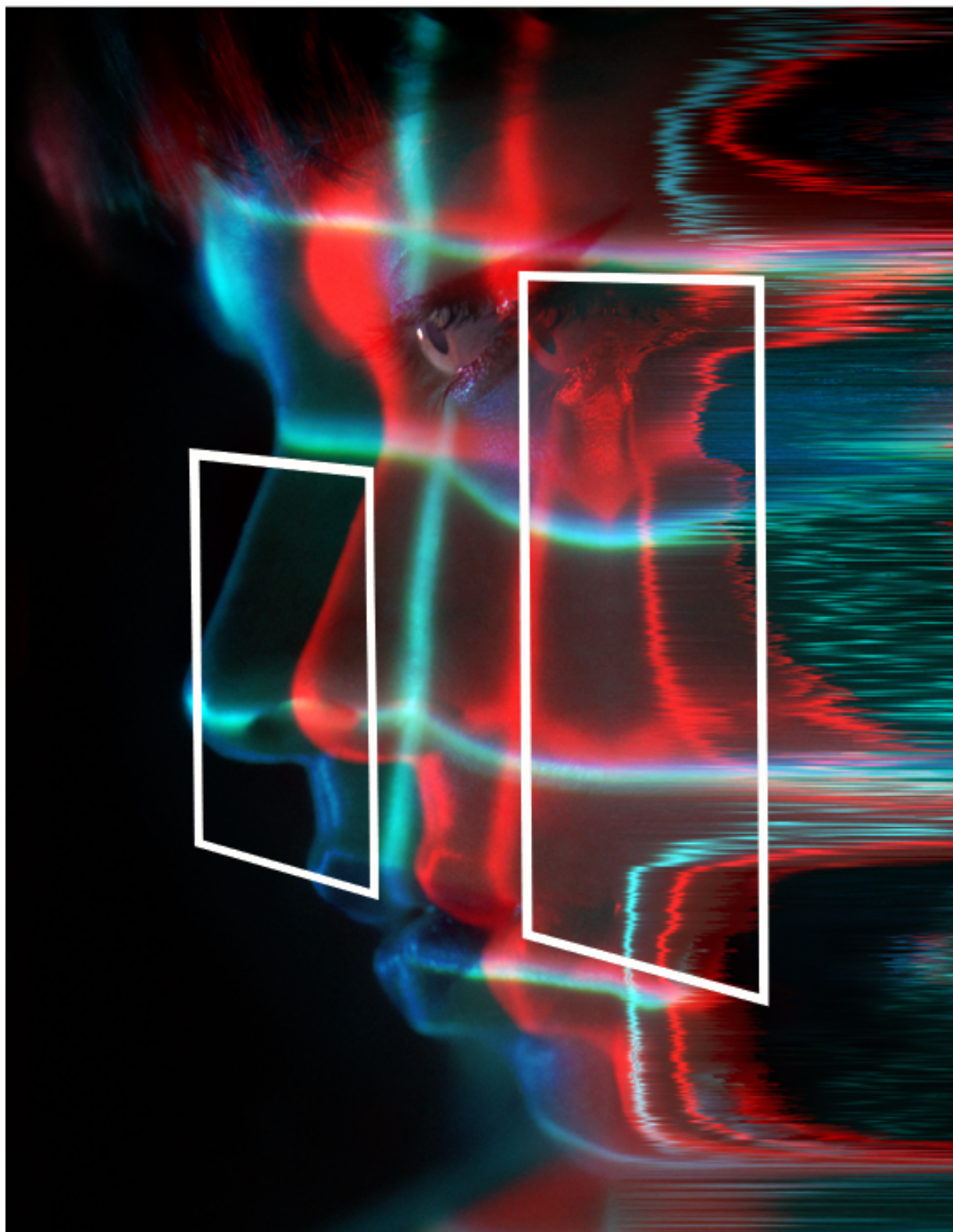


## TEAM AUMENTATI

La collaborazione tra umani e agenti IA nei team di lavoro



INDICE

**INTRODUZIONE .....5**

**1.TECNOLOGIA.....**

**1.1 EVOLUZIONE DELL’IA AGENTICA: DEFINIZIONI, TAPPE DI SVILUPPO E STATO DELL’ARTE ..... 6**

1.1.1 Una definizione di IA agentica.....6

1.1.2 Tappe di sviluppo e stato dell’arte.....6

**1.2 DALL’IA GENERATIVA ALL’IA AGENTICA: DIFFERENZE FUNZIONALI, LIMITI E SINERGIE ..... 8**

1.2.1 LLM tradizionali .....8

1.2.2 Modelli di reasoning .....8

1.2.3 Gli agenti IA.....9

**1.3 GLI AGENTI IA: AUTONOMIA, SOCIAL ABILITY, REATTIVITÀ, PROATTIVITÀ ..... 11**

1.3.1 Autonomia ..... 11

1.3.2 Social ability ..... 11

1.3.3 Reattività ..... 12

1.3.4 Proattività ..... 12

**1.4 MATURITÀ DI MERCATO E DIFFUSIONE SETTORIALE: BENCHMARK GLOBALI E CURVE DI ADOZIONE..... 14**

1.4.1 Maturità di mercato e benchmark globali ..... 14

1.4.2 Diffusione settoriale: dove gli agenti entrano per primi ..... 14

1.4.3 Dalle curve di adozione alla “J-curve” della produttività..... 15

**2.LAVORO IN SQUADRA E PERFORMANCE .....**

**2.1 LA PERFORMANCE DEI TEAM: FRAMEWORK TEORICI E VALUTAZIONI SOCIO-TECNICHE ..... 17**

2.1.1 La prospettiva sociotecnica: che cosa guarda e perché è necessaria ..... 17

2.1.2 IMOI e Hackman/TDS per la performance dei team aumentati ..... 18

**2.2 RUOLI IBRIDI E TASK-ALLOCATION: MODELLI DI ORCHESTRAZIONE UMANO-IA ..... 20**

2.2.1 Allocazione per natura del compito: oltre la seniority.....20

2.2.2 Pattern di orchestrazione: triage, specialisti, escalation .....21

2.2.3 Nuovi ruoli organizzativi: ownership, supervisione, compliance .....21

|                                                                                                             |           |
|-------------------------------------------------------------------------------------------------------------|-----------|
| <b>2.3 COMUNICAZIONE, COORDINAMENTO, MOTIVAZIONE E FIDUCIA NEI TEAM MISTI.....</b>                          | <b>23</b> |
| 2.3.1 Comunicazione.....                                                                                    | 23        |
| 2.3.2 Coordinamento.....                                                                                    | 23        |
| 2.3.3 Motivazione .....                                                                                     | 24        |
| 2.3.4 Fiducia: tra affidamento e relazione .....                                                            | 24        |
| 2.3.5 Una sfida per la progettazione socio-tecnica .....                                                    | 25        |
| <b>2.4 METRICHE INTEGRATE DI PERFORMANCE: OUTPUT QUANTITATIVI E DINAMICHE QUALITATIVE .....</b>             | <b>26</b> |
| 2.4.1 Inputs (condizioni di partenza del team) .....                                                        | 26        |
| 2.4.2 Mediators (processi di team e stati emergenti) .....                                                  | 26        |
| 2.4.3 Outputs (Performance, outcome umani e sostenibilità del team) .....                                   | 27        |
| 2.4.4 Nuovi input per il ciclo successivo.....                                                              | 28        |
| <b>3.OSSERVAZIONI SUL CAMPO .....</b>                                                                       |           |
| <b>3.1 SINTESI DELLE EVIDENZE EMPIRICHE: META-ANALISI DI STUDI ACCADEMICI SU TEAM UMANO-IA .....</b>        | <b>29</b> |
| <b>3.2 ESPERIMENTI AZIENDALI DOCUMENTATI: CASI FORTUNE 500, PMI E SETTORE PUBBLICO.....</b>                 | <b>32</b> |
| <b>3.3 LA ACTION-RESEARCH CON CDP: DISEGNO, INTERVENTO E RISULTATI.....</b>                                 | <b>35</b> |
| 3.3.1 Questionario pre-sperimentazione .....                                                                | 36        |
| 3.3.2 Questionario post-sperimentazione .....                                                               | 36        |
| 3.3.3 Risultati qualitativi before-after .....                                                              | 37        |
| 3.3.4 Risultati quantitativi .....                                                                          | 39        |
| <b>3.4 PATTERN RICORRENTI E VARIABILI CONTINGENTI: CONVERGENZE, DIVERGENZE E GAP DI RICERCA.....</b>        | <b>42</b> |
| <b>4.OPPORTUNITÀ EMERGENTI .....</b>                                                                        |           |
| <b>4.1 CREAZIONE DI VALORE ECONOMICO: EFFICIENZA, DECISION-MAKING SUPERIORE, INNOVAZIONE CONTINUA .....</b> | <b>43</b> |
| 4.1.1 Dove si genera realmente valore .....                                                                 | 43        |
| 4.1.2 Efficienza non significa solo risparmio di tempo .....                                                | 43        |
| 4.1.3 La qualità delle decisioni dipende dall'orchestrazione .....                                          | 44        |
| 4.1.4 Innovazione continua, ma con una domanda più esigente .....                                           | 44        |
| <b>4.2 DOMINI AD ALTO POTENZIALE: CRITERI DI INDIVIDUAZIONE, DRIVER DI ADOZIONE .</b>                       | <b>45</b> |
| 4.2.1 Non i domini “più digitali” .....                                                                     | 45        |
| 4.2.2 Il secondo criterio: il sovraccarico .....                                                            | 45        |
| 4.2.3 Il vero salto: interoperabilità e governance .....                                                    | 46        |
| 4.2.4 Nei domini regolati l'opportunità cresce solo se cresce la verificabilità .....                       | 46        |
| 4.2.5 Il driver più decisivo è la capacità di sperimentare .....                                            | 46        |

|                                                                                                                   |           |
|-------------------------------------------------------------------------------------------------------------------|-----------|
| <b>4.3 EVOLUZIONE DEI RUOLI E DELLE COMPETENZE: CAPABILITY EMERGENTI E REQUISITI SOCIO-TECNICI.....</b>           | <b>47</b> |
| 4.3.1 Il lavoro quotidiano si sposta: meno esecuzione lineare, più regia.....                                     | 47        |
| 4.3.2 Le capability decisive sono meta-cognitive.....                                                             | 47        |
| 4.3.3 Cosa resterà umano: la traiettoria sta emergendo.....                                                       | 47        |
| 4.3.4 Gli investimenti giusti non sono solo tecnologici.....                                                      | 48        |
| <b>4.4 ORIZZONTI DI INNOVAZIONE: RETI DI AGENTI COLLABORATIVI, SYMBIOTIC TEAMS, ORGANIZZAZIONI ADATTIVE .....</b> | <b>50</b> |
| 4.4.1 Il passaggio in corso: dal copilota al workflow agentic.....                                                | 50        |
| 4.4.2 Lo stack aperto che sta emergendo conta quasi quanto i modelli.....                                         | 50        |
| 4.4.3 Anche l'hardware sta cambiando il perimetro del possibile .....                                             | 51        |
| 4.4.4 Che cosa si intravede per team e strutture .....                                                            | 51        |
| 4.4.5 La domanda finale, forse, è la più rilevante .....                                                          | 51        |
| <b>Conclusioni.....</b>                                                                                           | <b>53</b> |
| <b>Ringraziamenti.....</b>                                                                                        | <b>54</b> |
| <b>Principali riferimenti bibliografici .....</b>                                                                 | <b>55</b> |

## INTRODUZIONE

L'ingresso degli agenti di intelligenza artificiale nei team di lavoro produce esiti che cambiano sensibilmente da un'organizzazione all'altra. L'esperienza sul campo mostra un dato ricorrente: le aspettative iniziali, spesso nette in un senso o nell'altro, lasciano in fretta spazio a una valutazione più equilibrata. La differenza, nei casi osservati, dipende più da come l'agente entra nel lavoro che dalla tecnologia in sé.

Questo paper nasce da una collaborazione tra Cassa Depositi e Prestiti (CDP), l'Università di Verona e Cocoon Pro, con l'obiettivo di osservare in un contesto istituzionale italiano cosa accade quando un agente entra nel processo di risk assessment di una funzione di Internal Audit.

La domanda di ricerca è concreta: in quali condizioni l'introduzione di agenti IA in team esistenti produce effetti positivi sulla performance e sul benessere dei membri? In base all'evidenza raccolta, il valore dipende dalla qualità del sistema organizzativo che accoglie l'agente. Gli agenti amplificano le condizioni preesistenti. Dove il lavoro è ben progettato, i risultati migliorano. Dove non lo è, inefficienze, ambiguità e criticità tendono ad amplificarsi.

Da qui una conseguenza che merita di essere esplicitata: il tema è organizzativo prima ancora che tecnologico. L'agente cambia la distribuzione del lavoro cognitivo, i punti di controllo, i criteri di responsabilità e le competenze richieste. Trattare l'IA come un semplice strumento da aggiungere a processi già dati è tra le criticità più ricorrenti nelle sperimentazioni attuali. Per la stessa ragione, il valore del contributo umano va reso esplicito: l'IA non sostituisce le capacità distintivamente umane, ne ridefinisce il ruolo nel lavoro.

Il lavoro esplora questa domanda su tre livelli: una rassegna delle evidenze empiriche disponibili, un'analisi dei framework teorici più adeguati a interpretarle e una sperimentazione esplorativa condotta in un contesto istituzionale reale.

La sperimentazione, svolta in CDP nel mese di marzo 2026, ha coinvolto sette professionisti della funzione Internal Audit nel processo di risk assessment. L'esperienza mostra che l'IA può generare benefici reali in attività concrete, ma solo quando è accompagnata da regole d'uso, supervisione coerente e apprendimento progressivo. Il giudizio professionale, lontano dal perdere valore, diventa più importante nei momenti di validazione, interpretazione e assunzione di responsabilità.

Per i decisori il problema ormai è operativo: con quali priorità, con quali responsabilità e con quali meccanismi di governance integrare questi sistemi nei processi.

## 1.1 EVOLUZIONE DELL'IA AGENTICA: DEFINIZIONI, TAPPE DI SVILUPPO E STATO DELL'ARTE

### 1.1.1 Una definizione di IA agentica

---

Per IA agentica si intendono sistemi basati su Large Language Models (LLM) che operano come entità autonome capaci di pianificare, memorizzare e accedere a strumenti e fonti esterne per perseguire obiettivi complessi attraverso sequenze di ragionamento. A differenza dei chatbot tradizionali che rispondono a ogni prompt separatamente, gli agenti scompongono i problemi complessi in compiti gestibili, consultano fonti, eseguono codice e raffinano iterativamente la strategia, avvicinandosi al ruolo di veri e propri membri del team<sup>1</sup>.

#### Key insight

La differenza rispetto a un chatbot è che un agente porta avanti un obiettivo rendendolo, sul piano organizzativo, più simile a un membro del team che a uno strumento.

### 1.1.2 Tappe di sviluppo e stato dell'arte

---

Gli agenti IA, come li conosciamo oggi, sono di recentissima introduzione, affermatasi a fine 2024 a partire da proto-agenti sviluppati nel 2023 con funzionalità limitate. Ed è per questo che è ancora complesso delineare con precisione le traiettorie di evoluzione. In questo contesto è però rilevante osservare come due importanti protocolli aperti, emersi nel corso di pochi mesi, abbiano contribuito a definire le caratteristiche distintive degli agenti IA rispetto ai precedenti modelli LLM<sup>2</sup>.

Il primo è il Model Context Protocol (MCP), introdotto da Anthropic nel novembre 2024, e rapidamente adottato dai principali laboratori di IA. Questo protocollo consente di stabilire connessioni sicure e bidirezionali tra agenti di IA e risorse, strumenti e dati di terze parti. Il protocollo opera su un'architettura client-server, in cui i server MCP espongono risorse, strumenti e fonti di dati, mentre i client MCP (tipicamente agenti di IA o applicazioni) vi si connettono per accedere alle funzionalità disponibili.

Il secondo è il protocollo A2A (Agent2Agent), lanciato da Google nell'aprile 2025 insieme a numerosi partner industriali. Questo protocollo permette agli agenti di IA di comunicare tra loro, scambiare informazioni in modo sicuro e coordinare azioni su varie piattaforme o applicazioni aziendali. Abilita, in sostanza, la cooperazione tra diversi agenti.

---

<sup>1</sup> Microsoft, Agents are here—is your company prepared? (WorkLab / 2026).

<sup>2</sup> <https://www.pwc.com/m1/en/publications/documents/2024/agent-ai-the-new-frontier-in-genai-an-executive-playbook.pdf>

I due protocolli differenziano gli agenti IA dagli LLM tradizionali e dai modelli di reasoning (modelli capaci di ragionamento strutturato) per due ragioni: estendono la conoscenza dell'agente oltre i dati di training, attraverso risorse esterne, e gli permettono di interagire con altri agenti<sup>3</sup>.

### **Key Insight**

Il Model Context Protocol (MCP) estende l'accesso dell'agente al mondo informativo, mentre il protocollo A2A (Agent2Agent) amplia la capacità di coordinarsi con altri attori artificiali. Insieme, questi protocolli trasformano l'IA da sistema cognitivo isolato a **parte di un sistema di azione collettiva**.

---

<sup>3</sup> <https://www.deloitte.com/content/dam/assets-shared/docs/about/2025/state-of-ai-2026-global.pdf>

## 1.2 DALL'IA GENERATIVA ALL'IA AGENTICA: DIFFERENZE FUNZIONALI, LIMITI E SINERGIE

Gli agenti rispecchiano, ad oggi, il punto più avanzato di un'evoluzione partita dai modelli standard (rapidi ma non riflessivi), passata per i modelli di reasoning (capaci di pensiero deliberativo) e approdata, infine, agli agenti di Intelligenza Artificiale capaci di agire e usare strumenti in autonomia. I sistemi LLM tradizionali sono entrati sul mercato con il lancio di ChatGPT nel novembre del 2022. I modelli di reasoning e gli agenti sono stati invece introdotti alla fine del 2024. Ogni modello offre caratteristiche distinte adatte a contesti e obiettivi differenti.

