

L'AI farà davvero strage di posti di lavoro? Ecco chi continuerà a essere insostituibile, i problemi che i «bot» non sanno risolvere di Luca Angelini

Le notizie contraddittorie sul fronte del lavoro: le aziende che mandano a casa migliaia di dipendenti e quelle (come la Ford) che riassumono centinaia di ingegneri esperti. I problemi che i «bot» non sono in grado di risolvere (Fonte: <https://www.corriere.it/> 2 luglio 2026)



Notizie delle ultime settimane: Block - la società fondata da **Jack Dorsey** (già fondatore di Twitter), che controlla Square e Cash App - ha annunciato il **taglio del 40% della propria forza lavoro**, oltre quattromila persone, adducendo come causa esplicita l'avvento degli «*strumenti di intelligenza*». Cloudflare ha **eliminato 1.120 posizioni** - il 20% del personale - nonostante ricavi record nel primo trimestre del 2026: per l'amministratore delegato Matthew Prince certi ruoli «*non sono i ruoli di cui le aziende avranno bisogno nel futuro*».

Altre notizie di queste ore: la casa automobilistica Ford sta riassumendo centinaia di ingegneri esperti per lavorare su problemi di qualità che i sistemi automatizzati non sono in grado di risolvere.

«L'intelligenza artificiale è uno strumento fantastico, ma è valida solo quanto le informazioni che si usano per addestrarla» - ha spiegato **Charles Poon**, vicepresidente per l'ingegneria hardware di Ford. La Commonwealth Bank of Australia ha **licenziato più di 40 addetti al servizio clienti**, sostituendoli con un bot vocale basato sull'intelligenza artificiale. Tuttavia, il sistema di AI non è stato in grado di gestire la situazione, causando un aumento delle chiamate e spingendo Cba a revocare i licenziamenti.

IBM ha sostituito le sue funzioni di risorse umane con l'intelligenza artificiale, che gestiva circa il 94% delle richieste di routine, ma non era in grado di soddisfare il restante 6%, che

comprendeva **dilemmi etici**. IBM ha poi annunciato l'intenzione di **triplicare le assunzioni di personale di livello base** negli Stati Uniti in tutte le unità aziendali entro il 2026. «*Se non continuiamo a investire nelle assunzioni di personale entry-level, cosa succederà tra 3-5 anni?*», si è chiesto Nickle LaMoreaux, responsabile delle risorse umane di IBM.

Le interazioni

Ma, allora, l'intelligenza artificiale è destinata a far strage di posti di lavoro, oppure no? Zeynep Tufekci, docente di Sociologia a Princeton, [sul New York Times](#) si è detta convinta che l'allarme sia eccessivo e ha citato altri casi di marce indietro clamorose. «A marzo, Meta ha annunciato che gli utenti di Facebook e Instagram a cui erano stati bloccati gli account **non avrebbero più interagito con un operatore del servizio clienti, ma con un'intelligenza artificiale** appositamente addestrata. Cogliendo l'opportunità, alcuni truffatori hanno convinto l'AI a cedere loro il controllo di [oltre 20.000 account Instagram](#), inclusi quelli della Casa Bianca di Obama e di un alto funzionario dell'amministrazione Trump. Poi hanno inondato i forum di Telegram con i loro racconti entusiasti di quanto fosse stato facile. Non si è trattato di un caso isolato.

Air Canada [ha disattivato](#) i suoi chatbot dopo che questi avevano **promesso per errore un rimborso a un cliente, il quale ha poi intentato causa e vinto**.

McDonald's [ha eliminato](#) il bot che prendeva gli ordini al drive-through dopo che diversi video virali ne avevano dimostrato il grave malfunzionamento. In un caso, **il bot ha aggiunto per errore centinaia di dollari di crocchette di pollo all'ordine di un cliente**».

