

Perché i progetti di agentic AI falliscono (e come farli funzionare davvero)

8 Settembre 2026

Un'analisi di ciò che ancora separa la sperimentazione dall'adozione nelle imprese

Portare gli agenti di intelligenza artificiale dalla sperimentazione all'uso concreto in azienda è una sfida che va ben oltre la scelta della tecnologia: entrano in gioco processi, qualità e accessibilità dei dati, integrazione con i sistemi esistenti, governance e competenze e molto altro. Il passaggio da un prototipo efficace a una soluzione affidabile e sostenibile richiede quindi di considerare insieme dimensioni tecnologiche e organizzative. **Massimiliano Luca**, ricercatore del centro Augmented Intelligence, e **Alberto Purinan**, del Corporate Relationship Office di FBK, mettono a confronto la prospettiva della ricerca e quella del rapporto con le imprese per analizzare gli ostacoli all'adozione degli agenti e le condizioni necessarie per generare valore concreto

Un **agente di intelligenza artificiale** è un sistema progettato per perseguire un obiettivo attraverso una sequenza di attività. Per esempio, un AI agent può raccogliere informazioni, valutare i passaggi necessari, consultare applicazioni e basi dati, generare un risultato o avviare azioni entro limiti definiti.

Consideriamo un'azienda manifatturiera. Alla domanda: "Quanti microchip abbiamo in magazzino?", un chatbot si limita a fornire il dato aggiornato. Un **agente**, invece, può riconoscere che quel numero è sceso sotto la soglia minima e, anziché limitarsi alla risposta, agisce di conseguenza: verifica il budget disponibile sul gestionale, confronta i prezzi tramite le API dei fornitori, precompila l'ordine d'acquisto e lo sottopone al responsabile per l'approvazione finale. L'agente in questo esempio agisce entro un perimetro ben definito di deleghe e controlli progettati su misura, facendosi carico di attività circoscritte e sempre verificabili. Il suo vero valore emerge soprattutto nei processi aziendali più articolati, che richiedono di orchestrare fonti informative eterogenee, snodare una sequenza variabile di passaggi e adattarsi con flessibilità a casi mai del tutto identici.

Oggi sul mercato sono presenti sia molte soluzioni già pronte all'uso per lo più generaliste, sia molte sperimentazioni a livello di "*proof of concept*", mentre il passaggio alla messa in produzione su scala resta una sfida aperta. Secondo l'indagine State of AI 2025 di McKinsey, il 23% delle organizzazioni intervistate sta estendendo almeno un sistema di agentic AI, mentre il 39% è ancora nella fase sperimentale. La quota di organizzazioni che dichiara di aver portato con successo l'agente a regime non supera il 10%. Anche se i dati citati sono del 2025 e la tecnologia

corre veloce, i motivi di tale scoglio di messa a regime/adozione permangono. I “proof of concept” infatti funzionano perché operano spesso su dati selezionati e percorsi lineari in un ambiente protetto. Le soluzioni commerciali generaliste, invece, devono fare i conti con la realtà aziendale: permessi di accesso diversi da un reparto all’altro, sistemi legacy e archivi non sempre completi o organizzati. Il risultato è che spesso mancano di trasparenza, faticano a integrarsi con i dati aziendali e rendono difficile prevedere i costi. Per superare questi scogli occorre che tecnologie e organizzazioni crescano insieme.

Uno dei principali errori strategici è quindi trattare l’agentic AI come un prodotto da installare, anziché come un **processo da costruire** insieme all’organizzazione. Per quanto avanzata, nessuna tecnologia può conoscere da sé le eccezioni, le regole informali e i criteri di qualità che ogni azienda si porta dietro. Prima di introdurre un agente è quindi essenziale un lavoro preparatorio che definisca gli obiettivi strategici, ridisegni i flussi di lavoro e stabilisca quali attività si possono delegare, chi interviene nei casi anomali e con quali indicatori si misura il successo. Senza questa analisi si corre il rischio di automatizzare le inefficienze senza accorgersene. L’agente può infatti funzionare fin troppo bene su un processo che, in realtà, non andava replicato. Gli output dei sistemi di IA si presentano sempre in forma ordinata e convincente, ma la cura della forma rischia di mascherare l’incorrettezza del contenuto. Sviluppare questa capacità critica nelle funzioni operative richiede formazione e gestione del cambiamento come condizioni necessarie e fondanti perché l’innovazione produca il risultato atteso.