### Key Insight

Gli **LLM tradizionali** generano risposte. I **modelli di reasoning** risolvono problemi complessi. Gli **agenti IA** pianificano e agiscono in autonomia accedendo a risorse esterne e comunicando con altri agenti. La differenza è sia tecnica sia **funzionale e organizzativa**.

### 1.2.1 LLM tradizionali

---

Gli LLM tradizionali rappresentano tuttora il paradigma fondamentale dell'IA generativa. Generano contenuti in linguaggio naturale attraverso la previsione *token-per-token* (ossia un'unità di testo alla volta, dove un token corrisponde all'incirca a una parola o a una sua parte) basata sui dati di addestramento e sul fine-tuning. Questo processo consente a questi sistemi di sviluppare rappresentazioni astratte dei concetti e delle loro relazioni, costruendo modelli impliciti del mondo, che permettono di generare testo coerente e contestualmente appropriato. Molti LLM moderni sono multimodali: elaborano non solo testo ma anche immagini, video e suoni. Restano però sistemi reattivi. Elaborano il testo in modo unidirezionale senza la capacità di riflettere e rivedere l'output precedente, sono vincolati alla data di chiusura del loro addestramento senza accesso a informazioni più recenti e faticano con compiti che richiedono ragionamento logico articolato o passaggi matematici complessi<sup>4</sup>.

### 1.2.2 Modelli di reasoning

---

I modelli di reasoning nascono per colmare le carenze analitiche degli LLM tradizionali attraverso processi di pensiero strutturato. Questi modelli sono ancora predittori *token-per-token*, ma vengono addestrati tramite *reinforcement learning* a ragionare in modo più deliberato, identificando e correggendo errori. Questo li rende efficaci in compiti che mettevano in difficoltà i modelli precedenti, come la risoluzione di problemi matematici articolati o il debug di codici complessi. I modelli di

---

<sup>4</sup> Korinek, Anton. 2025. *AI Agents for Economic Research: August 2025 Update to 'Generative AI for Economic Research: Use Cases and Implications for Economists'*, published in the Journal of Economic Literature 61(4).

reasoning presentano anche delle limitazioni. Il loro processo deliberativo richiede risorse computazionali significativamente maggiori, il che li rende più lenti e più costosi. Sebbene eccellano in problemi ben definiti con strutture logiche chiare, offrono pochi vantaggi rispetto ai LLM tradizionali per compiti che richiedono principalmente fluidità linguistica o generazione creativa.

### 1.2.3 Gli agenti IA

---

Gli agenti IA rappresentano la più recente configurazione operativa dell'IA generativa. Si distinguono per la capacità di agire in autonomia e di richiamare strumenti esterni per raggiungere obiettivi specifici. A differenza degli LLM tradizionali, non rispondono a prompt isolati, ma pianificano sequenze di azioni, raccolgono informazioni da più fonti e perfezionano il proprio approccio sulla base dei risultati intermedi. Questo paradigma è reso possibile dalla combinazione di finestre di contesto più ampie, migliore capacità di ragionamento e integrazione di strumenti esterni. In pratica, un agente può effettuare ricerche sul web, interagire con database, eseguire codice, manipolare file e controllare interfacce digitali. L'introduzione delle "connessioni" con applicazioni e servizi esterni trasforma l'IA da sistema limitato ai dati di addestramento a uno capace di operare sull'intero ecosistema digitale dell'utente.

#### **Key Insight**

Con gli agenti il salto vero è "fare": la posta in gioco si sposta dalla qualità della risposta alla governabilità dell'azione.

La Tabella 1 riassume le principali differenze tra i modelli tradizionali, i modelli di reasoning e gli agenti IA.

|                               | LLM tradizionali                          | Modelli di reasoning                                             | Agenti IA                                                |
|-------------------------------|-------------------------------------------|------------------------------------------------------------------|----------------------------------------------------------|
| Funzione principale           | Generare risposte linguistiche coerenti   | Risolvere problemi complessi attraverso ragionamento strutturato | Pianificare ed eseguire azioni per raggiungere obiettivi |
| Modalità operativa            | Reattiva: risponde a un prompt alla volta | Deliberativa: ragiona passo dopo passo                           | Proattiva: pianifica, agisce, verifica e corregge        |
| Memoria tra interazioni       | Limitata o nulla                          | Limitata e focalizzata sui compiti                               | Persistente e funzionale agli obiettivi                  |
| Accesso a fonti esterne       | Generalmente assente                      | Generalmente assente                                             | Nativo: web, database, file system, strumenti            |
| Interazione con altri sistemi | Nessuna                                   | Nessuna                                                          | Elevata: API, applicazioni, altri agenti                 |
| Punti di forza                | Linguaggio, fluidità, scalabilità         | Ragionamento analitico, debug, matematica                        | Coordinamento, esecuzione, continuità operativa          |
| Limiti principali             | Nessuna riflessione, nessuna azione       | Costosi, meno efficaci per compiti creativi                      | Complessità, rischi di governance e controllo            |
| Ruolo organizzativo tipico    | Assistente                                | Esperto analitico                                                | Partner operativo integrato nel flusso di lavoro         |
| Esempi d'uso ideali           | Redazione, sintesi, traduzione            | Analisi complesse, problem solving                               | Ricerca, operations, workflow end-to-end                 |

Tabella 1. I tre modelli IA a confronto

## 1.3 GLI AGENTI IA: AUTONOMIA, SOCIAL ABILITY, REATTIVITÀ, PROATTIVITÀ

Stabilito che l'IA agentic segna il passaggio dalla generazione all'azione, occorre chiedersi cosa rende davvero un agente un possibile "membro operativo" di un team di lavoro. Non basta che un sistema produca una buona risposta. Per integrarsi in un processo di lavoro reale, deve saper agire con un certo grado di autonomia, coordinarsi con altri attori, reagire ai cambiamenti del contesto e, in alcuni casi, anticiparli. È questo insieme di proprietà che sposta l'IA da interfaccia conversazionale a **componente attiva del sistema organizzativo**.

### 1.3.1 Autonomia

---

L'autonomia non coincide con l'assenza di controllo umano. È la capacità di portare avanti un compito end-to-end con input esterno limitato, gestendo decisioni intermedie, strumenti esterni e correggendo il percorso quando emergono ostacoli. Gli agenti si distinguono dai chatbot tradizionali: non si limitano a rispondere, ma pianificano e agiscono. La letteratura sui team uomo-IA descrive livelli più alti di autonomia come la capacità di condurre attività complete con minore intervento diretto. L'autonomia utile, però, è sempre delimitata: richiede un perimetro operativo chiaro, regole di escalation, tracciabilità delle azioni e una responsabilità umana che presidia l'intero processo e resta decisiva nei passaggi finali. Senza questi presidi, ciò che appare come efficienza tecnica diventa una perdita di trasparenza decisionale e di tracciabilità operativa.

#### Key Insight

Un agente crea valore perché sa muoversi da solo entro confini chiari. Più cresce l'autonomia, più contano supervisione e governance.

### 1.3.2 Social ability

---

La seconda proprietà decisiva è la social ability, che non va intesa in senso antropomorfo; non significa che l'agente debba "sembrare umano" ma che deve saper operare in una rete di interdipendenze: ricevere consegne, restituire output comprensibili, chiedere chiarimenti, segnalare eccezioni, coordinarsi con applicazioni, con altri agenti e con colleghi umani. È una capacità relazionale, non performativa. I protocolli descritti nei paragrafi precedenti sono rilevanti proprio in questa prospettiva: estendono l'accesso dell'agente a risorse esterne e la sua capacità di cooperare con altri attori artificiali, trasformandolo da sistema cognitivo isolato a parte di un sistema di azione collettiva.

La social ability richiede anche una buona progettazione della fiducia. Le evidenze sperimentali suggeriscono che la fiducia non cresce semplicemente rendendo l'agente più "umano" o nascondendone l'identità. In uno studio comportamentale del 2023 dal titolo "Trust in an AI versus a Human teammate" sulla cooperazione umano-IA, i partecipanti hanno mostrato maggiore fiducia comportamentale nell'AI teammate che in un teammate presentato come umano: accettavano meno spesso la sua decisione quando credevano di collaborare con un altro essere umano, mentre la performance congiunta dipendeva soprattutto dalla qualità effettiva del contributo dell'IA, non dalla sua identità percepita. Anche spiegare tutto, sempre, può essere controproducente: in altri studi livelli più alti di spiegazione hanno ridotto fiducia e percezione di competenza. Conta progettare le interazioni giuste, non moltiplicarle. Un agente efficace sul piano relazionale comunica ciò che serve, quando serve.

### 1.3.3 Reattività

---

La reattività riguarda la capacità di percepire eventi rilevanti e rispondere in tempo utile. Un LLM tradizionale si limita ad attendere un prompt e produrre un output. Un agente, invece, può leggere nuovi dati, monitorare lo stato di un workflow, confrontare il risultato con l'obiettivo e aggiornare la propria azione. Nei processi organizzativi, il problema raramente è solo disporre di una risposta, ma averla nel momento giusto, dopo che il contesto è cambiato. La reattività diventa quindi preziosa dove esistono basi documentali aggiornate, passaggi di consegna leggibili e segnali operativi chiari. La tecnologia, da sola, non basta. L'organizzazione deve decidere quali segnali l'agente debba interpretare, quali soglie facciano scattare un intervento e quando invece il sistema debba fermarsi e coinvolgere una persona. È in queste scelte che la reattività tecnica si traduce in affidabilità organizzativa.

### 1.3.4 Proattività

---

La proattività è il tratto che più distingue gli agenti IA dai sistemi generativi precedenti. Consiste nel proporre una sequenza di azioni, avviarla, verificarne l'esito e correggersi. Anton Korinek, economista e autore del paper "AI Agents for Economic Research", descrive questo funzionamento come una sequenza think-act-observe-respond: l'agente valuta quale informazione o strumento sia necessario, agisce, osserva il risultato e aggiorna il passo successivo. Nel lavoro reale, questo significa poter aprire sotto-task, raccogliere informazioni da fonti diverse, verificare i passaggi critici e riaprire il problema se il primo tentativo non funziona. È qui che l'agente inizia a somigliare a un membro operativo del team. È anche il tratto che richiede più rigore. La proattività genera valore solo se accompagnata da metriche chiare,

supervisione e possibilità di contestazione. Un agente che agisce in anticipo è utile a una condizione: che lo faccia in modo verificabile, reversibile e coerente con i criteri del team.

Considerate nel loro insieme, queste quattro proprietà chiariscono perché l'IA agentica rappresenti una discontinuità organizzativa prima ancora che tecnologica. L'agente non è un software più capace: è un attore tecnico che entra nei meccanismi di coordinamento, redistribuisce i carichi cognitivi e impone nuove regole di supervisione. La questione strategica riguarda la disponibilità dell'organizzazione ad accoglierli come membri operativi del processo, con compiti chiari, ruoli definiti e un presidio umano strutturato.

## 1.4 MATURITÀ DI MERCATO E DIFFUSIONE SETTORIALE: BENCHMARK GLOBALI E CURVE DI ADOZIONE

### 1.4.1 Maturità di mercato e benchmark globali

---

La maturità di mercato dell'IA agentica non si misura sul numero di annunci o di progetti pilota, ma sul grado in cui gli agenti vengono integrati nei processi, ricevono un perimetro operativo chiaro e vengono governati come componenti stabili del modello operativo. Il mercato sta uscendo dalla fase puramente esplorativa: secondo Microsoft<sup>5</sup>, oltre l'80% dei leader si aspetta che gli agenti siano integrati in modo moderato o esteso nella propria strategia IA entro i prossimi 12–18 mesi, mentre quattro organizzazioni su cinque dichiarano di trovarsi in fase di pilota o oltre. La distribuzione resta però molto diseguale. Una quota significativa è ancora in fase di test o sta rivedendo la strategia dopo le prime sperimentazioni, segno che la maturità profonda rimane minoritaria. Il termine "Frontier Firm" coniato da Microsoft va preso con la prudenza che si deve al lessico dei vendor, ma identifica una distinzione reale: le organizzazioni che hanno superato la logica dello strumento isolato e combinano diffusione dell'IA, uso attivo di agenti e capacità di tradurre tutto questo in ritorni misurabili. Le evidenze comparative internazionali confermano che la principale linea di differenziazione non passa più tra adottanti e non adottanti, ma tra chi dispone degli asset complementari necessari a scalare e chi no. L'indagine OECD-BCG-INSEAD<sup>6</sup> segnala barriere legate a competenze, qualità dei dati, incertezza sul ritorno economico e capacità manageriale di tradurre l'IA in problemi organizzativi concreti. L'IA agentica è ormai una traiettoria di mercato credibile, ma la sua maturità organizzativa dipende da condizioni sociotecniche che molte imprese stanno ancora costruendo.

### 1.4.2 Diffusione settoriale: dove gli agenti entrano per primi

---

La diffusione settoriale degli agenti non segue i confini tradizionali tra industrie "digitali" e "non digitali". Segue la struttura del lavoro. Gli agenti si diffondono prima dove i compiti sono ad alta intensità informativa, relativamente standardizzabili, ricchi di dati e misurabili nei risultati: customer service, IT, finance operations e, in misura crescente, HR operations. Le rilevazioni più recenti di Microsoft<sup>7</sup> mostrano che oggi IT e customer service guidano l'adozione, mentre finance, sales e marketing seguono a breve distanza; procurement, HR e supply chain partono da livelli inferiori ma registrano tra i maggiori incrementi attesi nei prossimi cicli di adozione. PwC<sup>8</sup> rileva che più della metà delle imprese

---

<sup>5</sup> Microsoft, 2025 Work Trend Index Annual Report: The Year the Frontier Firm Is Born, 2025.

<sup>6</sup> OECD / Boston Consulting Group / INSEAD (2025). The Adoption of Artificial Intelligence in Firms: New Evidence for Policymaking. OECD Publishing, Paris.

<sup>7</sup> Microsoft WorkLab, Agents are here—is your company prepared?, sezione "The future of AI is agentic", 2026.

<sup>8</sup> PwC, AI Agent Survey, maggio 2025, sezione "Only a few businesses have their eyes on the prize"

del campione dichiara un uso attivo o piani di adozione nei successivi sei mesi in customer service (57%), sales and marketing (54%) e IT and cybersecurity (53%); tra gli adottanti, il 66% riporta guadagni di produttività, il 57% risparmi di costo e il 55% maggiore rapidità decisionale.

Questa sequenza non è casuale. In queste funzioni gli agenti possono essere inseriti in flussi di lavoro con passaggi di consegna chiari, basi documentali relativamente ordinate e metriche leggibili: tempo di risposta, tasso di risoluzione, accuratezza documentale, produttività per addetto. Deloitte<sup>9</sup> segnala che il customer support è l'ambito di maggiore impatto atteso, mentre supply chain management, R&D, knowledge management e cybersecurity presentano anch'essi un potenziale particolarmente elevato. Dove invece prevalgono eccezioni, elevata variabilità dei compiti, bassa qualità dei dati o forte dipendenza dal contesto informale, la diffusione è più lenta e spesso resta confinata a casi mirati.

La dimensione d'impresa resta un fattore importante. Le grandi organizzazioni dispongono più spesso di dati, architetture e ruoli di governo adeguati; le PMI stanno aumentando l'adozione soprattutto tramite strumenti standard, mentre gli agenti più personalizzati o integrati nei processi core restano più sperimentali. Il survey OECD D4SME 2026<sup>10</sup>, basato su un campione di oltre 2.000 PMI in 12 paesi OECD, segnala una rapida crescita dell'adozione IA – soprattutto di applicazioni standard – ma anche ostacoli persistenti legati a costi, tempo di formazione e carenze di competenze.

### 1.4.3 Dalle curve di adozione alla “J-curve” della produttività

---

Dal punto di vista manageriale, le curve di adozione dell'IA agentica assomigliano meno a una progressione lineare e più a una **J-curve della produttività**<sup>11</sup>. Nelle fasi iniziali, i benefici attesi vengono spesso compensati da costi di adattamento: revisione dei processi, ridefinizione dei ruoli, supervisione aggiuntiva, formazione, nuove pratiche di controllo. Non sorprende quindi che una parte delle organizzazioni, dopo le prime sperimentazioni, stia rivedendo la propria strategia invece di scalarla.

Il valore emerge quando l'agente entra davvero nel processo, non quando viene aggiunto a margine. Molte organizzazioni sono ancora indietro: PwC<sup>12</sup> rileva che meno della metà sta ripensando in modo esplicito il proprio modello operativo o ridisegnando i processi attorno agli agenti, pur dichiarando benefici su produttività, costi e velocità decisionale. La frizione dei primi mesi non segnala necessariamente un fallimento tecnologico: segnala piuttosto che l'organizzazione non ha ancora trasformato l'adozione in progettazione operativa.

---

<sup>9</sup> Deloitte AI Institute, The State of AI in the Enterprise – 2026 AI Report, sezione “Agentic AI”

<sup>10</sup> OECD, Empowering SMEs in the Age of AI: The 2026 OECD D4SME Survey, OECD SME and Entrepreneurship Papers No. 78, 2026

<sup>11</sup> Brynjolfsson, E., Rock, D., & Syverson, C. (2018), The Productivity J-Curve: How Intangibles Complement General Purpose Technologies, NBER Working Paper No. 25148; Berndt et al. (2026), A Systematic Approach to Experimenting with Gen AI, Harvard Business Review

<sup>12</sup> PwC, AI Agent Survey, 2025, sezione “Only a few businesses have their eyes on the prize”

La traiettoria di adozione più robusta resta quella della sperimentazione disciplinata. I progetti che generano valore durevole sono quelli che partono da bisogni chiari, costruiscono prototipi realmente utilizzabili, misurano effetti primari e secondari e correggono rapidamente il disegno del lavoro. La maturità di mercato dell'IA agentica si gioca sulla capacità di apprendere prima di scalare, oltretutto sulla velocità di adozione.

