IA non sia al servizio dell'ottimizzazione dei costi umani. Mentre facevo un po' di ricerca su questo argomento, ho scoperto che c'era un economista che aveva già detto questa cosa molto meglio di me, 6 anni prima. È la tesi di Daron Acemoglu, l'economista turco-americano del MIT che nel 2024 ha vinto il Nobel per l'Economia insieme a Simon Johnson e James Robinson. In [*"The Wrong Kind of AI?"*](#) (2020), scritto con Pascual Restrepo, Acemoglu distingue due direzioni possibili per l'intelligenza artificiale: una che automatizza — sostituisce il lavoro umano nei compiti che già svolgeva — e una che aumenta, che crea nuovi compiti in cui il lavoro umano ha un vantaggio comparato. La prima comprime la domanda di lavoro; la seconda la espande. E c'è un corollario che riporta al sospetto di prima: quell'efficienza promessa, spesso, non arriva nemmeno. Acemoglu e Restrepo la chiamano “so-so automation” — automazione “così così”, adottata non perché renda davvero di più, ma perché è appena abbastanza a buon mercato da soppiantare il lavoro umano. È la peggiore delle combinazioni: si perde il lavoro senza incassare la produttività promessa. E il punto di fondo — che Acemoglu e Johnson allargheranno in [*Potere e progresso*](#) (2023, trad. it. il Saggiatore, 2024) — è che nulla, in questa biforcazione, è tecnicamente predeterminato. È una scelta politica ed economica. Stiamo scegliendo, quasi senza deciderlo, l'IA che sostituisce mentre potremmo scegliere quella che moltiplica il bisogno di lavoro umano qualificato.

Progettare un'IA non produttivista e “conviviale” è tecnicamente possibile, ma nell'attuale congiuntura storica, sembra improbabile: viviamo in un momento dominato dall'etica della prestazione, dall'ottimizzazione come valore in sé,

dall'idea che qualsiasi processo debba essere reso più rapido, più scalabile, più misurabile. Immaginare un'IA progettata per rallentare le decisioni invece di accelerarle sembra quasi una provocazione. Eppure è in questo punto che i sentieri del futuro si biforcano. Finché l'obiettivo sarà fare di più con meno, l'IA non potrà che produrre meno medici, meno professori, meno relazioni, meno tempo.

In copertina, fotografia di [Immo Wegmann](#) - Unsplash.

Se continuiamo a tenere vivo questo spazio è grazie a te. Anche un solo euro per noi significa molto.

Torna presto a leggerci e [SOSTIENI DOPPIOZERO](#)

