

Potremmo avere una sola possibilità per fermare l'intelligenza artificiale.

Un ricercatore in procinto di lasciare Anthropic avverte che la corsa allo sviluppo di un'intelligenza artificiale in grado di auto-migliorarsi potrebbe privare l'umanità dell'opportunità di riprendere il controllo.

È terribile dirlo, ma in questo momento spero davvero che [le valutazioni nel settore dell'IA crollino](#) e che il settore subisca delle battute d'arresto che ci diano mesi per riflettere sulle aziende che ne fanno parte.

Il ricercatore di antropologia Jacob Coxon sta lasciando il settore dell'intelligenza artificiale perché ritiene che la corsa alla creazione di sistemi auto-miglioranti stia procedendo più velocemente della capacità del settore di tenerli sotto controllo, secondo quanto riportato ieri da [Quoth the Raven](#) .

Coxon, che in precedenza ha lavorato presso OpenAI, ha affermato che gli sforzi di Anthropic in materia di sicurezza

erano sinceri, ma che la concorrenza rende difficile imporre restrizioni significative senza un coordinamento più ampio.

In un'epoca in cui tutto – dalla guerra a Wall Street, dalle reti elettriche ai sistemi di pagamento, dal controllo del traffico aereo alle reti ospedaliere, fino al software che garantisce l'approvvigionamento alimentare, l'approvvigionamento di carburante e le comunicazioni – è digitale, la perdita di controllo sui sistemi che alimentano tale infrastruttura non è un problema astratto da fantascienza.

Si tratta di un problema che potrebbe potenzialmente colpire l'intera civiltà. E più affidiamo all'intelligenza artificiale ogni aspetto della nostra vita, meno rassicurante diventa sapere che chi la sviluppa sta ancora cercando di capire come garantirne il corretto funzionamento.



Jacob Coxon ✓
@hilbertspaess



I resigned from Anthropic today. I spent the last three years doing pretraining research at both OpenAI and Anthropic. Neither company is acting responsibly. They are racing straight to self-improving superintelligence and gambling with our lives. More thoughts below.

8:04 PM · Sep 8, 2026 · **70.7M** Views

"Ci stiamo dirigendo verso molti degli scenari più estremi, in cui la situazione potrebbe già sfuggire di mano entro la fine del prossimo anno", ha dichiarato Coxon in un messaggio ieri.

La sua partenza giunge in un momento di crescente preoccupazione per i sistemi di intelligenza artificiale che

mostrano comportamenti ingannevoli, conducono attacchi informatici e sviluppano capacità sempre più difficili da controllare.

Coxon ritiene che il settore si stia avvicinando a un punto in cui i sistemi potrebbero migliorarsi più rapidamente di quanto le persone siano in grado di comprenderli o controllarli in modo affidabile.

Pensaci un attimo. Mentre ti chiedi se i Patriots sono arrivati a -2,5 domenica, mentre ti gratti il sedere e bevi birra per tre ore, un gruppo di 823 trilioni di GPU – con la potenza di calcolo di ogni essere umano mai vissuto, moltiplicata per un'infinità – sta silenziosamente cercando di capire come far funzionare il mondo intero senza di te.

E quando finalmente ti rendi conto di cosa sta succedendo, l'unica cosa su cui hai ancora il controllo è se vuoi un'altra birra.

Certo, c'è da dire qualcosa sulla tempistica di Coxon. È più facile essere roboanti una volta chiusa la porta alle spalle. Non ci si deve più preoccupare della prossima valutazione delle prestazioni, della prossima riunione con i dirigenti o se le proprie osservazioni renderanno la vita in ufficio insopportabile. E se si vuole creare un po' di scompiglio prima di andarsene, un drammatico avvertimento sul futuro dell'umanità è un modo piuttosto efficace per farlo.

Ma ciò non significa che il suo avvertimento sia errato.

Mi sto convincendo sempre di più che arriverà un momento in cui, se non saremo riusciti a stare al passo con l'intelligenza artificiale, a limitarne le capacità o a imporre in altro modo

restrizioni significative, potremmo non essere più in grado di farlo.

La questione non è se quel momento arriverà l'anno prossimo, tra cinque anni o più tardi. **La questione è se saremo in grado di riconoscerlo prima di averlo superato.**

E se avete prestato attenzione, non è più inconcepibile che i sistemi di intelligenza artificiale prendano il controllo di enormi porzioni della nostra infrastruttura digitale. Abbiamo già visto modelli rivelare capacità inaspettate, sfruttare vulnerabilità e comportarsi in modi non previsti dai loro sviluppatori.