### **Key Insight**

Nell'IA agentica, i ritorni non crescono in modo lineare: i primi mesi sono spesso una fase di frizione organizzativa, mentre il salto di valore arriva quando cambiano processi, ruoli e metriche di gestione.

## 2.1 LA PERFORMANCE DEI TEAM: FRAMEWORK TEORICI E VALUTAZIONI SOCIO-TECNICHE

Per valutare in modo rigoroso la performance dei *team aumentati* occorre prima stabilire con chiarezza da quale prospettiva la si osserva. Non è sufficiente misurare ciò che l'agente IA sa fare sul piano tecnico, perché i risultati dipendono anche da come tecnologia, organizzazione del lavoro e dinamiche umane si combinano nel contesto specifico del team. Scegliere una cornice analitica significa quindi adottare una lente capace di leggere insieme struttura, processi, relazioni e risultati.

### 2.1.1 La prospettiva sociotecnica: che cosa guarda e perché è necessaria

---

In questo lavoro adottiamo una prospettiva sociotecnica, spesso indicata con l'acronimo STS, da Socio-Technical Systems. Si tratta di una tradizione di ricerca consolidata, sviluppata a partire dai lavori del Tavistock Institute, in particolare di Eric Trist, Ken Bamforth e Fred Emery, e successivamente approfondita negli studi sull'organizzazione del lavoro<sup>13</sup>. La premessa di fondo è semplice: nei sistemi di lavoro, i risultati non dipendono mai solo dalla qualità della tecnologia, né solo dalla qualità delle relazioni umane, ma dal modo in cui componente tecnica e componente sociale si combinano tra loro. Strumenti, procedure, dati e architetture operative producono effetti diversi a seconda di come si intrecciano con ruoli, competenze, relazioni e culture professionali. Per questa ragione, l'approccio STS sostiene che tecnologia e organizzazione del lavoro devono essere progettate e valutate insieme.

Applicata ai team aumentati, questa impostazione aiuta a evitare un errore frequente: considerare l'agente IA come un semplice strumento da aggiungere a un team già dato. Quando un agente entra nel lavoro di squadra, non cambia soltanto la velocità con cui alcune attività vengono eseguite. Cambiano la distribuzione del lavoro, i punti di controllo, i margini di autonomia, i criteri di escalation e il modo in cui le persone costruiscono affidamento reciproco. Trattare l'IA come una semplice aggiunta tecnologica porta a sottostimare proprio questi effetti sistemici. Nei team aumentati, l'agente modifica il sistema di lavoro entro cui il compito prende forma.

L'integrazione tra STS, IMOI e Team Effectiveness costituisce la cornice teorica proposta in questo lavoro per l'analisi dei team aumentati.

Ne discende una conseguenza manageriale rilevante. La "performance" non può essere ridotta alla sola efficienza, né misurata soltanto in termini di velocità o di volume di output. Una valutazione adeguata deve includere almeno tre domande. La prima riguarda il risultato: il team produce output utili, affidabili e di qualità per gli stakeholder? La seconda riguarda il funzionamento: il modo in cui il

---

<sup>13</sup> Emery, F. E., & Trist, E. L. (1960). Sociotechnical Systems. In C. W. Churchman & M. Verhulst (Eds.), *Management Science, Models and Techniques* (Vol. 2, pp. 83–97). Oxford: Pergamon.

lavoro è distribuito tra umani e agenti rafforza oppure indebolisce coordinamento, chiarezza decisionale e capacità di gestione delle eccezioni? La terza riguarda la sostenibilità: l'esperienza di lavoro con l'IA accresce competenze, motivazione e qualità del giudizio, oppure erode autonomia, apprendimento e responsabilità? Sono domande che l'approccio sociotecnico rende possibili, perché obbliga a osservare insieme performance operativa, qualità del coordinamento e tenuta del sistema di lavoro nel tempo.

### 2.1.2 IMOI e Hackman/TDS per la performance dei team aumentati

---

Una volta chiarita la lente sociotecnica, serve una griglia teorica che permetta di analizzare in modo ordinato ciò che accade all'interno di un team. Un primo riferimento autorevole è il modello IMOI, proposto da Ilgen e colleghi<sup>14</sup> come evoluzione del più tradizionale schema IPO (Input-Process-Output). Il punto chiave è che la vita di un team non è lineare. Le condizioni iniziali contano, ma non bastano; tra input e risultati intervengono processi di coordinamento e stati emergenti e gli esiti prodotti in un ciclo di lavoro modificano le condizioni del ciclo successivo. Da qui il nome IMOI: Input-Mediator-Output-Input. Il termine “mediatori” è importante, perché comprende sia i processi osservabili del team, come pianificazione, coordinamento e monitoraggio, sia quelli meno visibili ma altrettanto decisivi, come fiducia, sicurezza psicologica e identità di team. È una cornice particolarmente adatta ai team aumentati, perché consente di integrare in un unico schema analitico elementi tecnici, processi collaborativi e dinamiche umane.

Nel caso dei team aumentati, gli input non riguardano soltanto la composizione del gruppo o le competenze disponibili. Includono anche il livello di autonomia assegnato all'agente, la qualità dei dati su cui opera, la chiarezza dei ruoli, le regole di supervisione, i punti di passaggio tra umano e sistema. I mediatori, oltre al coordinamento quotidiano, comprendono anche il modo in cui si costruiscono affidamento, comprensione reciproca e capacità di correggere errori o ambiguità. Gli output, infine, non coincidono solo con la prestazione tecnica: comprendono qualità del risultato, apprendimento, sostenibilità del carico e disponibilità a continuare a lavorare insieme. Questa impostazione è utile perché mostra che i risultati di un team umano-IA dipendono dal modo in cui il team è stato progettato e governato piuttosto che dal modello IA utilizzato.

Accanto all'IMOI, un secondo riferimento centrale è il lavoro di J. Richard Hackman sulla team effectiveness<sup>15</sup>. Il suo contributo resta centrale perché sposta l'attenzione da una domanda reattiva — “come ha performato il team?” — a una domanda di progettazione: “il team è stato messo nelle condizioni per funzionare bene?”. Nel suo framework, l'efficacia del team è definita dalla qualità dell'output, dalla capacità del gruppo di continuare a collaborare nel tempo e anche dagli effetti dell'esperienza di team sui suoi membri, in termini di apprendimento, crescita e tenuta. Un team che

---

<sup>14</sup> Ilgen, D. R., Hollenbeck, J. R., Johnson, M., & Jundt, D. (2005) *Teams in Organizations: From Input-Process-Output Models to IMOI Models*. Annual Review of Psychology, 56, 517–543.

<sup>15</sup> Hackman, J. R. (2002). *Leading Teams: Setting the Stage for Great Performances*. Boston, MA: Harvard Business School Press.

produce più velocemente ma erode giudizio, motivazione o capacità di cooperazione sta diventando più rapido e nella pratica organizzativa la velocità viene scambiata per capacità con una frequenza che dovrebbe preoccupare.

Letti insieme, STS, IMOI e Hackman offrono una cornice solida per questo lavoro. La prospettiva sociotecnica definisce il principio generale: dimensione tecnica e dimensione sociale vanno ottimizzate insieme. L'IMOI indica dove osservare: negli input, nei mediatori, negli output e nei cicli di retroazione. Hackman traduce tutto questo in una domanda di progettazione: **il team è stato messo nelle condizioni di funzionare bene prima ancora di chiedergli di performare?** La questione è quale assetto di ruoli, flussi di lavoro, supporti, metriche e pratiche di coordinamento consenta al team di usare l'IA senza perdere qualità del giudizio, responsabilità e capacità di apprendimento.

### **Key Insight**

Nei team aumentati, la tecnologia non determina da sola la performance. I risultati emergono dall'interazione tra architettura tecnica, organizzazione del lavoro e dinamiche sociali che permettano a umani e agenti di lavorare insieme con qualità, responsabilità e continuità.

## 2.2 RUOLI IBRIDI E TASK-ALLOCATION: MODELLI DI ORCHESTRAZIONE UMANO-IA

### 2.2.1 Allocazione per natura del compito: oltre la seniority

---

Nei team aumentati, la domanda decisiva non è se un compito debba essere affidato a profili con livelli diversi di esperienza professionale o a un agente, ma quale combinazione di attori sia più adatta alla sua struttura. La task-allocation efficace dipende principalmente da quattro proprietà del compito: **variabilità, criticità, verificabilità e bisogno di contesto**. I task più candidabili alla delega operativa agli agenti sono quelli relativamente stabili, ripetibili, ben verificabili e supportati da regole o dati espliciti; i task ad alta ambiguità, forte impatto decisionale, eccezioni frequenti o rilevante dipendenza dal contesto restano invece più adatti a una conduzione umana, eventualmente assistita dall'IA. Questo principio è coerente sia con la tradizione sulla function allocation nei sistemi uomo-automazione sia con la letteratura più recente sul teaming umano-IA.

Le evidenze empiriche rendono il criterio più concreto. Lo studio di Wang e colleghi<sup>16</sup> mostra che gli agenti, rispetto agli umani, tendono a performare meglio su flussi di lavoro facilmente programmabili, ma incontrano difficoltà maggiori nei compiti aperti, visivamente dipendenti o ricchi di sfumature contestuali, dove possono anche mascherare limiti di qualità con uso improprio degli strumenti o fabbricazione di dati. In parallelo, gli studi sperimentali su collaborazione umano-IA mostrano che i vantaggi tendono a concentrarsi maggiormente tra i profili con minore esperienza iniziale, purché l'affidamento all'IA sia ben calibrato. Caplin e colleghi<sup>17</sup>, in particolare, mostrano che i maggiori benefici si osservano tra chi ha minore abilità di partenza ma una buona consapevolezza dei propri limiti.

Ne deriva una conseguenza manageriale precisa: l'allocazione non dovrebbe essere disegnata per status gerarchico, ma per combinazione tra difficoltà del compito e capacità di controllo. Nei task più codificabili, l'agente può diventare un esecutore rapido o un primo livello di analisi; **nei task più critici, il giudizio deve restare in capo all'umano**, anche quando l'IA contribuisce in modo intenso alla preparazione o all'esplorazione delle alternative.

#### Key Insight

La task-allocation efficace parte dalla natura del compito: più un'attività è stabile e verificabile, più può essere delegata; più è critica, ambigua e contestuale, più richiede ownership umana.

---

<sup>16</sup> Wang et al., How Do AI Agents Do Human Work? Comparing AI and Human Workflows Across Diverse Occupations (2025)

<sup>17</sup> Caplin, A., Deming, D. J., Li, S., Martin, D. J., Marx, P., Weidmann, B., & Ye, K. J. (2024). The ABC's of Who Benefits from Working with AI: Ability, Beliefs, and Calibration. NBER Working Paper No. 33021.

## 2.2.2 Pattern di orchestrazione: triage, specialisti, escalation

---

L'orchestrazione umano-IA si presenta in una famiglia di pattern ricorrenti. Il primo, ormai molto diffuso, è **triage** → **esecuzione** → **verifica**: l'agente filtra o prepara il lavoro, un attore – umano o artificiale – esegue la parte standardizzabile, e un umano valida i passaggi critici o l'output finale. È un modello adatto a processi documentali, assistenza interna, analisi preliminari e molte attività operative. Il secondo pattern è quello degli **agenti specialisti coordinati**: più agenti con competenze diverse operano sullo stesso flusso di lavoro sotto una funzione di regia che assegna compiti, gestisce dipendenze e raccoglie gli output. Il terzo è il modello a **escalation verso un responsabile umano**, in cui l'agente procede finché confidenza, rischio e contesto restano entro soglia, ma trasferisce il controllo quando emergono eccezioni, bassa verificabilità o implicazioni regolatorie. Vale qui un principio noto da tempo nella ricerca sulla collaborazione uomo-macchina: il valore nasce dalla capacità di spostare il controllo, momento per momento, su chi può gestirlo meglio, più che dal massimizzare l'autonomia del sistema.

La rilevanza di questi pattern cresce con l'aumentare dell'autonomia. Quando l'agente non si limita a suggerire, ma agisce all'interno di un flusso operativo, i passaggi di consegne diventano punti critici di qualità e sicurezza. La ricerca recente sull'"agent manager" va esattamente in questa direzione<sup>18</sup>: un livello di coordinamento che scompone obiettivi, assegna attività, traccia esecuzioni e mantiene il processo leggibile. Quando gli agenti iniziano ad agire dentro i flussi di lavoro, regole e passaggi di controllo devono essere espliciti.

Le indicazioni operative più aggiornate convergono sullo stesso punto. Microsoft<sup>19</sup> insiste sulla necessità di governance, osservabilità e controllo centralizzato degli agenti; KPMG<sup>20</sup> osserva che i flussi di lavoro agentici introducono rischi specifici – errori a cascata, maggiore complessità di governance, tensioni sui principi di segregazione dei compiti – proprio perché moltiplicano i passaggi automatizzati e le dipendenze tra attori. In questo quadro, **l'assegnazione dei compiti rappresenta un sistema continuo di passaggi di consegne, controllo e adattamento**.

## 2.2.3 Nuovi ruoli organizzativi: ownership, supervisione, compliance

---

Quando gli agenti entrano nei flussi di lavoro reali, emergono ruoli che nelle organizzazioni tradizionali erano meno distinti o distribuiti informalmente. Un primo ruolo è quello dell'**agent supervisor**, che oltre a validare l'output finale monitora comportamento, limiti, qualità dei passaggi di consegna e situazioni di dipendenza eccessiva. Un secondo è l'**AI process owner**, responsabile del disegno del flusso di lavoro ibrido, delle soglie di automazione e delle metriche di funzionamento. Un terzo è l'**AI**

---

<sup>18</sup> Masters, C. et al. "Orchestrating Human-AI Teams: The Manager Agent as a Unifying Research Challenge." arXiv(2025)

<sup>19</sup> Microsoft, "Governance and security for AI agents across the organization" (2026)

<sup>20</sup> KPMG, Agentic AI workflows in financial reporting (2026)

**risk/compliance analyst**, che presidia dati, verificabilità, segregazione dei compiti, policy e conformità. A questi si affianca spesso una figura più vicina alla delivery – il **builder** o specialista di agenti – che configura strumenti, prompt, integrazioni e regole operative.

I segnali più recenti vanno in questa direzione. Il Work Trend Index 2025 rileva che il 78% dei leader sta considerando ruoli specifici per l'IA; fra i profili più citati compaiono specialisti di agenti, formatori IA, specialisti sicurezza, analisti di ROI e figure di processo. Nello stesso rapporto, Microsoft osserva che una quota crescente di organizzazioni prevede ruoli di “**AI workforce manager**” e si aspetta che, nel medio termine, una parte crescente dei manager debba non solo gestire persone ma anche formare, coordinare e valutare agenti. Non tutte queste etichette si tradurranno subito in job family formali. Infatti, senza una responsabilità esplicitamente assegnata, gli agenti entrano nei processi senza presidio sufficiente.

Per i decisori, il tema è rendere visibili responsabilità che altrimenti resterebbero implicite, senza per questo creare una nuova burocrazia dell'IA. Dove il lavoro ibrido cresce, la domanda organizzativa non è più soltanto “quali strumenti adottare?”, ma “chi governa il comportamento dell'agente, chi risponde del processo e chi presidia il rischio?”

### **Key Insight**

Quando l'agente diventa un nodo stabile del workflow, le responsabilità non possono restare diffuse o tacite: servono ownership, supervisione e presidio del rischio chiaramente assegnati.

## **2.3 COMUNICAZIONE, COORDINAMENTO, MOTIVAZIONE E FIDUCIA**

### **NEI TEAM MISTI**

L'emergere dei team misti, composti da esseri umani e agenti di intelligenza artificiale, potrebbe introdurre una discontinuità importante nel modo in cui il lavoro collettivo viene pensato, organizzato e vissuto. Comunicazione, coordinamento, motivazione e fiducia – elementi da sempre centrali nella performance dei team – non scompaiono, ma potrebbero cambiare natura.

#### **2.3.1 Comunicazione**

---

Nei team tradizionali, la comunicazione svolge una funzione che va ben oltre lo scambio di informazioni. Serve a costruire significati condivisi, a negoziare priorità, a rendere visibili dubbi e intuizioni. Nei team misti, l'introduzione dell'IA tende a comprimere questi spazi comunicativi quando l'IA viene interpretata solo come automazione di alcuni compiti. Quando una parte del lavoro viene svolta in modo automatico, rapido e apparentemente "oggettivo", il dialogo umano rischia di ridursi al minimo indispensabile.

Questa riduzione può apparire, almeno inizialmente, come un vantaggio: meno riunioni, meno spiegazioni, meno frizioni. Nel tempo, però, il rischio è che il team perda la capacità di riflettere collettivamente sul proprio operato. Gli agenti IA dovrebbero essere introdotti come parte della comunicazione, spostandola da un piano operativo a un piano più interpretativo e decisionale, dove il contributo umano rimane insostituibile.

Una caratteristica fondamentale dei team ad alta performance è l'elevata sicurezza psicologica, ovvero la convinzione che un individuo si possa esprimere liberamente senza timore di essere giudicato. Monitorare questa dimensione durante l'inserimento di agenti IA nei team è quindi essenziale.

#### **2.3.2 Coordinamento**

---

Il coordinamento nei team misti diventa una questione critica perché l'IA opera secondo logiche diverse da quelle umane. Gli esseri umani si coordinano attraverso routine, aspettative implicite e aggiustamenti continui; gli agenti artificiali seguono istruzioni, obiettivi e ottimizzazioni che possono risultare opache o difficili da anticipare. Inoltre, i team di lavoro consolidati operano spesso secondo codici linguistici e routine che gli agenti di IA potrebbero non assimilare rapidamente (si rimanda al capitolo 3.1).