Un buon Large Language Model da solo non fa un buon agente. Le **difficoltà tecniche** nascono soprattutto dal modo in cui l’agente lo usa nel processo.

Un primo problema, comune anche alla dimensione organizzativa, è l’**errore cumulativo**. Un agente può eseguire correttamente diversi passaggi e produrre comunque un risultato inadeguato se interpreta male un’informazione, usa una fonte sbagliata o sceglie lo strumento improprio. Servono quindi validazioni e punti di controllo nei passaggi critici.

Un secondo tema è la **tracciabilità**, ossia capire perché un agente ha prodotto un certo risultato. Serve registrare fonti, istruzioni e strumenti usati, così da correggere gli errori e rispondere a esigenze di audit e sicurezza.

Anche la **memoria** va progettata con attenzione perché le informazioni aziendali non possono dipendere dal ricordo di conversazioni passate, ma devono provenire da fonti aggiornate e governate, con permessi definiti.

Infine c’è la scelta del **modello più adatto** al compito. I grandi modelli generalisti offrono ampie capacità ma sono spesso poco efficienti su attività ripetitive o specialistiche, dove modelli più piccoli garantiscono risposte più rapide e costi contenuti. La soluzione migliore spesso combina più approcci tra modelli linguistici, ricerca documentale, basi dati e regole deterministiche. Il principio guida resta uno, ovvero usare solo la complessità necessaria al risultato richiesto.

Nei prossimi tre anni è realistico attendersi un passaggio da strumenti isolati a servizi coordinati, con ruoli definiti e interazioni controllate tra più componenti software. I sistemi **multi-agent**, in cui più agenti specializzati collaborano e si coordinano per raggiungere un obiettivo comune, saranno utili quando esisterà una chiara divisione dei compiti: per esempio, una componente per la ricerca documentale, una per la verifica delle regole e una per la preparazione di una proposta da validare. Senza questa distribuzione funzionale, l’aumento del numero di agenti rischia di aumentare costi e

complessità senza migliorare il risultato.

Il lavoro delle persone si concentrerà maggiormente sulla supervisione, sulla gestione delle eccezioni, sulla qualità dei dati e sulla progettazione dei processi. Un agente potrà ridurre il tempo dedicato alla raccolta e al trattamento preliminare delle informazioni, ma ogni decisione su cosa sia un risultato accettabile, quale eccezione meriti un trattamento diverso e dove il processo vada corretto continuerà a dipendere da competenze di dominio e responsabilità professionale che il sistema non possiede.

Di fronte a un mercato che offre soluzioni generaliste pronte all'uso ma fatica a portare le sperimentazioni in produzione, il valore distintivo di un centro di ricerca come FBK non è fornire strumenti preconfezionati, ma **avanzare la frontiera della conoscenza con soluzioni specializzate e su misura**, in un processo che fa crescere insieme tecnologia e organizzazione. L'impresa mette in campo i propri problemi concreti, i dati e i vincoli operativi. FBK mette competenze scientifiche, architetture specializzate e metodi di valutazione rigorosi. L'obiettivo è fornire le evidenze e le competenze necessarie a portare in produzione sistemi efficienti, integrati con i dati aziendali e pienamente rispettosi dei requisiti di conformità e sovranità. È questa collaborazione tra ricerca e impresa a fare la differenza tra un progetto di agentic AI che fallisce e uno che funziona davvero.

LINK

<https://magazine.fbk.eu/it/news/perche-i-progetti-di-agentic-ai-falliscono-e-come-farli-funzionare-davvero/>

TAG

- #agenti
- #agentic ai
- #aziende
- #intelligenzaartificiale
- #intelligenzaaugmentata
- #llm
- #mobs
- #multi-agent

AUTORI

- Massimiliano Luca
- Alberto Purinan