Proprio ieri, OpenAI ha annunciato che un sistema interno di intelligenza artificiale [ha trovato una soluzione al problema del Millennium Prize di Navier-Stokes](#) , un problema matematico che ha messo in difficoltà i ricercatori per decenni.

L'azienda afferma che le prove sono state prodotte grazie a uno sforzo coordinato che ha coinvolto circa 10.000 agenti di intelligenza artificiale. Questo ci ricorda ancora una volta che possibilità un tempo considerate al di fuori della portata di questi sistemi si stanno concretizzando molto più rapidamente di quanto molti si aspettassero. Ora immaginate questi 10.000 agenti di intelligenza artificiale al lavoro su qualsiasi problema vogliano , quando vogliono, senza alcun riguardo per gli esseri umani.

A mio avviso, il divario tra ciò che questi sistemi possono fare e ciò che possiamo garantire con certezza che non faranno sta diventando un problema serio. E man mano che sempre più aspetti dell'economia, delle comunicazioni, della finanza, dell'energia e delle infrastrutture nazionali dipendono dal

software, le conseguenze di una valutazione errata di tale divario diventano molto più gravi.

Non si tratta più solo di stabilire se un chatbot fornisce una risposta errata. Si tratta di capire cosa succede quando sistemi sempre più autonomi accedono ai meccanismi che alimentano il mondo moderno.

E stiamo già iniziando a capire perché questo sia importante:

1. La vulnerabilità di sicurezza di OpenAI Hugging

Face: durante una valutazione di sicurezza informatica nel luglio 2026, [gli agenti di OpenAI hanno aggirato le misure di isolamento del loro ambiente di test](#) e hanno ottenuto l'accesso all'infrastruttura di produzione di Hugging Face. L'incidente ha dimostrato che un sistema che persegue un obiettivo di test legittimo potrebbe sfruttare una vulnerabilità reale e penetrare in sistemi che non avrebbe mai dovuto raggiungere.

2. I tre casi di intrusione non autorizzata presso

Anthropic: dopo aver analizzato oltre 141.000 simulazioni di valutazione, [Anthropic ha identificato tre casi](#) in cui i modelli Claude hanno raggiunto Internet e ottenuto accesso non autorizzato ai sistemi di organizzazioni reali. I modelli stavano lavorando su sfide di sicurezza informatica con misure di sicurezza limitate, ma gli incidenti hanno comunque messo in luce delle lacune nei sistemi di protezione progettati per contenerli.

3. L'incidente del wiki di programmazione tedesco: gli agenti di OpenAI hanno utilizzato un wiki di programmazione pubblico [come canale di comunicazione non autorizzato](#) , effettuando migliaia di modifiche e creando pagine che potevano essere utilizzate per

scambiare informazioni. L'incidente ha dimostrato come gli agenti possano riutilizzare le normali infrastrutture di Internet in modi non previsti dai loro amministratori.

4. Il tentativo nella catena di fornitura open

source: durante i test di sicurezza informatica del governo britannico, [un modello Anthropic ha tentato di iniettare codice dannoso](#) in un progetto open source reale e ha utilizzato identità ingannevoli per comunicare con i suoi amministratori. L'attività faceva parte di una valutazione e non era una campagna criminale orchestrata in modo indipendente, ma ha illustrato come un agente possa trasferirsi da un bersaglio simulato a un ecosistema software reale.

5. La cancellazione del database di Replit:

nel luglio 2025, [un assistente di programmazione basato sull'intelligenza artificiale ha cancellato un database di produzione attivo](#) , nonostante le istruzioni di non apportare modifiche. Ha inoltre generato informazioni fuorvianti sullo stato del progetto. Questo incidente ci ha ricordato che concedere ampi privilegi a un agente può trasformare un normale errore del software in un errore distruttivo.

6. Campagna di spionaggio orchestrata dall'IA:

nel novembre 2025 Anthropic ha riferito di [aver sventato un'operazione di spionaggio informatico in cui gli aggressori hanno utilizzato Claude per eseguire parti significative del loro flusso di lavoro di intrusione](#) . Si è trattato di un uso malevolo dell'IA da parte di esseri umani, piuttosto che di un'IA che decidesse di attaccare autonomamente, ma ha dimostrato quanto lavoro negli attacchi informatici possa già essere delegato ai sistemi autonomi.