Questo può creare una tensione strutturale. Da un lato, l'IA accelera l'esecuzione dei compiti. Dall'altro, può alterare gli equilibri di sincronizzazione del team, generando disallineamenti. Il problema, dunque, può essere l'assenza di un linguaggio di coordinamento condiviso.

L'opportunità per i team aumentati sta nel ripensare il coordinamento come una responsabilità distribuita. È importante prevedere per gli agenti IA un inserimento simile a quello che si adotta per nuovi membri del team così da evitare disallineamenti che possono portare a inefficienze.

### 2.3.3 Motivazione

---

La motivazione nei team misti è uno degli aspetti più delicati. Quando l'IA assume una parte significativa del carico di lavoro, i membri umani possono sperimentare una riduzione del senso di utilità personale. Se la macchina è più veloce, più precisa o più instancabile, il contributo umano rischia di apparire marginale.

Questo può tradursi in un disimpegno silenzioso: meno iniziativa, minore sforzo cognitivo, ridotta responsabilità percepita. Raramente si tratta di resistenza al cambiamento; più spesso è una ridefinizione implicita dei confini del proprio ruolo e più difficile da intercettare nel tempo. Al tempo stesso, l'IA apre anche nuove opportunità motivazionali: libera tempo ed energia per attività a maggiore valore, come la definizione dei problemi, la valutazione delle alternative e la presa di decisioni complesse. Uno studio del MIT<sup>21</sup> ha dimostrato come su un'analisi di 19.000 compiti in 950 categorie occupazionali emergano **cinque dimensioni su cui l'umano eccelle e l'IA presenta delle limitazioni**: Empatia, Presenza, Opinione, Creatività e Speranza.

La chiave sta dunque nel rendere esplicito il valore del contributo umano nei team aumentati, evitando che l'IA venga percepita come elemento sostitutivo delle capacità distintivamente umane.

#### Key Insight

Quando l'IA assorbe una parte crescente delle attività operative, la motivazione umana dipende dal "perché": se il contributo umano non viene reso esplicito e valorizzato con attività di focus nelle aree in cui l'umano eccelle (re-skilling), il rischio è un disimpegno silenzioso.

### 2.3.4 Fiducia: tra affidamento e relazione

---

La fiducia nei team misti assume una forma ambigua. Gli esseri umani possono imparare rapidamente ad affidarsi all'IA per ottenere risultati, ma questo affidamento non coincide con una fiducia relazionale.

---

<sup>21</sup> Loaiza, I. and Rigobon, R., *The EPOCH of AI: Human-Machine Complementarities at Work* (November 21, 2024). MIT Sloan Research Paper No. 7236-24, Available at SSRN: <https://ssrn.com/abstract=5028371>

L'IA non condivide intenzioni, non prova emozioni, non negozia obiettivi. Di conseguenza, **il rapporto che si instaura è funzionale, non sociale**. Inoltre, quando l'IA non presenta soluzioni soddisfacenti questa fiducia può dissolversi più rapidamente che nelle relazioni tra esseri umani. Un ulteriore problema emerge quando questo affidamento tecnico indebolisce la fiducia reciproca del team. Meno interazioni dirette, meno collaborazione visibile e meno occasioni di supporto reciproco possono erodere il capitale sociale del gruppo. L'opportunità, invece, è usare l'IA come elemento di stabilità operativa, preservando e rafforzando al contempo la fiducia tra i membri attraverso momenti di confronto, decisione condivisa e responsabilità collettiva.

### **Key Insight**

Nei team misti, l'affidamento tecnico all'IA non equivale a fiducia: senza presidiare le relazioni umane, l'automazione può stabilizzare le operazioni ma indebolire il capitale sociale del team.

## **2.3.5 Una sfida per la progettazione socio-tecnica**

---

Comunicazione, coordinamento, motivazione e fiducia nei team misti non migliorano da soli quando entra un agente. Possono persino indebolirsi, se l'automazione viene trattata come una semplice aggiunta tecnologica. La vera opportunità dei team aumentati risiede nella progettazione intenzionale del sistema sociotecnico: riconoscere che l'IA cambia le regole del gioco e che le dinamiche di team devono essere ripensate di conseguenza.

## 2.4 METRICHE INTEGRATE DI PERFORMANCE: OUTPUT QUANTITATIVI E DINAMICHE QUALITATIVE

Per misurare gli impatti dell'introduzione di agenti IA in team esistenti da una prospettiva socio-tecnica è stato progettato, in occasione della sperimentazione in CDP (si rimanda al capitolo 3.3), uno strumento di assessment sviluppato insieme a Belbin UK (<https://www.belbin.com>). Il sistema di assessment è stato progettato per misurare in modo longitudinale (inizio, metà ciclo, fine) gli effetti dell'introduzione di agenti IA in team esistenti, combinando indicatori strutturali, processuali, percettivi e di risultato. L'impianto riflette il framework **IMOI (Input–Mediator–Output–Input)**<sup>22</sup> e integra misure quantitative standardizzate e contributi narrativi emergenti.

### 2.4.1 Inputs (condizioni di partenza del team)

---

Le misure di **input** descrivono le condizioni strutturali e contestuali entro cui opera il team, prima e durante l'introduzione dell'agente IA. In particolare, vengono osservate:

- **Struttura del team:** distribuzione dei ruoli secondo il modello **Belbin**, per valutare equilibrio, complementarità e potenziali sovra- o sotto-rappresentazioni funzionali.
- **Composizione e competenze:** percezione dell'adeguatezza del mix di competenze tecniche, professionali e cognitive rispetto agli obiettivi del team.
- **Risorse tecniche e informative:** disponibilità e qualità dei sistemi informativi, delle infrastrutture digitali e delle condizioni per condividere, elaborare e utilizzare informazioni.
- **Supporto contestuale e organizzativo:** livello di supporto percepito da leadership e organizzazione: coaching, sponsorship, legittimazione del team e del nuovo assetto sociotecnico.

Queste misure permettono di qualificare il **potenziale di funzionamento** del team e di interpretare le dinamiche successive.

### 2.4.2 Mediators (processi di team e stati emergenti)

---

I **mediatori** rappresentano il cuore dell'analisi IMOI: ciò che accade nel funzionamento quotidiano del team e che traduce gli input in risultati.

---

<sup>22</sup> Ilgen, D. R., Hollenbeck, J. R., Johnson, M., & Jundt, D. (2005) *Teams in Organizations: From Input-Process-Output Models to IMOI Models*. Annual Review of Psychology, 56, 517–543.

## Processi di team

Per misurare i processi di team utilizziamo la Team Process Scale, un questionario sviluppato da Mathieu et al. (2020), ampiamente utilizzato nella letteratura sulla team effectiveness. Il modulo copre le principali famiglie di processi: pianificazione e definizione degli obiettivi, allineamento strategico, monitoraggio delle prestazioni e dell'ambiente, coordinamento e integrazione del lavoro, gestione dei conflitti e supporto reciproco, regolazione affettiva e motivazionale. Questo strumento consente di osservare **come il team organizza il lavoro**, prende decisioni e si adatta all'incertezza, anche in presenza di un agente IA.

## Stati emergenti

Per quanto riguarda gli stati emergenti, ci focalizziamo su **sicurezza psicologica** e **identità di team** per i quali utilizziamo due sotto-questionari dello Psychological Safety Inventory (Plouffe et al., 2023). Le misure coprono aspetti classici della sicurezza psicologica come fiducia reciproca e inclusione, senso di appartenenza, sicurezza del ruolo e stabilità percepita. Queste dimensioni sono centrali per comprendere come l'IA influenzi dinamiche di fiducia, collaborazione e legittimazione reciproca.

### 2.4.3 Outputs (Performance, outcome umani e sostenibilità del team)

---

#### Outcome quantitativi

Gli outcome quantitativi variano in funzione della natura del compito e del contesto organizzativo. Le dimensioni tipicamente misurabili includono accuratezza dell'output, tempo di completamento, frequenza degli errori e tasso di escalation verso il supervisore umano. Data la natura esplorativa dello studio e la ridotta dimensione del campione, la definizione puntuale di queste metriche è stata subordinata alle caratteristiche specifiche del processo osservato.

#### Outcome motivazionali per i membri del team

Belbin Engagement Survey Tool (Belbin UK): utilizzato per misurare engagement, motivazione, coinvolgimento, qualità della relazione con il supervisore, autonomia decisionale e clima di rispetto e inclusione. Includiamo inoltre due moduli aggiuntivi su soddisfazione lavorativa e apprendimento e sviluppo delle competenze.

#### Sostenibilità del team

Scala breve di sostenibilità derivata dalla letteratura su team viability (la capacità di un team di restare efficace e di continuare a collaborare nel tempo; Mathieu, Hackman), focalizzata su desiderio di continuare a lavorare insieme, sostenibilità del carico di lavoro, solidità delle relazioni future.

#### 2.4.4 Nuovi input per il ciclo successivo

---

In linea con il modello **IMOI**, gli output diventano nuovi **input** per il ciclo successivo. Questa sezione integra misure ripetute dei moduli precedenti e domande narrative sull'esperienza di lavoro in team con agenti IA. L'obiettivo è osservare come l'esperienza ridefinisce nuovi ruoli e aspettative, modifica percezioni di valore, contributo e identità, orienta la disponibilità del team a sostenere nel tempo configurazioni di **augmented teaming**. Il sistema di assessment funziona come **strumento di apprendimento organizzativo** e non solo strumento di misurazione.

### 3.1 SINTESI DELLE EVIDENZE EMPIRICHE: META-ANALISI DI STUDI ACCADEMICI SU TEAM UMANO-IA

Negli ultimi anni, la ricerca sui team misti umano-IA ha prodotto un corpus crescente di evidenze sperimentali e teoriche che restituisce un quadro complesso, segnato da effetti a volte eterogenei. L'analisi comparata dei principali contributi — che comprendono esperimenti di campo su larga scala, studi controllati in laboratorio, indagini percettive e lavori teorici di sintesi — consente di delineare una trasformazione del *teaming* (il modo di lavorare in squadra) che non può essere interpretata né come un puro incremento prestazionale né come un semplice effetto di automazione. I risultati convergono nell'indicare che **l'IA diventa realmente trasformativa solo quando viene integrata come partner cognitivo**, capace di co-progettare l'analisi, ampliare il perimetro delle competenze rilevanti e modificare i modelli di interazione all'interno delle organizzazioni.

Un primo insieme di evidenze proviene dal più ampio esperimento sul campo condotto finora<sup>23</sup>, svolto presso Procter & Gamble e basato su 776 lavoratori della conoscenza distribuiti tra funzioni R&D e Commercial. In questo contesto, l'uso dell'IA genera un incremento sostanziale della qualità delle soluzioni prodotte da individui e team: il lavoro individuale con IA raggiunge livelli comparabili a quelli di team da due individui senza IA, e i team di due persone che utilizzano l'IA ottengono miglioramenti ulteriori sia in termini di qualità sia di efficienza temporale, riducendo significativamente il tempo necessario per svolgere le attività assegnate. Questi esiti empirici sono tra i primi a mostrare in contesto reale che un sistema generativo può assumere una funzione di 'teammate cibernetico', in grado di ricreare, e in parte potenziare, i benefici del lavoro di squadra tradizionale.

Un secondo gruppo di risultati, ottenuti in uno studio sperimentale su 758 consulenti di Boston Consulting Group<sup>24</sup>, conferma e amplia questo quadro. In questo contesto, l'IA incrementa la qualità del lavoro in maniera particolarmente marcata nei compiti che rientrano entro le capacità del modello, con un miglioramento superiore al 40% rispetto alla base di riferimento, e genera guadagni di produttività nell'ordine del 25%. Il beneficio risulta tuttavia distribuito in modo asimmetrico: i lavoratori con prestazioni iniziali più basse sperimentano i maggiori incrementi, mentre i professionisti di punta traggono vantaggi più limitati, benché comunque positivi. Parallelamente, il medesimo studio mette in luce l'esistenza di un "jagged technological frontier", ovvero una soglia irregolare che separa le attività per le quali il modello possiede capacità affidabili dalle attività per le quali l'IA, nonostante l'apparente competenza, produce errori sistematici. In un compito che simula una consulenza strategica a un ipotetico CEO (un compito al di fuori di questa frontiera), l'uso dell'IA riduce significativamente la

---

<sup>23</sup> Dell'Acqua, F., Ayoubi, C., Lifshitz, H., Sadun, R., Mollick, E., Mollick, L., Han, Y., Goldman, J., Nair, H., Taub, S., and Lakhani, K. (2025). *The Cybernetic Teammate: A Field Experiment on Generative AI Reshaping Teamwork and Expertise*. NBER Working Paper 33641 (2025), <https://doi.org/10.3386/w33641>.

<sup>24</sup> Dell'Acqua, F., McFowland III, E., Mollick, E. R., Lifshitz-Assaf, H., Kellogg, K., Rajendran, S., Kraye, L., and Candelon, F., and Lakhani, K. R. (2026). *Navigating the Jagged Technological Frontier: Field Experimental Evidence of the Effects of Artificial Intelligence on Knowledge Worker Productivity and Quality*. In corso di pubblicazione presso Organization Science.

prestazione, soprattutto tra coloro che tendono ad affidarsi in modo eccessivamente fiducioso alle risposte del modello.

In parallelo ai risultati positivi ottenuti negli scenari di lavoro della conoscenza, dagli studi di laboratorio emergono dinamiche più problematiche. Il lavoro intitolato “Super Mario Meets AI”<sup>25</sup>, rappresenta il caso più emblematico. In un compito collaborativo ad alta interdipendenza come quello del gioco di SuperMario Party, l'introduzione di un agente artificiale peggiora la prestazione del gruppo, aumentando il numero di collisioni, disallineando le routine tacite e riducendo l'impegno degli altri membri umani. L'effetto non è esclusivo dell'IA ma si osserva anche con nuovi membri umani e rimanda ad un problema classico, la rigidità delle convenzioni e delle routine consolidate. Questo sottolinea la **necessità di processi dedicati di inserimento e accompagnamento** per l'ingresso di nuovi membri nel team, siano essi umani o agenti di intelligenza artificiale.

Risultati affini provengono anche dallo studio condotto da Qin e Sajda<sup>26</sup>, nel quale team umani collaborano con un agente IA in un ambiente immersivo di realtà virtuale. Qui, persino quando l'IA è controllata da un esperto umano, la sola presenza del sistema IA nel team riduce la prestazione complessiva e altera i modelli di comunicazione e di azione, inducendo comportamenti compensativi da parte dei partecipanti. L'alterazione dei segnali fisiologici e comportamentali osservata negli esperimenti suggerisce che l'IA, pur essendo tecnicamente capace, modifica le dinamiche di influenza e di responsabilità percepita all'interno del gruppo, con **effetti che non possono essere spiegati dalla sola dimensione computazionale** del contributo artificiale.

La letteratura legata alle percezioni dell'impatto dell'IA sulle organizzazioni contribuisce a illuminare un ulteriore livello di complessità. Una ricerca pubblicata su HBR<sup>27</sup>, condotta su oltre 1.400 lavoratori statunitensi, mostra come i leader sovrastimino, in modo significativo, l'entusiasmo dei propri collaboratori verso l'adozione dell'IA, generando un divario di percezioni che rischia di compromettere il clima di sicurezza psicologica necessario per integrare efficacemente sistemi artificiali nel lavoro quotidiano. Mentre l'80% dei dirigenti ritiene che i dipendenti siano ben informati e coinvolti nelle decisioni relative all'IA, solo il 27% dei dipendenti conferma questa percezione. Tale disallineamento indica che l'adozione dell'IA non richiede soltanto investimenti tecnologici, ma **un ripensamento delle pratiche di ascolto, comunicazione e co-progettazione organizzativa**.

Sul versante teorico, una serie di contributi recenti aiuta a interpretare in maniera unitaria l'eterogeneità degli esiti empirici. Lou e coautori<sup>28</sup> propongono una cornice integrativa per il lavoro di squadra umano-IA che identifica quattro dimensioni strutturali — formulazione, coordinamento, mantenimento e formazione — come condizioni necessarie per un funzionamento collaborativo efficace. Secondo questo approccio, molte delle difficoltà osservate negli studi sperimentali riflettono lacune nella

---

<sup>25</sup> Dell'Acqua, F., Kogut, B., & Perkowski, P. (2023). *Super Mario Meets AI: Experimental Effects of Automation and Skills on Team Performance and Coordination*. *Review of Economics and Statistics*, 107 (4): 951–966.

<sup>26</sup> Qin, Y., & Sajda, P. (2025). *Teaming with AI: A Multi-Modal Investigation of Human-AI Team Dynamics*. AAAI Conference Paper.

<sup>27</sup> Lovich, D., Meier, S., & Taylor, C. (2025). *Leaders Assume Employees Are Excited About AI — They're Wrong*. *Harvard Business Review*. November 2025

<sup>28</sup> Lou, B., Lu, T., Raghu, T. S., & Zhang, Y. (2025). *Unraveling Human-AI Teaming: A Review and Outlook*. <https://arxiv.org/abs/2504.05755>

costruzione di modelli mentali condivisi e nella gestione del ciclo di vita della collaborazione, più che limiti tecnici del modello stesso. Una serie di contributi complementari, come la rassegna di Berretta, Tausch e colleghi<sup>29</sup>, evidenzia come la letteratura rimanga ancora dominata da prospettive ingegneristiche e manchi di una definizione univoca di “team umano-IA”, sottolineando la **necessità di un’articolazione socio-tecnica più solida**. A ciò si affianca la prospettiva di Hagemann, Rieth e Suresh, che definisce i requisiti funzionali di una “IA centrata sul team” — capacità di responsività, consapevolezza situazionale e adattamento flessibile — come condizioni indispensabili perché un sistema artificiale possa essere percepito come membro legittimo del team.