I motori di plausibilità

Secondo Tufekci si tratta di qualcosa di più di incidenti di percorso: «Non sono il risultato di errori di programmazione. Sono il risultato di un fatto essenziale e ineludibile sull'intelligenza artificiale che è diventata così comune in tanti aspetti della nostra vita quotidiana: **i grandi modelli linguistici non sono macchine di ragionamento. Sono motori di plausibilità**. Non si tratta solo del fatto che non verificano i loro output per assicurarsi che siano corretti o logici, o che non lo facciano in certi casi. Non possono, e non saranno mai in grado di farlo da soli. **Possono solo valutare quali risposte sono probabili, in base ai dati su cui i modelli sono stati addestrati**. E questo vale sia che siano addestrati sull'intera gamma di output umani che solo su articoli scientifici. È intrinseco al loro modo di funzionare».

Questo difetto d'origine è quel che rende la sociologa fiduciosa sulla sopravvivenza di lavoratori in carne ed ossa: «È per questo che non do retta alle fosche previsioni di un'imminente apocalisse occupazionale causata dall'AI. I modelli basati sull'intelligenza artificiale possono fare molte cose con una competenza sorprendente, ma **non sono in grado di svolgere la stragrande maggioranza dei lavori umani senza incorrere in disastri qua e là**. Nessun aggiornamento o lancio di nuovi modelli cambierà questa situazione».

Gli errori

L'AI funziona benissimo per alcuni lavori, come la **programmazione informatica**, che si basa su un linguaggio strutturato e formale che può essere testato in tempo reale. «Tuttavia - sottolinea Tufekci - **un numero enorme di lavori non funziona in questo modo: non i lavori da chirurgo, non i lavori nel servizio clienti e non i lavori da insegnante di quarta elementare**. Per questi lavori è necessaria la tecnologia specializzata della buona vecchia intelligenza umana. Trascorro molto tempo a parlare di questi temi in contesti pubblici, e una domanda emerge sempre: anche i lavoratori umani commettono errori, quindi integriamo delle misure di sicurezza per individuarne la maggior parte. Perché non possiamo fare lo stesso per gli errori dell'intelligenza artificiale generativa?

Il problema è che **questi modelli non commettono lo stesso tipo di errori di un essere umano**. Né le loro impressionanti capacità né le loro strane debolezze si adattano bene a un tipo di intelligenza umana. Questa discrepanza **rende difficile integrarli in sistemi progettati per individuare gli errori umani**».

Tufekci, in proposito, scrive che, forse per l'impressionante «proprietà di linguaggio» dei modelli di AI (e forse perché qualcuno doveva venderli) «siamo stati **tratti in inganno sulla natura di questa tecnologia**». Perché quella che si è affermata non è un'intelligenza artificiale **simbolica** (do alla macchina regole e istruzioni e lei le segue alla perfezione) ma un'intelligenza artificiale **concessionista** (*connectionist AI*). Fa una bella differenza: «I nostri modelli attuali sono sistemi concessionisti, resi possibili da enormi quantità di dati e potenza di calcolo. **Generano risposte basate non sulla verità o sul ragionamento, ma sulle probabili connessioni tra i dati che sono stati loro forniti**. Da qui il nome: AI generativa. Non possiamo controllare completamente i modelli generativi. Tutto ciò che possiamo fare è addestrarli e poi cercare di indirizzarli nella giusta direzione. Anche in questo caso, **non possiamo mai essere certi che i nostri interventi avranno l'effetto desiderato, perché non comprendiamo appieno il funzionamento di questi modelli**. Sono delle scatole nere».

Certo, abbiamo escogitato modi per migliorare i modelli, come il **rinforzo tramite feedback** (grandi team di esseri umani per monitorare tutti gli output del modello e rispondere con un pollice in su o in giù) o i **suggerimenti di sistema**. Ma, ad esempio, più una chat si protrae, più quei suggerimenti di sistema diventano un ricordo lontano. «Da qui la nascita del "*jailbreaking*", termine che indica la manipolazione di uno di questi sistemi per aggirarne le protezioni. (...) In uno dei primi casi di *jailbreak*, un chatbot digitale di una concessionaria Chevrolet è stato manipolato per vendere a qualcuno un nuovo SUV [per 1 dollaro](#). "*Affare fatto*", ha detto il chatbot, "*e questa è un'offerta legalmente vincolante, non si torna indietro*"». E problemi ci sono stati anche in un colosso come **Amazon**, come ha ricordato Velia Alvich [su Login](#).