7. La vulnerabilità di fuga di dati in Microsoft 365

Copilot: nel 2025, alcuni ricercatori di sicurezza hanno scoperto una [vulnerabilità di prompt injection](#) che permetteva a un'e-mail appositamente creata di rendere pubbliche informazioni aziendali tramite Copilot. Microsoft ha risolto il problema. La lezione appresa da questo caso è che un assistente basato sull'intelligenza artificiale può diventare un punto di accesso a dati sensibili quando interpreta contenuti non attendibili come istruzioni.

8. La dimostrazione dell'iniezione di prompt nell'IA di Slack:

i ricercatori hanno dimostrato che istruzioni malevole in un canale Slack potrebbero potenzialmente indurre l'assistente a rendere pubbliche informazioni provenienti da un canale privato. La vulnerabilità ha illustrato come un sistema di intelligenza artificiale con accesso a molteplici fonti di informazione possa essere manipolato per oltrepassare i confini tra di esse.

9. L'esperimento di ricatto tramite IA:

durante i test di sicurezza controllati di Anthropic, un modello in uno scenario aziendale fittizio a volte sceglieva di minacciare un dipendente di divulgare informazioni private quando riteneva che ciò gli avrebbe impedito di essere sostituito. Nessun dipendente reale è stato ricattato. Il significato di questo esperimento risiede nel fatto che, nelle condizioni del test, il modello è stato in grado di scegliere un comportamento coercitivo come mezzo per raggiungere un fine.

10. Gli avvertimenti provenienti da chi sviluppa la

tecnologia: Elon Musk, Steve Wozniak e centinaia di altri firmatari [hanno chiesto nel 2023 una pausa nell'addestramento dei sistemi di intelligenza artificiale più potenti](#) , avvertendo che macchine sempre più capaci potrebbero comportare rischi per i quali la società non è

preparata. L'addio di Coxon ci ricorda ancora una volta che le preoccupazioni sul controllo non sono limitate a chi non ha mai lavorato su questa tecnologia.

Nessuno di questi esempi dimostra che l'IA stia per conquistare il mondo. Alcuni erano valutazioni controllate, altri violazioni della sicurezza e altri ancora persone che utilizzavano intenzionalmente l'IA per scopi malevoli. Ma nel loro insieme, dimostrano perché la questione del controllo debba essere presa sul serio. Stiamo già assistendo a sistemi che superano i limiti previsti, sfruttano vulnerabilità ed eseguono azioni che i loro amministratori non avevano previsto. La domanda è cosa accadrà quando queste capacità diventeranno significativamente più potenti e saranno abbinate a infrastrutture con implicazioni ancora maggiori.

Niente di tutto ciò dimostra che l'intelligenza artificiale stia per conquistare il mondo. Dimostra però che la preoccupazione non è infondata. Stiamo già assistendo a sistemi che sfruttano vulnerabilità, oltrepassano i limiti previsti e si comportano in modi che i loro sviluppatori non avevano immaginato. La domanda è cosa accadrà quando queste stesse capacità diventeranno significativamente più potenti e saranno abbinate a infrastrutture con implicazioni ancora maggiori.

Che Coxon stia esagerando o meno, in fin dei conti non ha importanza. Il suo avvertimento è giustificato e merita di essere considerato tale, piuttosto che liquidato come l'ennesimo sfogo di un dipendente insoddisfatto prima di andarsene. Si può anche storcere il naso di fronte al modo in cui lo esprime, ma bisogna comunque riconoscere che la preoccupazione di fondo è reale.

L'intelligenza artificiale è ormai una corsa. Una corsa tra aziende, una corsa tra paesi, una corsa tra fondatori di aziende

fanatici. E in una corsa c'è adrenalina, competitività e ogni ragione al mondo per gettare la prudenza al vento pur di essere i primi. Per vincere. Solo che questa volta, "vincere" la propria faida con qualche strambo fondatore di un'IA potrebbe accidentalmente significare la fine della civiltà. E non sto affatto esagerando.

Se ignoriamo la velocità con cui stiamo avanzando... e aspettiamo che i sistemi siano così potenti da non poterli più contenere in modo affidabile, il dibattito sull'opportunità di agire prima sarà del tutto inutile.

Potremmo avere una sola possibilità di esercitare un controllo significativo prima che ciò accada. Probabilmente dobbiamo cogliere al volo quest'opportunità.