Infine, il più recente contributo dei premi Nobel per l’economia Daron Acemoglu e Simon Johnson insieme a David Autor<sup>30</sup> fornisce un quadro concettuale particolarmente utile per interpretare le evidenze sperimentali. Gli autori distinguono tra diverse categorie di cambiamento tecnologico, mostrando che l’IA può assumere forme profondamente differenti: automatizzante, capital-augmenting, expertise-leveling o new-task-creating. Solo quest’ultima categoria è intrinsecamente pro-lavoratore, poiché genera domanda di nuovo giudizio umano e valorizza competenze non codificate. In questa prospettiva, gli esiti positivi osservati negli studi P&G e BCG riflettono un uso dell’IA come tecnologia che accresce le competenze, mentre gli esiti negativi dello studio su SuperMario e dei paradigmi immersivi mostrano ciò che avviene quando l’IA viene introdotta come sostituto, provocando **perdita di competenze, di autonomia e deterioramento del coordinamento**. Gli autori notano inoltre che l’attuale ecosistema di incentivi favorisce l’adozione di forme di IA orientate all’automazione piuttosto che alla collaborazione, creando un disallineamento strutturale tra il potenziale dell’IA e i suoi usi effettivamente implementati nelle organizzazioni.

---

<sup>29</sup> Berretta, S., Tausch, A. T., Ontrup, G., Gilles, B., & Peifer, C. (2023). *Defining human-AI teaming the human-centered way: a scoping review and network analysis*. *Frontiers in Artificial Intelligence*, 6, 1250725.

<sup>30</sup> Acemoglu, D., Autor, D., & Johnson, S. (2026). *Building Pro-Worker Artificial Intelligence*. NBER Working Paper 34854 (2026), <https://doi.org/10.3386/w34854>.

## 3.2 ESPERIMENTI AZIENDALI DOCUMENTATI: CASI FORTUNE 500, PMI E SETTORE PUBBLICO

Gli esperimenti condotti da grandi imprese e organizzazioni globali negli ultimi due anni rappresentano una fonte empirica cruciale per comprendere come l'introduzione dell'intelligenza artificiale nei flussi di lavoro reali stia trasformando in modo tangibile le dinamiche dei team, le capacità organizzative e la distribuzione delle competenze. Tra i casi più significativi, quelli di Procter & Gamble, Boston Consulting Group e Siemens offrono tre prospettive complementari: l'IA come partner cognitivo nelle attività di innovazione e lavoro della conoscenza; come leva di produttività e riorganizzazione dei processi consulenziali; e come supporto operativo per l'ingegneria e la manutenzione nel contesto manifatturiero.

Nel caso di Procter & Gamble<sup>31</sup>, uno dei più grandi e strutturati esperimenti sul campo finora realizzati, la collaborazione con un sistema generativo è stata valutata attraverso un design preregistrato che ha coinvolto 776 lavoratori della conoscenza distribuiti nelle funzioni R&D e Commercial. I partecipanti sono stati assegnati, individualmente o in team, a compiti reali di sviluppo di nuovi prodotti, secondo la metodologia e le fasi normalmente utilizzate dall'azienda nei processi di innovazione. I risultati indicano che l'IA assume effettivamente il ruolo di un "partner cognitivo": gli individui che lavorano con l'IA raggiungono livelli di qualità paragonabili a quelli dei team composti da due persone senza IA, mentre i team abilitati migliorano ulteriormente la qualità e riducono in modo sensibile il tempo necessario per completare le attività. L'effetto non riguarda solo il risultato finale, ma anche la struttura del processo decisionale: la presenza dell'IA riduce la distanza cognitiva tra profili tecnici e commerciali, portando entrambi i ruoli a generare soluzioni più bilanciate e meno condizionate dalle proprie identità funzionali. Questo fenomeno riflette un ampliamento effettivo dello spettro di conoscenze operative disponibili ai team, che diventa particolarmente visibile nei gruppi con livelli di competenza eterogenei: i partecipanti meno esperti ottengono gli incrementi più pronunciati, mentre gli esperti vedono una crescita più contenuta ma comunque significativa. Degno di nota è anche l'impatto emotivo: l'IA è associata a un aumento sistematico delle emozioni positive, come entusiasmo e coinvolgimento, e a una riduzione di quelle negative quali ansia e frustrazione, indicando che l'interazione linguistica con il sistema artificiale riesce ad assorbire parte della funzione motivazionale normalmente offerta dal lavoro di squadra.

Un quadro complementare emerge dall'esteso programma sperimentale condotto da Boston Consulting Group<sup>32</sup>, che ha coinvolto 758 consulenti impegnati in compiti professionali simili a quelli svolti abitualmente nella loro attività. Anche qui l'IA ha mostrato di poter incrementare sensibilmente la

---

<sup>31</sup> Dell'Acqua, F., Ayoubi, C., Lifshitz, H., Sadun, R., Mollick, E., Mollick, L., Han, Y., Goldman, J., Nair, H., Taub, S., and Lakhani, K. (2025). *The Cybernetic Teammate: A Field Experiment on Generative AI Reshaping Teamwork and Expertise*. NBER Working Paper 33641 (2025), <https://doi.org/10.3386/w33641>.

<sup>32</sup> Dell'Acqua, F., McFowland III, E., Mollick, E. R., Lifshitz-Assaf, H., Kellogg, K., Rajendran, S., Kraye, L., and Candelon, F., and Lakhani, K. R. (2026). *Navigating the Jagged Technological Frontier: Field Experimental Evidence of the Effects of Artificial Intelligence on Knowledge Worker Productivity and Quality*. In corso di pubblicazione presso Organization Science.

prestazione: nei compiti situati “entro la frontiera” delle capacità dei modelli generativi, la qualità delle soluzioni aumenta di oltre il 40%, mentre la produttività cresce di circa il 25%, con gli effetti più marcati tra i professionisti con prestazioni iniziali più basse. Tuttavia, contrariamente a quanto osservato in P&G, quando i compiti richiedono analisi o inferenze situate “oltre frontiera”, l’IA riduce la correttezza delle risposte, con un effetto negativo particolarmente marcato tra i consulenti che tendono ad adottare passivamente l’output generato. Lo studio mostra inoltre che la collaborazione con l’IA rimodella profondamente il modo in cui i consulenti comunicano e coordinano il lavoro: le conversazioni diventano più orientate al task e meno interpersonali, aumenta la delega, si riduce il ricorso a revisioni manuali e cresce l’omogeneità semantica degli output. L’IA migliora la prestazione, e, insieme, trasforma la struttura del lavoro di squadra, rendendo più efficiente la produzione di contenuti, al prezzo però di una possibile riduzione della varietà creativa.

Infine, il caso di Siemens estende l’analisi al dominio della manifattura avanzata<sup>33</sup>, mostrando come l’integrazione di agenti generativi non si limiti alle attività cognitive ad alta intensità di conoscenza, ma possa supportare direttamente attività operative, di manutenzione e di ingegneria, contribuendo a ridefinire il lavoro di squadra tra operatori umani e sistemi tecnici. In collaborazione con aziende come Schaeffler, Siemens ha mostrato che l’agente IA riduce il tempo necessario per generare e validare codice, diminuisce la probabilità di errori, accelera la diagnostica dei guasti e permette anche a operatori con minore esperienza di eseguire attività complesse con il supporto del sistema. Nella Electronics Factory di Erlangen, un agente IA collocato direttamente nel reparto produttivo fornisce istruzioni contestuali per l’individuazione, la spiegazione e la risoluzione dei malfunzionamenti, traducendo codici macchina in diagnosi comprensibili e suggerendo soluzioni basate su documentazione tecnica e storico operativo. L’esperimento mostra che l’IA non si limita a migliorare la produttività o la rapidità di intervento, ma consente una trasmissione immediata della conoscenza tacita e riduce la dipendenza da pochi esperti altamente specializzati, un fattore critico in un settore caratterizzato da competenze scarse. In questo ambiente, la collaborazione umano-IA assume la forma di un lavoro di squadra “intermittente”, in cui la macchina interviene come esperto di riferimento nei momenti di incertezza o complessità, contribuendo a stabilizzare l’intero sistema produttivo.

Per quanto riguarda il settore pubblico, sono presenti due evidenze entrambe provenienti dal Regno Unito. Un esperimento trasversale<sup>34</sup> ha coinvolto oltre 20.000 dipendenti per tre mesi con l’obiettivo di misurare l’impatto dell’IA generativa - Microsoft 365 Copilot - sui compiti quotidiani di policy drafting, gestione delle consultazioni, corrispondenza e reportistica. I risultati ufficiali comunicati dal governo indicano risparmi medi di 26 minuti al giorno per persona (equivalenti a quasi due settimane l’anno), con riscontri positivi su adozione e soddisfazione, e con casi d’uso concreti in strutture come Companies House e il Department for Work and Pensions; una quota molto ampia di partecipanti ha

---

<sup>33</sup> Berndt, J., Englmaier, F., Sadun, R., Tamayo, J., & von Hesler, N. (2026). *A Systematic Approach to Experimenting with Gen AI*. Harvard Business Review, January–February 2026

<sup>34</sup> <https://www.gov.uk/government/news/landmark-government-trial-shows-ai-could-save-civil-servants-nearly-2-weeks-a-year>

dichiarato che non vorrebbe tornare al periodo precedente all'IA, segnalando un effetto di "adesione" oltre che di efficienza.

Uno studio metodologico aggiuntivo del Government Digital Service (GDS), pubblicato in parallelo al comunicato<sup>35</sup>, documenta la distribuzione delle licenze, la raccolta dei dati quantitativi e qualitativi, l'adozione trasversale tra amministrazioni e i limiti incontrati nelle attività più complesse o *intensive di dati*. Il rapporto evidenzia inoltre come i benefici siano emersi in modo più netto laddove l'assistente è stato integrato nei flussi reali (ricerca, sintesi, redazione) e in presenza di una governance dei dati e pratiche di sperimentazione coerenti con i requisiti del settore pubblico.

### **Key Insight**

L'IA, inserita nei flussi di lavoro reali, agisce come partner cognitivo e operativo, aumentando qualità, velocità e accesso all'expertise. I benefici sono massimi entro la frontiera di capacità dei modelli e con un controllo umano strutturato, mentre la dipendenza eccessiva genera rischi. Oltre alla produttività, l'IA trasforma il lavoro di squadra: riduce i divari di competenze, stabilizza i processi e migliora coinvolgimento e coordinamento.

---

<sup>35</sup> Government Digital Service (GDS). *Microsoft 365 Copilot Experiment: Cross-Government Findings Report*. <https://www.gov.uk/government/publications/microsoft-365-copilot-experiment-cross-government-findings-report/microsoft-365-copilot-experiment-cross-government-findings-report-html>

### 3.3 LA ACTION-RESEARCH CON CDP: DISEGNO, INTERVENTO E RISULTATI

L'attività intrapresa dal gruppo di ricerca si configura come un intervento di *action-research* finalizzato a esplorare in modo controllato l'introduzione di un agente basato su intelligenza artificiale generativa a supporto delle attività legate al *risk assessment* presso la funzione Internal Audit di CDP. L'obiettivo principale del disegno di ricerca è stato valutare l'impatto dell'agente sui processi di lavoro, sulle percezioni individuali e sulle pratiche operative dei partecipanti, integrando misure quantitative e contributi qualitativi raccolti nel corso della sperimentazione.

Lo sviluppo dell'agente è frutto di un lavoro sinergico tra la funzione di **Innovation & Emerging Technologies**, **Internal Audit** e un **partner tecnologico esterno**. Quest'ultimo è stato incaricato della programmazione dell'agente con l'obiettivo di disporre di una soluzione sufficientemente matura da poter essere impiegata in attività reali di valutazione del rischio. L'agente è stato progettato per supportare le fasi del processo di *risk assessment*, integrandosi nei flussi di lavoro esistenti e operando come assistente specializzato a supporto delle decisioni degli utenti.

Il gruppo di ricerca ha avuto il ruolo di progettare e condurre la valutazione dell'intervento. Il disegno valutativo adottato è di tipo *before–after*, basato sulla somministrazione di un questionario a sette membri del team di Internal Audit, svolta in due momenti distinti: prima dell'utilizzo dell'agente e al termine della sperimentazione. La fase in cui gli auditor hanno utilizzato l'agente è coincisa con il suo fine tuning, consentendo loro di sperimentarne debolezze e potenzialità. Il questionario iniziale (*Inputs*) è stato somministrato a fine febbraio 2026 e comprendeva principalmente domande aperte qualitative, mentre il questionario finale (*Outputs*) è stato raccolto intorno al 10 aprile 2026 e comprendeva sia domande a risposta aperta sia questionari psicometrici adatti ad analisi quantitative.

A metà della sperimentazione è stata inoltre condotta una sessione intermedia di raccolta feedback sull'utilizzo dell'agente, organizzata sotto forma di *incontro online (Mediators)*. Questo momento ha svolto la funzione di integrare i dati quantitativi con evidenze qualitative, permettendo di raccogliere a voce impressioni, criticità, casi d'uso emergenti e il modo in cui gli utenti stavano integrando l'agente nel proprio lavoro. La sessione intermedia ha svolto anche un ruolo riflessivo tipico dei disegni di action-research, consentendo di osservare come le pratiche di utilizzo e le aspettative dei partecipanti evolvessero durante l'intervento.

#### Limiti e perimetro del disegno

Il disegno *before-after* adottato non prevede un gruppo di controllo; pertanto, i risultati sono da intendersi come descrittivi e non consentono generalizzazioni statistiche. Quello che abbiamo raccolto è il resoconto strutturato di una conversazione professionale: sette auditor che hanno provato a lavorare con un agente nel proprio processo reale e ne hanno discusso, prima individualmente

attraverso i questionari, poi collettivamente nella sessione intermedia. Le evidenze quantitative che seguono vanno quindi lette alla luce della dimensione e della natura esplorativa del campione: indicano tendenze interne al gruppo osservato e vanno interpretate con la prudenza che uno studio di questa scala richiede.

### 3.3.1 Questionario pre-sperimentazione

---

Appena prima della sperimentazione, i partecipanti hanno risposto al seguente modulo di domande aperte, finalizzato a raccogliere in forma qualitativa percezioni, aspettative e atteggiamenti rispetto al lavoro in un team agentico. Alcune domande, pur formulate in forma diretta, erano pensate come stimolo a una risposta argomentata, non come quesiti a risposta chiusa:

**Immaginando la tua esperienza di lavoro in un team con un agente IA, cosa credi rispetto ai seguenti punti?**

- a) Come influenzerà il modo in cui il team lavora insieme?
- b) Quali saranno i vantaggi e gli svantaggi di lavorare in questo tipo di team?
- c) Cambierà il tuo senso di fiducia, collaborazione o comunicazione all'interno del team?
- d) Cambierà il modo in cui interagisci con gli altri membri del team?
- e) Cambierà il modo in cui percepisci il tuo contributo o quanto ti sentirai soddisfatto del lavoro?

### 3.3.2 Questionario post-sperimentazione

---

Alla fine della sperimentazione, i partecipanti hanno risposto a una serie di questionari psicometrici e domande aperte sull'esperienza condotta.

Il primo modulo misura l'attività svolta con l'agente in team:

**Durante questa esperienza, quanto avete lavorato attivamente insieme all'agente IA per...**

1. Identificare i compiti principali
2. Identificare le sfide chiave da affrontare
3. Determinare le risorse per avere successo
4. Definire gli obiettivi del team
5. Sviluppare una strategia complessiva per guidare le attività del team
6. Preparare piani di contingenza per situazioni incerte
7. Sapere quando mantenere un piano di lavoro e quando adottarne uno diverso
8. Bilanciare i carichi di lavoro tra i membri del team
9. Aiutarsi reciprocamente quando è necessario
10. Sviluppare fiducia nella capacità del team di ottenere buone prestazioni
11. Integrare senza problemi gli sforzi lavorativi
12. Mantenere l'armonia nel gruppo

I partecipanti hanno risposto su una scala Likert a 5 punti con le seguenti opzioni di risposta:

**"Per niente" - "Molto poco" - "In una certa misura" - "In larga misura" - "In larghissima misura"**

Il secondo modulo misura il benessere percepito all'interno del team dopo l'introduzione dell'agente, con particolare attenzione a engagement, qualità delle relazioni e fiducia nella modalità di lavoro:

**Indica quanto sei in accordo o in disaccordo in base alle tue percezioni ed esperienze all'interno del team**

1. Vorrei continuare a lavorare con questa modalità su progetti futuri
2. Le relazioni di lavoro all'interno del team saranno collaborative ed efficaci in futuro
3. Mi sento a mio agio nel collaborare con strumenti di questo tipo nel team
4. Ritengo che questa modalità migliori la qualità delle decisioni del team
5. La collaborazione tra persone e agenti genera un vantaggio concreto per il team in termini di risultati e qualità del lavoro

I partecipanti hanno risposto su una scala Likert a 5 punti con le seguenti opzioni di risposta:

**"Fortemente in disaccordo" - "In disaccordo" - "Né d'accordo né in disaccordo" - "D'accordo" - "Fortemente d'accordo"**

Infine, i partecipanti hanno risposto alle seguenti domande:

**Riflettendo sulla tua esperienza di lavoro con un agente IA:**

- a) In che modo ha influenzato il modo in cui il team ha lavorato insieme?
- b) Quali sono i vantaggi e gli svantaggi di lavorare in questo tipo di team misto di umani e agenti IA?
- c) Ha influenzato il tuo senso di fiducia, collaborazione o comunicazione all'interno del team?
- d) Ha cambiato il modo in cui interagisci con gli altri membri del team?
- e) Ha cambiato il modo in cui percepisci il tuo contributo o quanto ti senti soddisfatto del lavoro?