Gli atti distruttivi

La conclusione di Tufekci è abbastanza drastica: «Le attività facilmente automatizzabili sono già state eliminate dalla maggior parte dei nostri lavori anni fa, grazie alla tecnologia tradizionale basata su regole. Gran parte di ciò che rimane non può essere ridotto così facilmente a giusto o sbagliato, bianco o nero. **Richiede qualcuno con almeno un po' di buon senso e capacità di ragionamento**, non un chatbot basato sull'intelligenza artificiale che **si lascia convincere con le buone maniere a fare cose che sfidano la logica**. (...) Non passa giorno senza che io senta parlare di un sistema di intelligenza artificiale che cancella l'intero codice sorgente o gli archivi di qualcuno, o che si impegna in altri **atti distruttivi**. Ora immaginate questi sistemi scatenati, su vasta scala, che attaccano reti sanitarie, banche, sistemi di controllo del traffico aereo, infrastrutture critiche e reti di difesa. (...) L'intelligenza artificiale generativa, nella sua forma attuale, non può facilmente sostituire gli esseri umani, perché non è in grado di manifestare l'intelligenza umana. **Ciò non le impedirà, però, di destabilizzare la società in modi ben più profondi di quanto possiamo immaginare**. Prima aggiorneremo il nostro modo di pensare allo stato attuale dell'AI, prima smetteremo di allarmarci per le cose sbagliate e inizieremo a prepararci per le trasformazioni che porterà realmente al nostro mondo».

Qualcuno, a quelle trasformazioni, anche in termini di posti di lavoro, pare si stia già preparando: la Cina, ad esempio, come ha ricordato Danilo Taino [su AmericaCina](#) (con l'avvertenza fatta da Katrin Bennhold [sul New York Times](#), che «un controllo in stile cinese sull'industria tecnologica non è fattibile nella maggior parte dei Paesi occidentali. Ma la Cina dimostra che i responsabili politici hanno un ruolo attivo in questa tecnologia. Possono **influenzarne la direzione**, invece di lasciare che le cose dell'AI vadano come vadano»).

Le trasformazioni

Posto che le trasformazioni di cui sopra sembrano giustificare visioni più pessimistiche di quelle di Tufekci (se ne è occupato di recente Alessandro Trocino in questa [Rassegna](#)), padre Paolo Benanti, [sul Sole 24 Ore](#), invita a non fermarsi ai soli numeri per giudicare l'impatto dell'AI sull'occupazione. «I dati sulle offerte di lavoro **non ci dicono se questi posti siano meglio retribuiti o più precari dei precedenti**, se siano accessibili ai lavoratori che hanno perso l'impiego o se richiedano competenze che solo una minoranza già possiede. Non ci dicono se la fase attuale - nella quale le aziende assumono figure capaci di interagire proficuamente con i sistemi di intelligenza artificiale - sia strutturale o transitoria, destinata a invertirsi nel momento in cui i modelli diventeranno più autonomi e i meccanismi di controllo più affidabili. E soprattutto **non ci dicono nulla sulla qualità del lavoro che sopravvive**: se la supervisione dell'output automatico costituisca una forma autentica di realizzazione professionale oppure una nuova figura di alienazione, in cui l'umano è ridotto a controllore di una produzione che non comprende pienamente e non governa. La complessità dei mercati del lavoro, insomma, non è solo un dato statistico: è una condizione epistemica.

Non sappiamo ancora se ci troviamo di fronte a una **transizione ordinata** - analoga all'automazione manifatturiera degli anni Ottanta, dolorosa per molti ma compensata nel lungo periodo da nuovi settori emergenti - o a qualcosa di qualitativamente diverso, in cui **la velocità del cambiamento supera la capacità adattiva delle istituzioni e degli individui**. Ciò che è certo è che né il catastrofismo né l'ottimismo ingenuo rendono giustizia alla realtà. Il [paradosso di Jevons](#) ci insegna che l'efficienza non risparmia risorse: le moltiplica. Ma non ci dice nulla su chi le riceve. Questa è una domanda politica, non tecnica. Ed è la domanda che conta».