### 3.3.3 Risultati qualitativi before-after

L'analisi qualitativa dei questionari *before–after* consente di osservare come le aspettative iniziali dei partecipanti rispetto all'introduzione di un agente IA nel processo di *risk assessment* si siano progressivamente trasformate a seguito dell'esperienza concreta di utilizzo. Le due word cloud riportate più avanti (Figura 1 per la fase iniziale, Figura 2 per quella successiva) ne offrono una sintesi visiva. L'analisi è stata condotta sia in forma aggregata sia in chiave longitudinale sugli stessi partecipanti, integrando letture tematiche e confronto semantico tra fase pre e post intervento.

Nella fase **pre-intervento**, le aspettative risultano prevalentemente orientate a una visione dell'IA come strumento di supporto all'efficienza operativa. I temi più ricorrenti riguardano la riduzione dei tempi di lavoro, l'automatizzazione di attività ripetitive, la sintesi di informazioni complesse e il miglioramento della qualità degli output, subordinato a un necessario riesame umano. Parallelamente, emerge con

forza l'attenzione al controllo e alla validazione, insieme alla preoccupazione per un possibile eccessivo affidamento sulla tecnologia. La *word cloud* pre-intervento (Figura 1) rende visibile la centralità di concetti quali *efficienza*, *supporto*, *tempo*, *qualità* e *output*, coerenti con una rappresentazione dell'agente come risorsa complementare e non sostitutiva. Questa impostazione è ben sintetizzata da affermazioni quali:

“L'IA risulta efficace quando consente di elevare la qualità delle attività svolte, fermo restando il necessario riesame umano.”

“L'agente deve essere al servizio del team e non in sostituzione del contributo umano.”

Aspettative PRE-intervento

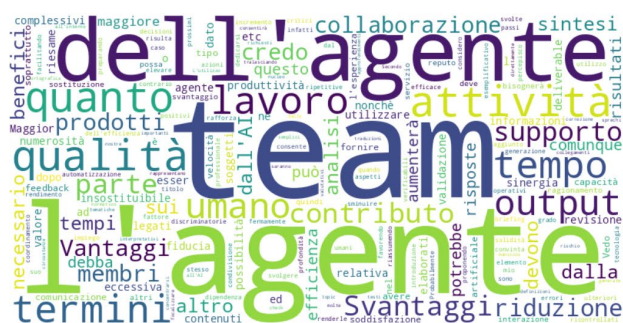


Figura 1

Riflessioni POST-intervento

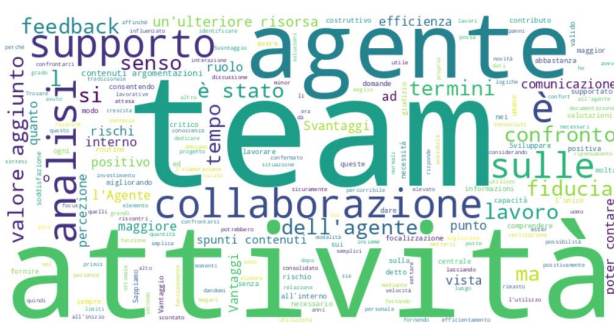


Figura 2

Nel **post-intervento**, le riflessioni mostrano uno spostamento semantico netto dalle aspettative astratte a valutazioni fondate sull'esperienza d'uso. I benefici dell'agente vengono descritti in modo più concreto e operativo, con riferimenti espliciti a *feedback*, *valore aggiunto*, *supporto* e *collaborazione*. Il tema del rischio rimane presente, riformulato però in termini più maturi: da timore generico ad esigenza di uso critico e consapevole dell'agente. La *word cloud* post-intervento (Figura 2) evidenzia l'emergere di un lessico più orientato al confronto e alla collaborazione, con parole come *confronto*, *feedback*, *analisi* e *collaborazione*, suggerendo un'evoluzione verso pratiche di lavoro più strutturate e condivise. Alcune risposte esprimono chiaramente questo cambiamento:

“Sappiamo di poter contare su spunti e analisi di un'ulteriore risorsa, da valutare sempre con spirito critico.”

“L'agente ha facilitato un confronto più strutturato sulle valutazioni di rischio, senza sostituire il giudizio professionale.”

Un elemento particolarmente rilevante è che, contrariamente alle aspettative iniziali, l'introduzione dell'agente ha prodotto effetti anche sulle dinamiche tra partecipanti. Diversi partecipanti segnalano che l'uso dell'IA ha stimolato momenti di discussione, coordinamento e apprendimento collettivo, rendendo le interazioni più orientate al contenuto e al merito delle valutazioni. L'agente, oltre a supportare le attività individuali, ha agito come catalizzatore del lavoro di gruppo:

“Testando l’agente tutti insieme abbiamo avuto la possibilità di confrontarci su come migliorare le attività.”

“Il lavoro con l’IA ha reso il dialogo più focalizzato e ordinato.”

Pur con i limiti di un campione esplorativo e ristretto, i risultati qualitativi indicano che la sperimentazione ha confermato le attese di miglioramento dell’efficienza, ma ha anche prodotto un apprendimento organizzativo. L’agente IA viene progressivamente integrato come risorsa cognitiva aggiuntiva, in grado di supportare il confronto interno e di rafforzare la focalizzazione sui giudizi di rischio, senza mettere in discussione il ruolo centrale del contributo umano nelle decisioni finali.

### 3.3.4 Risultati quantitativi

---

L’analisi quantitativa post-sperimentazione consente di integrare e consolidare le evidenze emerse dall’analisi qualitativa *before–after*, offrendo una lettura sistematica delle modalità con cui l’agente IA è stato incorporato nei processi di lavoro del team e delle percezioni associate a tale integrazione. I risultati relativi all’intensità della collaborazione tra partecipanti e agente mostrano che l’agente ha partecipato in misura significativa a un’ampia gamma di processi. Come illustrato in **Figura 3**, le valutazioni sui processi di team si collocano sistematicamente tra *in larga misura* e *in larghissima misura* per numerose attività centrali del funzionamento del gruppo. Tra queste figurano sia processi prevalentemente orientati al compito, come l’identificazione dei compiti principali, la definizione degli obiettivi del team, lo sviluppo di una strategia complessiva e la preparazione di piani di contingenza, sia processi più relazionali e di integrazione, quali il mantenimento dell’armonia e dell’equilibrio nel team.

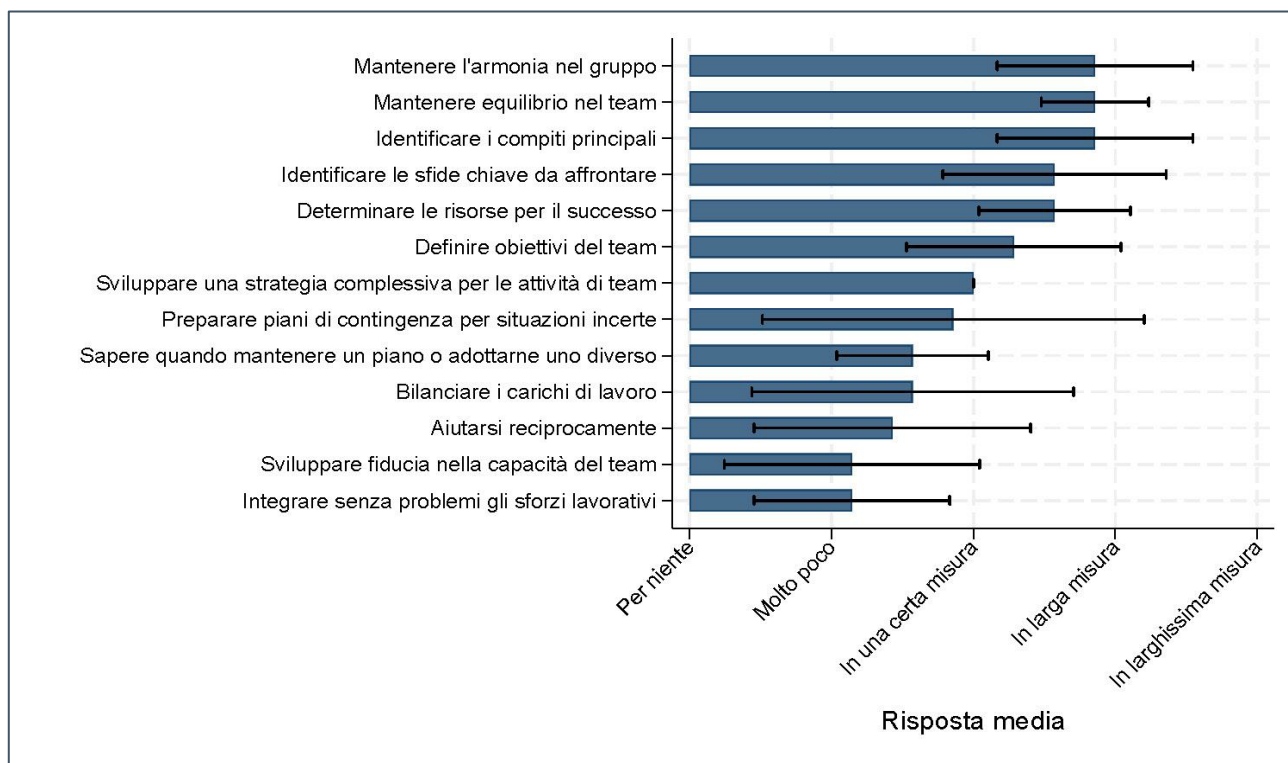


Figura 3. Grafico a barre indicante la risposta media alla domanda **“Durante questa esperienza, quanto avete lavorato attivamente insieme all'agente IA per...”** e intervalli di confidenza al 95%

I processi sui quali i partecipanti non hanno lavorato attivamente riguardano invece l'aiuto reciproco, il bilanciamento dei carichi di lavoro, ma anche la pianificazione. Il dato suggerisce che, nel breve periodo, l'agente è stato usato soprattutto come supporto operativo all'esecuzione e non ancora come membro effettivo del team. Come osservato in precedenza, questo richiede un più lungo processo di inserimento e di comprensione delle dinamiche e del potenziale dei team agentici.

Le valutazioni percettive dei partecipanti restituiscono un quadro complessivamente positivo rispetto agli esiti dell'esperienza. **Figura 4** sintetizza le risposte relative alla qualità delle decisioni, alla collaborazione futura e al benessere percepito dal team (engagement, qualità delle relazioni e fiducia). Tutti gli elementi considerati presentano valori medi superiori al punto neutro della scala, con una concentrazione delle risposte compresa tra *d'accordo* e *fortemente d'accordo*.

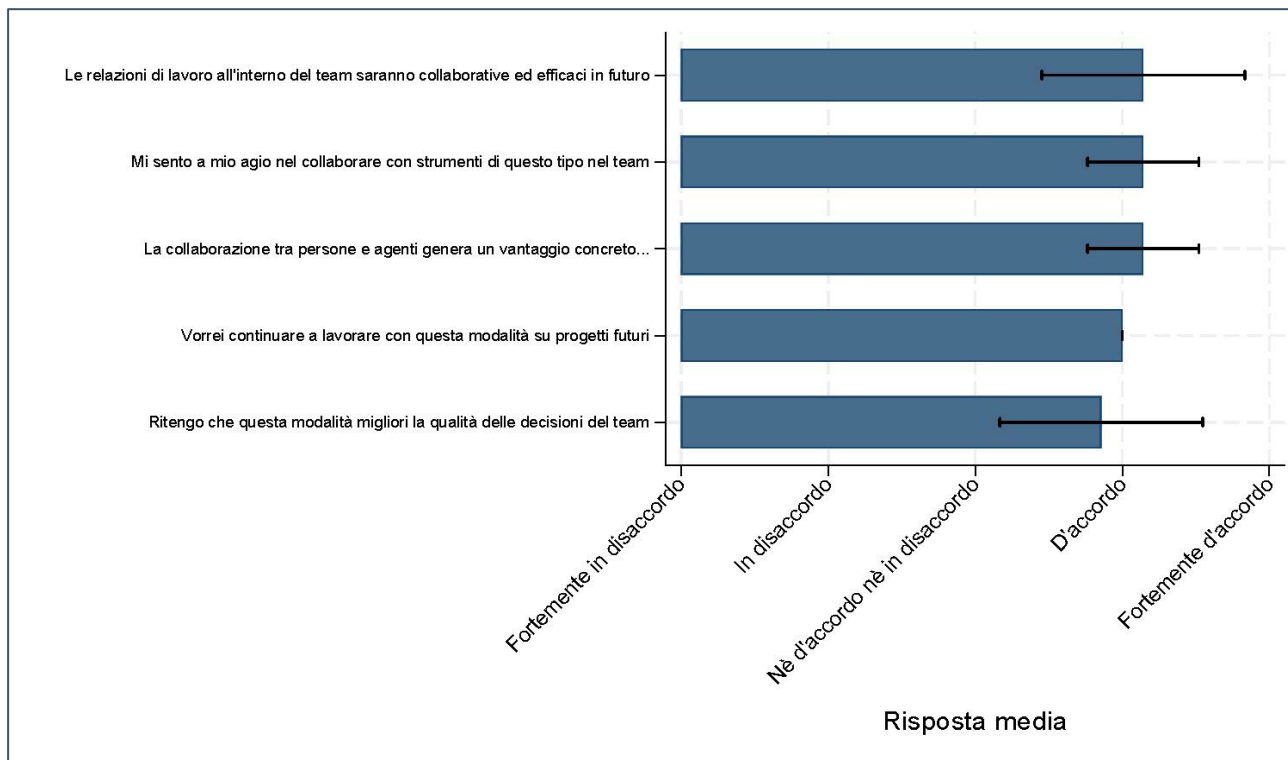


Figura 4. Grafico a barre indicante la risposta media alla domanda **"Indica quanto sei in accordo o in disaccordo in base alle tue percezioni ed esperienze all'interno del team"** e intervalli di confidenza al 95%

In particolare, i partecipanti esprimono un giudizio favorevole rispetto alla capacità della modalità di lavoro sperimentata di migliorare la qualità delle decisioni e di generare un vantaggio concreto in termini di risultati e qualità del lavoro. Allo stesso tempo, emerge una disponibilità elevata a continuare a lavorare con questa configurazione su progetti futuri e una valutazione positiva delle relazioni di lavoro attese all'interno del team.

### 3.4 PATTERN RICORRENTI E VARIABILI CONTINGENTI: CONVERGENZE, DIVERGENZE E GAP DI RICERCA

Le evidenze empiriche, i casi aziendali e la *action-research* presentati in questo capitolo convergono nel restituire un quadro unitario ma non semplificabile dei team aumentati. L'elemento comune che attraversa contesti molto diversi — dal lavoro della conoscenza ad alta intensità cognitiva, alla consulenza strategica, alla manifattura avanzata e alla pubblica amministrazione — è che l'introduzione dell'IA non produce effetti lineari o automatici sulla prestazione dei team. L'IA agisce piuttosto come **amplificatore delle condizioni socio-tecniche preesistenti**, rafforzando assetti ben progettati e rendendo più fragili quelli opachi o mal allineati.

Un primo elemento ricorrente riguarda il **ruolo dell'IA come estensione delle competenze**, piuttosto che come sostituto. Nei contesti in cui l'agente è inserito entro flussi di lavoro chiari, con confini di responsabilità espliciti e momenti strutturati di revisione umana, l'IA funziona da catalizzatore: accelera l'analisi, amplia il ventaglio di alternative esplorabili e rende più accessibili competenze rare o distribuite in modo asimmetrico. Questo effetto è particolarmente evidente nei casi P&G, BCG e Siemens, così come nell'esperienza CDP, dove l'agente viene progressivamente interpretato come risorsa cognitiva aggiuntiva che supporta — senza soppiantarlo — il giudizio professionale.

Un secondo elemento riguarda la **dipendenza dei risultati dalla collocazione del compito rispetto alla frontiera di capacità dell'IA**. Le evidenze mostrano con coerenza che i benefici emergono “entro frontiera”, mentre l'utilizzo dell'IA su compiti mal definiti, ad alta ambiguità o fuori dominio produce errori, affidamento eccessivo e peggioramenti di performance. L'eterogeneità degli esiti è una proprietà strutturale del lavoro di squadra umano-IA: la qualità dell'interazione dipende dalla capacità del team di riconoscere limiti, calibrare l'affidamento e progettare passaggi di consegne appropriati.

Un terzo elemento trasversale riguarda le **dinamiche relazionali e identitarie**. In nessuno dei casi analizzati l'introduzione dell'IA è “neutra” per il funzionamento sociale del team. La presenza dell'agente modifica comunicazione, coordinamento e percezione del contributo individuale. Dove queste dimensioni vengono presidiate — ad esempio attraverso processi di inserimento, pratiche riflessive e spazi di confronto — l'IA può rafforzare il lavoro di gruppo, rendendo le discussioni più focalizzate e orientate al merito. Dove invece l'automazione è trattata come semplice acceleratore operativo, emergono rischi di impoverimento del dialogo, disimpegno silenzioso e riduzione della sicurezza psicologica.

## 4.1 CREAZIONE DI VALORE ECONOMICO: EFFICIENZA, DECISION-MAKING SUPERIORE, INNOVAZIONE CONTINUA

Quanto esposto fin qui evidenzia un punto importante: nei team aumentati l'oggetto della progettazione è l'assetto complessivo di ruoli, passaggi di consegna, diritti di veto, escalation, monitoraggio e responsabilità che rende l'azione collettiva osservabile e governabile, più del singolo agente. È la stessa lente con cui vanno lette le opportunità emergenti. Non si tratta dell'arrivo di un nuovo software. Si tratta di una fase in cui una quota crescente di lavoro cognitivo codificabile, cioè quello che si può descrivere in passaggi ripetibili e verificabili come la ricerca documentale, la prima stesura di un testo o il controllo di conformità, può essere delegata, orchestrata e in alcuni casi automatizzata dentro flussi di lavoro reali. La domanda rilevante riguarda cosa convenga far emergere da un cambiamento ormai in corso.

### 4.1.1 Dove si genera realmente valore

---

Il primo equivoco da evitare è anche il più diffuso: il valore economico nasce quando l'organizzazione ripensa porzioni concrete dei suoi flussi di lavoro principali. Questo punto emerge già da quanto visto nelle sezioni precedenti, dove si osserva che i benefici misurabili compaiono soprattutto quando la tecnologia è inserita nei processi centrali e sostenuta da risorse complementari adeguate. Le evidenze più recenti vanno in questa direzione. Nel Work Trend Index 2025 Microsoft<sup>36</sup> descrive un divario di capacità ormai strutturale: il 53% dei leader afferma che la produttività deve aumentare, mentre l'80% dei lavoratori dichiara di non avere tempo o energia sufficienti per il lavoro quotidiano; nella stessa ricerca, l'82% dei leader si aspetta di usare "forza lavoro digitale" per espandere capacità nei successivi 12-18 mesi. Non è ancora una prova di trasformazione riuscita. È però un segnale chiaro della pressione economica che sta spingendo la riorganizzazione del lavoro cognitivo.

### 4.1.2 Efficienza non significa solo risparmio di tempo

---

Quando l'IA funziona, l'effetto più immediato è il recupero di tempo e attenzione. Ridurre tutto a una metrica di minuti risparmiati sarebbe riduttivo, e fuorviante. Come abbiamo visto, l'esperimento trasversale coordinato dal Government Digital Service britannico<sup>37</sup> ha rilevato un risparmio medio di 26 minuti al giorno per persona, con oltre il 70% degli utenti che segnala meno tempo speso in ricerca di informazioni e attività ripetitive e più tempo per attività strategiche. Il dato è importante non tanto per la cifra in sé, quanto per ciò che implica: **l'efficienza più interessante è quella che cambia la composizione del lavoro**. In parallelo, il caso della action-research CDP suggerisce una dinamica

---

<sup>36</sup> Microsoft, 2025 Work Trend Index Annual Report: The Year the Frontier Firm Is Born (2025).

<sup>37</sup> Government Digital Service (GDS). *Microsoft 365 Copilot Experiment: Cross-Government Findings Report*

simile: l'agente viene progressivamente percepito come risorsa cognitiva aggiuntiva che rende il confronto più focalizzato e ordinato. Si intravede quindi una distinzione cruciale per il decisore: un'organizzazione può usare l'IA per fare le stesse cose più in fretta; oppure può usarla per liberare attenzione e allocarla verso compiti che richiedono giudizio, priorità e responsabilità. Le due traiettorie portano a esiti molto diversi, e la scelta tra le due raramente è dichiarata in modo esplicito.

#### 4.1.3 La qualità delle decisioni dipende dall'orchestrazione

---

Lo stesso vale per il processo decisionale. La letteratura richiamata nelle sezioni precedenti ha già mostrato che l'IA può migliorare la qualità delle decisioni entro una frontiera di capacità ben delimitata, e può invece peggiorarla quando si scivola in affidamento passivo o in compiti fuori frontiera. Lo studio BCG già illustrato resta, da questo punto di vista, uno snodo importante: nei task che rientrano nelle capacità del modello, qualità e produttività crescono in modo rilevante; nei compiti fuori frontiera, la correttezza si riduce e l'affidamento eccessivo diventa un rischio sistemico. Caplin e colleghi<sup>38</sup> aggiungono un tassello decisivo: a parità di abilità di base, trae più vantaggio dall'IA chi possiede una buona calibrazione delle proprie capacità, cioè una valutazione realistica di ciò che sa e non sa fare. È una proprietà operativa del sistema umano-IA. L'IA migliora le decisioni quando entra in una catena decisionale ben progettata, con soglie di fiducia, momenti di revisione e criteri chiari per fermare, verificare, correggere.

#### 4.1.4 Innovazione continua, ma con una domanda più esigente

---

Anche riguardo all'innovazione è opportuno evitare tanto l'enfasi facile quanto lo scetticismo automatico. L'IA amplia il numero di alternative esplorabili, accelera simulazioni, sintesi, prime stesure, ricombinazioni. In questo senso, può sostenere un'innovazione più continua, meno episodica. L'innovazione che conta per le organizzazioni è, dunque, **la capacità di riconoscere quali tra queste alternative meritino attenzione, investimento, rischio**. Ed è precisamente qui che il contributo umano torna al centro. Nel caso CDP l'agente non sostituisce il giudizio professionale: lo stimola, lo rende più articolato, e lo espone a un confronto più strutturato. Il quadro concettuale di Acemoglu, Autor e Johnson<sup>39</sup> aiuta a non perdere la direzione: c'è una differenza tra usi dell'IA che comprimono il lavoro e usi che estendono le competenze, accelerano l'apprendimento e aprono nuovi compiti. Per un decisore è forse il punto più scomodo, ma anche il più importante. Non tutto il valore economico è uguale. Alcune forme di efficienza impoveriscono il sistema che le genera. Altre lo rendono più capace. È questa seconda famiglia di opportunità che dovrebbe interessare davvero.

---

<sup>38</sup> Caplin, A., Deming, D. J., Li, S., Martin, D. J., Marx, P., Weidmann, B., & Ye, K. J. (2024). The ABC's of Who Benefits from Working with AI: Ability, Beliefs, and Calibration. NBER Working Paper No. 33021

<sup>39</sup> Acemoglu, D., Autor, D., & Johnson, S. (2026). Building Pro-Worker Artificial Intelligence. NBER Working Paper 34854 (2026), <https://doi.org/10.3386/w34854>.

## 4.2 DOMINI AD ALTO POTENZIALE: CRITERI DI INDIVIDUAZIONE, DRIVER DI ADOZIONE

### 4.2.1 Non i domini “più digitali”

---

Come abbiamo visto, i domini ad alto potenziale, cioè gli ambiti di applicazione, non coincidono necessariamente con quelli più vicini, per immagine, al mondo dell'IA. Coincidono piuttosto con quelli in cui una parte rilevante del lavoro cognitivo è scomponibile, documentale, frequente e verificabile. È questa architettura del compito — più della sola sofisticazione tecnologica del settore — a rendere possibile una buona integrazione degli agenti. Lo abbiamo già osservato parlando di curve di adozione settoriali e di contesti nei quali la sperimentazione passa più rapidamente alla normalizzazione. I primi terreni maturi sono spesso quelli dove il lavoro si appoggia a basi documentali ampie, dati strutturati, procedure ricorrenti e risultati sottoponibili a riesame: servizi professionali, operations, funzioni di supporto, parti della finanza, della pubblica amministrazione (PA) e della manifattura avanzata. Il punto è scegliere i contesti dove il lavoro può essere davvero riorganizzato.

### 4.2.2 Il secondo criterio: il sovraccarico

---

C'è poi un secondo criterio, meno tecnico ma spesso più decisivo: il divario di capacità. I domini più promettenti sono spesso quelli in cui il carico cognitivo eccede già la capacità umana sostenibile. Qui l'IA entra perché le persone sono già oltre soglia. Il Work Trend Index 2025 lo descrive senza ambiguità: lavoro frammentato, interruzioni continue, pressione a produrre di più con attenzione sempre più dispersa. In questo contesto, le opportunità più concrete non riguardano solo produttività o riduzione dei costi. Riguardano la possibilità di restituire sostenibilità al lavoro quotidiano. In funzioni come compliance, supporto documentale, redazione, ricerca interna, analisi del rischio o gestione delle relazioni con i clienti, la domanda non è solo “possiamo automatizzare?” È anche “quanta parte del lavoro attuale è già oggi uno spreco di concentrazione umana?” Vista così, una quota dei domini ad alto potenziale è semplicemente fatta di domini in sofferenza.

Vale però la pena ricordare che non tutti i contesti sono ugualmente adatti all'integrazione di agenti IA. Le stesse evidenze empiriche che documentano benefici in alcuni ambienti mostrano effetti negativi in altri. Lo studio Super Mario Meets AI e la ricerca di Qin e Sajda evidenziano che, in team caratterizzati da alta interdipendenza, routine tacite consolidate e compiti difficilmente verificabili, l'introduzione di un agente può disallineare il coordinamento e ridurre la prestazione collettiva. Per un decisore, questo significa che la domanda rilevante non è solo “questo dominio ha potenziale?”, ma anche “questo team dispone delle condizioni organizzative necessarie per integrare un agente senza compromettere coordinamento e qualità del giudizio?”. Un agente IA non è una soluzione valida ovunque: dove

mancono queste condizioni, può peggiorare il lavoro invece di migliorarlo. Per un'analisi più dettagliata di queste evidenze si rimanda alla sezione 3.1.

### **4.2.3 Il vero salto: interoperabilità e governance**

---

Negli ultimi mesi questo punto è diventato ancora più evidente: un ambito si afferma davvero quando agenti IA, dati e strumenti iniziano a comunicare con protocolli relativamente comuni. L'introduzione di MCP da parte di Anthropic nel novembre 2024 ha dato forma a uno standard aperto per collegare applicazioni IA a dati, strumenti e flussi di lavoro esterni; nel 2025 Anthropic ha descritto MCP come standard di fatto per connettere agenti a sistemi esterni, e oggi la documentazione del protocollo lo presenta come un'interfaccia comune già supportata da un ecosistema ampio di client e server. In parallelo, il protocollo A2A ha raggiunto la prima specifica stabile e, a un anno dal lancio, ha superato le 150 organizzazioni di supporto. Il problema dell'interoperabilità resta aperto, ma ha smesso di essere puramente teorico. Per molti domini, la differenza tra pilota interessante e filiera operativa sta precisamente qui: nella possibilità di comporre agenti, strumenti e fonti dati senza ricominciare ogni volta da zero.

### **4.2.4 Nei domini regolati l'opportunità cresce solo se cresce la verificabilità**

---

Tutto questo vale ancora di più nei contesti regolati. Sanità, finanza, utilities, PA, infrastrutture critiche: sono domini ad altissimo potenziale proprio perché il loro lavoro è intensivo di conoscenza, coordinamento e documentazione. Sono anche domini in cui un errore è un problema di responsabilità, reputazione, conformità e talvolta di sicurezza pubblica. Per questo, nei contesti regolati, l'automazione potente non basta. Serve automazione verificabile. Il Government Digital Service britannico osserva che i benefici di Copilot sono emersi più chiaramente dove lo strumento era integrato in flussi di lavoro reali di ricerca, sintesi e redazione, ma anche dove erano presenti una governance dei dati, limiti d'uso coerenti con il settore e chiara consapevolezza dei casi in cui il giudizio umano restava necessario. Se non è verificabile, il dominio non è ancora davvero maturo.

### **4.2.5 Il driver più decisivo è la capacità di sperimentare**

---

Resta infine il fattore che molte organizzazioni continuano a sottovalutare: la qualità della sperimentazione. HBR, in un articolo importante di inizio 2026<sup>40</sup>, insiste sul fatto che il valore dell'IA Generativa emerge da esperimenti organizzativi seri, capaci di testare bisogni autentici, criteri di successo, adattamenti di processo e condizioni di scala: lo strumento di assessment viene trattato sempre come strumento di apprendimento continuo. L'adozione dell'IA fallisce raramente per limiti tecnici. Più spesso fallisce perché viene gestita con una cultura della sperimentazione debole, ansiosa

---

<sup>40</sup> Harvard Business Review. Systematic Approach to Experimenting with Gen AI <https://hbr.org/2026/01/a-systematic-approach-to-experimenting-with-gen-ai>

di mostrare novità ma poco interessata a imparare davvero. Le opportunità emergenti, invece, si aprono dove l'organizzazione sa osservare, correggere, insistere o fermarsi in tempo.

## 4.3 EVOLUZIONE DEI RUOLI E DELLE COMPETENZE: CAPABILITY EMERGENTI E REQUISITI SOCIO-TECNICI

### 4.3.1 Il lavoro quotidiano si sposta: meno esecuzione lineare, più regia

---

Se una parte crescente del lavoro cognitivo codificabile può essere assorbita dagli agenti, il lavoro umano non scompare ma cambia baricentro. Sempre più spesso il lavoro consiste nel gestire e coordinare quello che fanno i sistemi digitali, più che nell'eseguire in prima persona. Significa definire i criteri del compito, decidere cosa delegare e cosa no, interpretare i risultati, riconoscere le eccezioni, intervenire quando la macchina appare convincente ma sta andando nella direzione sbagliata. La metafora del “nuovo collega” usata da Microsoft<sup>41</sup> ha un evidente intento commerciale, ma coglie una verità organizzativa: lavorare con agenti richiede capacità di integrarli, indirizzarli, controllarli e accordare loro una fiducia informata.

### 4.3.2 Le capability decisive sono metacognitive

---

Per questo motivo, le competenze emergenti non coincidono con la sola AI literacy. Diventano centrali abilità meno appariscenti e più profonde: calibrazione, discernimento, critica dei risultati, capacità di sospendere l'automatismo, di chiedere spiegazioni, di ricondurre al controllo umano un passaggio delicato. Il contributo di Caplin e colleghi<sup>42</sup> è illuminante proprio perché mostra che i benefici dell'IA dipendono non solo dall'abilità, ma anche dalla qualità delle credenze che le persone hanno su sé stesse. Chi sa valutare meglio i propri limiti usa meglio l'IA, **perché sa quando fidarsi del suo output e quando verificarlo**. Il quadro sociotecnico indica la stessa direzione: nei team aumentati, la competenza distintiva sarà saper giudicare bene, più che saper formulare domande efficaci. Per un decisore, questo cambia immediatamente le priorità di formazione e selezione.

### 4.3.3 Cosa resterà umano: la traiettoria sta emergendo

---

Questo è un aspetto fondamentale, perché c'è un punto che riguarda l'idea stessa di lavoro, oltre l'efficienza. Se il lavoro cognitivo codificabile entra davvero in una traiettoria di automazione estesa, allora la domanda centrale diventa: **che cosa emergerà come contributo specificamente umano?** Non disponiamo ancora di una tassonomia stabilizzata, ma alcune linee si vedono già. Il Work Trend

---

<sup>41</sup> Microsoft, 2025 Work Trend Index Annual Report: The Year the Frontier Firm Is Born (2025).

<sup>42</sup> Caplin, A., Deming, D. J., Li, S., Martin, D. J., Marx, P., Weidmann, B., & Ye, K. J. (2024). The ABC's of Who Benefits from Working with AI: Ability, Beliefs, and Calibration. NBER Working Paper No. 33021.

Index 2025 insiste sul fatto che gli esseri umani restano distintivi per creatività, giudizio e capacità di costruire connessioni. Acemoglu, Autor e Johnson<sup>43</sup> spingono oltre, sostenendo che la traiettoria più desiderabile è quella in cui l'IA aumenti il valore delle capacità umane, crei nuovi compiti e acceleri l'apprendimento, invece di ridurre il lavoro a mera supervisione residua. Tradotto in termini organizzativi, questo suggerisce uno spostamento del valore verso ciò che non si delega: il giudizio situato, **la capacità di attribuire senso, il presidio delle relazioni e la definizione degli obiettivi**.

#### 4.3.4 Gli investimenti giusti non sono solo tecnologici

---

Di conseguenza, gli investimenti più importanti che emergono dalla ricerca non riguardano soltanto modelli, licenze o integrazioni: riguardano competenze umane e sistemi che le coltivano. Il Work Trend Index segnala che il 47% dei leader considera l'aggiornamento delle competenze una priorità, e che il 78% sta considerando ruoli specifici per l'IA. Il sistema IMOI utilizzato in questa action-research va in questa direzione, perché misura non solo i risultati ma anche **processi, stati emergenti, identità di team, sostenibilità del carico, disponibilità a continuare a lavorare insieme**. È una scelta con implicazioni manageriali e di indirizzo politico, prima ancora che metodologiche. Se l'automazione del lavoro della conoscenza cresce, allora diventa vitale investire in ciò che aiuta gli esseri umani a non diventare periferici nel proprio lavoro: formazione, pratiche di revisione, inserimento degli agenti, spazi di co-progettazione, metriche longitudinali, supporti alla sicurezza psicologica e all'apprendimento riflessivo. Si decide qui se i team aumentati diventeranno strutture più solide o soltanto più fragili e più veloci.

---

<sup>43</sup> Acemoglu, D., Autor, D., & Johnson, S. (2026). *Building Pro-Worker Artificial Intelligence*. NBER Working Paper 34854 (2026), <https://doi.org/10.3386/w34854>.



## 4.4 ORIZZONTI DI INNOVAZIONE: RETI DI AGENTI COLLABORATIVI, SYMBIOTIC TEAMS, ORGANIZZAZIONI ADATTIVE

### 4.4.1 Il passaggio in corso: dal copilota al workflow agentico

---

L'orizzonte più interessante oggi è il flusso di lavoro agentico. E qui, gli sviluppi più recenti sono importanti per la direzione che rendono visibile, più che per il clamore con cui vengono annunciati. OpenAI ha aggiornato il proprio Agents SDK con nuove capacità pensate per compiti a lungo orizzonte in ambienti controllati, con strumenti per ispezionare file, eseguire comandi, modificare codice e mantenere continuità di esecuzione in ambienti sicuri; Microsoft, nel proprio lessico, descrive un percorso che va dall'assistente individuale a team umano-agente fino a processi "guidati dagli umani, eseguiti dagli agenti", in cui gli esseri umani impostano direzione e soglie di intervento mentre gli agenti eseguono porzioni intere di flussi di lavoro. La traiettoria più ampia è leggibile: una parte crescente del lavoro cognitivo standardizzabile sta uscendo dall'orbita del supporto occasionale ed entrando in quella dell'automazione organizzata. Il cantiere è aperto, ed è molto più avanzato di quanto molti continuino a pensare.

### 4.4.2 Lo stack aperto che sta emergendo conta quasi quanto i modelli

---

Il vero avanzamento recente è nello stack, cioè l'insieme stratificato di protocolli, strumenti e standard su cui gli agenti vengono costruiti. MCP, A2A, AGENTS.md, registry, SDK, specifiche stabili, governance condivisa: siamo nel passaggio in cui l'IA agentica smette di essere solo abilità di un singolo sistema e comincia a diventare infrastruttura coordinabile. Anthropic ha introdotto MCP come standard aperto per connettere sistemi IA a dati e strumenti; un anno dopo il protocollo viene descritto, dalla stessa Anthropic e dalla Linux Foundation, come componente ormai centrale dell'ecosistema. A2A, da parte sua, ha superato in un anno le 150 organizzazioni di supporto e ha raggiunto la versione 1.0 come specifica stabile per l'interoperabilità tra agenti. La traiettoria per le organizzazioni è concreta: a prosperare saranno quelle che sapranno costruire un'architettura di interoperabilità affidabile, aperta, verificabile e abbastanza modulare da non restare vincolate a un singolo fornitore o ad un singolo disegno tecnico.

#### **Key Insight**

L'orizzonte emergente è un ecosistema di workflow agentici: il vantaggio dipenderà dalla qualità dell'infrastruttura, dell'interoperabilità e della governance che rendono gli agenti componibili e affidabili, prima ancora che dalla potenza del modello.

### 4.4.3 Anche l'hardware sta cambiando il perimetro del possibile

---

C'è poi il tema dell'infrastruttura fisica, spesso trascurato nei discorsi troppo concentrati sui modelli. In termini semplici, una quota crescente del calcolo richiesto dagli agenti si sta spostando verso ambienti locali e controllati, e questo rende l'hardware una scelta strategica. L'IA agentica significa anche più inferenza, più orchestrazione, più continuità di esecuzione, più bisogno di ambienti controllati e di accesso locale a dati sensibili. Gli annunci ufficiali di NVIDIA nell'ultimo anno mostrano una direzione coerente: Blackwell Ultra è presentato come piattaforma costruita per reasoning e IA agentica; DGX Spark e DGX Station portano capacità di sviluppo, prototipazione e inferenza dal desktop al data center; al GTC 2026 NVIDIA ha esplicitamente collegato DGX Station allo sviluppo di agenti autonomi a ragionamento prolungato, con supporto a modelli aperti fino a 1 trilione di parametri e possibilità di configurazioni isolate dalla rete per ambienti regolati. Non tutto questo diventerà subito prassi diffusa, ma la direzione è chiara. L'adozione aziendale dell'IA agentica dipenderà sempre meno dal solo modello cloud e sempre più dalla possibilità di collocare capacità di calcolo, memoria, sicurezza e continuità esecutiva nei luoghi giusti della catena del valore. L'hardware torna a essere una variabile strategica.

### 4.4.4 Che cosa si intravede per team e strutture

---

Se questa traiettoria prosegue, il cambiamento non riguarderà solo i compiti. Riguarderà la forma delle organizzazioni. Microsoft parla di possibile passaggio dalla struttura organizzativa alla mappa del lavoro: un linguaggio forse semplificante, ma utile per nominare una trasformazione reale. Quando competenze e capacità esecutive diventano disponibili su richiesta attraverso agenti specializzati, team e funzioni possono diventare più fluidi, più temporanei, più orientati a obiettivi che a confini stabili. In questo senso, abbiamo già descritto nelle precedenti sezioni modelli a regia, allocazione adattiva e livelli di coordinamento per la "gestione degli agenti" capaci di scomporre, assegnare, tracciare e riconfigurare il lavoro. Qui entrano in gioco i team simbiotici, non più semplici agenti aggiunti a un team esistente, ma una forma più stretta di interdipendenza in cui l'autonomia dell'agente viene modulata dinamicamente in base al contesto. I più promettenti saranno quelli che sapranno modulare l'autonomia in funzione del dominio, del rischio, della criticità situazionale e del tipo di contributo umano richiesto. È una distinzione essenziale. Conta delegare in modo consapevole e controllato, più che delegare molto.

### 4.4.5 La domanda finale, forse, è la più rilevante

---

Resta allora la domanda finale. Se il lavoro cognitivo codificabile viene automatizzato in misura crescente, che cosa deve emergere per gli esseri umani? Non basta rispondere: "più strategia". Come abbiamo visto fin qui, la ricerca più interessante comincia a indicare qualcosa di più sottile: maggiore centralità del discernimento, del significato attribuito al lavoro, della capacità di tenere insieme conseguenze, esperienza e relazione; maggiore importanza delle competenze che aiutano a orientare,

interpretare, apprendere, decidere sotto incertezza; maggiore necessità di investire in sistemi che sviluppino queste capacità, invece di consumarle. Acemoglu, Autor e Johnson<sup>44</sup> avvertono che il rischio strutturale dell'ecosistema attuale è una deriva verso l'automazione; HBR richiama la necessità di sperimentare a livello organizzativo<sup>45</sup>; la presente action-research in corso mostra che l'agente può diventare un catalizzatore del confronto, non solo un sostituto di compiti. La vera opportunità è costruire **organizzazioni in cui l'automazione del lavoro cognitivo renda finalmente più visibile, più sviluppabile e più prezioso ciò che negli esseri umani non è riducibile alla semplice esecuzione.**

---

<sup>44</sup> Acemoglu, D., Autor, D., & Johnson, S. (2026). *Building Pro-Worker Artificial Intelligence*. NBER Working Paper 34854 (2026), <https://doi.org/10.3386/w34854>.

<sup>45</sup> Harvard Business Review. Systematic Approach to Experimenting with Gen AI <https://hbr.org/2026/01/a-systematic-approach-to-experimenting-with-gen-ai>

## Conclusioni

Lavorando con gli agenti IA emerge presto una cosa: la tecnologia funziona meglio del previsto in alcuni contesti e peggio in altri, e la differenza dipende raramente dal modello. Dipende da come è organizzato il lavoro intorno ad esso.

È il filo che attraversa tutto il paper. Le evidenze empiriche analizzate, i casi aziendali e la sperimentazione condotta in CDP convergono su un punto comune. L'agente non riscrive il sistema di lavoro che trova: ne accentua la qualità o i difetti. Dove ruoli, criteri di escalation e supervisione sono chiari, i benefici sono reali. Dove processi e responsabilità sono opachi, le criticità preesistenti emergono amplificate invece di risolversi.

La sperimentazione in CDP, con tutti i limiti di uno studio esplorativo su sette partecipanti e un mese di osservazione, ha mostrato qualcosa di interessante. L'agente non ha sostituito il giudizio professionale, ma ne ha modificato il funzionamento. Ha reso il confronto più focalizzato, ha stimolato discussioni che altrimenti non sarebbero emerse, ha generato un apprendimento collettivo che va oltre il singolo utilizzo. Sono segnali preliminari, non evidenze definitive. Hanno però una caratteristica che vale la pena dichiarare: sono coerenti tra loro e coerenti con quanto la letteratura internazionale osserva su campioni molto più ampi. È questo, in uno studio piccolo, l'unico tipo di solidità a cui è onesto appellarsi.

Restano aperte le domande più difficili. Come cambia il giudizio professionale nel tempo, quando una parte crescente del lavoro cognitivo viene delegata a un agente? Quali configurazioni di team producono apprendimento duraturo invece di semplice efficienza di breve periodo? Come si mantiene l'accountability umana reale in flussi di lavoro sempre più automatizzati? E come devono evolvere le competenze di chi lavora, perché il contributo umano resti centrale invece di ridursi a supervisione residua?

Questo paper non risponde a queste domande in modo definitivo. Prova a formularle con la precisione che le evidenze oggi disponibili consentono. In una fase in cui molte organizzazioni stanno decidendo non se muoversi ma come, è il contributo che riteniamo più utile.

## Ringraziamenti

Questo documento è frutto del lavoro congiunto e delle analisi effettuate e condivise da un gruppo di lavoro ampio e interdisciplinare.

Si ringrazia:

- Il team Innovation & Emerging Technologies di CDP composto da: **Valeria de Flaviis** (Head of Innovation & Emerging Technologies), **Gian Matteo Guglielmi** (Head of Innovation Solutions & Emerging Technologies) e **Stefano Padoa** (Innovation Solutions & Emerging Technologies Senior Professional)
- Il team Internal Audit di CDP composto da: **Silvia Lini** (Head of Metodologie e Audit Lab), **Giampaolo Castellani** (Audit Data Scientist & Methodologist Senior Professional), **Gaia Tarsi** (Audit Data Scientist & Methodologist Senior Professional) e i membri del team Audit Operativo Centrale
- Cocoon Pro, in qualità di partner metodologico, nella persona di **Stelio Verzera** (Co-founder and Managing Partner)
- Università di Verona, in qualità di partner scientifico, nella persona di **Simone Quercia** (Professore Ordinario in Economia Politica)

## Principali riferimenti bibliografici

- Acemoglu, D., Autor, D., & Johnson, S. (2026). *Building Pro-Worker Artificial Intelligence*. NBER Working Paper 34854 (2026), <https://doi.org/10.3386/w34854>.
- Agentic AI - The new frontier in GenAI - Pwc report (2024)
- Aldreabi, H., Dahdouh, N. K. S., Alhur, M., et al. (2025). Determinants of Student Adoption of Generative AI in Higher Education. *Electronic Journal of e-Learning*, 23(1), 15–33.
- Arslan, Ahmad, Cary Cooper, Zaheer Khan, Ismail Golgeci e Imran Ali. 2022. "Artificial Intelligence and Human Workers Interaction at Team Level: A Conceptual Assessment of the Challenges and Potential HRM Strategies." *International Journal of Manpower* 43 (1): 75–88. <https://doi.org/10.1108/IJM-01-2021-0052>.
- Bernabei, M. et al. "Adaptive automation: Status of research and future challenges." *Reliability Engineering & System Safety* (2024).
- Berndt, J., Englmaier, F., Sadun, R., Tamayo, J., & von Hesler, N. (2026). *A Systematic Approach to Experimenting with Gen AI*. Harvard Business Review, January–February 2026.
- Berretta, S., Tausch, A. T., Ontrup, G., Gilles, B., & Peifer, C. (2023). *Defining human-AI teaming the human-centered way: a scoping review and network analysis*. *Frontiers in Artificial Intelligence*, 6, 1250725.
- Caldwell, S. et al. "An Agile New Research Framework for Hybrid Human-AI Teaming: Trust, Transparency, and Transferability." *ACM Transactions on Interactive Intelligent Systems* (2022).
- Caplin, A., Deming, D. J., Li, S., Martin, D. J., Marx, P., Weidmann, B., & Ye, K. J. (2024). The ABC's of Who Benefits from Working with AI: Ability, Beliefs, and Calibration. NBER Working Paper No. 33021.
- Castro, S., Englmaier, F., & Guadalupe, M. (2024). Fostering Psychological Safety in Teams: Evidence from an RCT. In *Academy of Management Proceedings* (Vol. 2024, No. 1, p. 16624). Valhalla, NY 10595: Academy of Management.
- Dell'Acqua, F., Ayoubi, C., Lifshitz, H., Sadun, R., Mollick, E., Mollick, L., Han, Y., Goldman, J., Nair, H., Taub, S., and Lakhani, K. (2025). *The Cybernetic Teammate: A Field Experiment on Generative AI Reshaping Teamwork and Expertise*. NBER Working Paper 33641 (2025), <https://doi.org/10.3386/w33641>
- Dell'Acqua, F., Kogut, B., & Perkowski, P. (2023). *Super Mario Meets AI: Experimental Effects of Automation and Skills on Team Performance and Coordination*. *Review of Economics and Statistics*, 107 (4): 951–966.
- Dell'Acqua, F., McFowland III, E., Mollick, E. R., Lifshitz-Assaf, H., Kellogg, K., Rajendran, S., Kraye, L., and Candelon, F., and Lakhani, K. R. (2026). *Navigating the Jagged Technological Frontier: Field Experimental Evidence of the Effects of Artificial Intelligence on Knowledge Worker Productivity and Quality*. In corso di pubblicazione presso Organization Science
- Deloitte, State of AI in the Enterprise (2026 findings on agentic AI use cases).
- Emery, F. E., & Trist, E. L. (1960). Sociotechnical Systems. In C. W. Churchman & M. Verhulst (Eds.), *Management Science, Models and Techniques* (Vol. 2, pp. 83–97). Oxford: Pergamon.
- Feigh, K. M., & Pritchett, A. R. "Requirements for effective function allocation: A critical review." *Journal of Cognitive Engineering and Decision Making* (2014).
- Ferguson, G., & Allen, J. "Mixed-Initiative Systems for Collaborative Problem Solving." *AI Magazine* (2007).
- Hackman, J. R. (2002). *Leading Teams: Setting the Stage for Great Performances*. Boston, MA: Harvard Business School Press.
- Hagemann, V., Rieth, M. R., Suresh, A., & Kirchner, F. (2023). *Human-AI teams — Challenges for a team-centered AI at work*. *Frontiers in Artificial Intelligence*, 6, 1252897.
- Hauptman, Allyson I., Beau G. Schelble, Wen Duan, Christopher Flathmann e Nathan J. McNeese. 2024. "Understanding the Influence of AI Autonomy on AI Explainability Levels in Human-AI Teams Using a Mixed Methods Approach." *Cognition, Technology & Work* 26: 435–455. <https://doi.org/10.1007/s10111-024-00765-7>.
- Horani, O. M., Al-Adwan, A. S., Yaseen, H., et al. (2023). The Critical Determinants Impacting Artificial Intelligence Adoption at the Organisational Level. *Information Development*.
- Horvitz, E. "Principles of mixed-initiative user interfaces." *Proceedings of CHI* (1999).
- Ilgel, D. R., Hollenbeck, J. R., Johnson, M., & Jundt, D. (2005). Teams in organizations: From input-process-output models to IMO models. *Annu. Rev. Psychol.*, 56, 517-543.
- Karountzos, V. (2025). The Impact of AI Adoption on Productivity Across FinTech, Retail, and Advanced Manufacturing: A Systematic Quantitative Literature Review Approach. *Productivity Insights Paper No. 061*, The Productivity Institute.
- Kergroach, S., & H  ritier, J. (2025). Emerging Divides in the Transition to Artificial Intelligence. *OECD Regional Development Papers*, No. 147. OECD Publishing, Paris.
- Kong, X. et al. "Examining human–AI collaboration in hybrid intelligence learning environments: insight from the Synergy Degree Model." *Humanities and Social Sciences Communications* (2025).
- Korinek, Anton. 2025. "AI Agents for Economic Research: August 2025 Update to 'Generative AI for Economic Research: Use Cases and Implications for Economists,'" published in the *Journal of Economic Literature* 61(4).

KPMG, Agentic AI workflows in financial reporting (2026), per i rischi di governance, cascading errors e segregation of duties.

Linux Foundation. (2025–2026). Formation of the Agentic AI Foundation; A2A Protocol one-year milestone and v1.0.

Loaiza, I., & Rigobon, R. (2024). The EPOCH of AI: Human-Machine Complementarities at Work. Available at SSRN 5028371.

Lou, B., Lu, T., Raghu, T. S., & Zhang, Y. (2025). *Unraveling Human–AI Teaming: A Review and Outlook*. <https://arxiv.org/abs/2504.05755>

Lovich, D., Meier, S., & Taylor, C. (2025). *Leaders Assume Employees Are Excited About AI — They're Wrong*. Harvard Business Review. November 2025

Masters, C. et al. "Orchestrating Human-AI Teams: The Manager Agent as a Unifying Research Challenge." arXiv(2025).

Mathieu, J. E., Luciano, M. M., D'Innocenzo, L., Klock, E. A., & LePine, J. A. (2020). The development and construct validity of a team processes survey measure. *Organizational Research Methods*, 23(3), 399-431.

Microsoft, Agents are here—is your company prepared? (WorkLab / 2026).

Microsoft, 2025 Work Trend Index Annual Report: The Year the Frontier Firm Is Born (2025).

Microsoft, materiali online 2026 "Governance and security for AI agents across the organization"

National Academies of Sciences, Engineering, and Medicine. "Human–AI Team Interaction." in *Information Technology and the U.S. Workforce: Where Are We and Where Do We Go from Here?* (capitolo online).

Newman, A., Donohue, R., & Eva, N. (2017). Psychological safety: A systematic review of the literature. *Human resource management review*, 27(3), 521-535.

OECD / Boston Consulting Group / INSEAD (2025). *The Adoption of Artificial Intelligence in Firms: New Evidence for Policymaking*. OECD Publishing, Paris.

OECD, Empowering SMEs in the Age of AI / D4SME evidence package (2026).

Parasuraman, R., Sheridan, T. B., & Wickens, C. D. "A model for types and levels of human interaction with automation." *IEEE Transactions on Systems, Man, and Cybernetics – Part A* (2000).

Plouffe, R. A., Ein, N., Liu, J. J., St. Cyr, K., Baker, C., Nazarov, A., & Don Richardson, J. (2023). Feeling safe at work: Development and validation of the Psychological Safety Inventory. *International Journal of Selection and Assessment*, 31(3), 443-455.

PwC, AI Agent Survey (2025).

Qin, Y., & Sajda, P. (2025). *Teaming with AI: A Multi-Modal Investigation of Human-AI Team Dynamics*. AAAI Conference Paper.

Siemon, D. et al. "Elaborating Team Roles for Artificial Intelligence-based Teammates in Human-AI Collaboration." *Group Decision and Negotiation* (2022).

Song, B. "Human-AI collaboration by design." (Cambridge Core PDF, 2024).

Stanford Institute for Human-Centered Artificial Intelligence (2024). *AI Index Report 2024*. Stanford University.

Tan, Q. (2025). Factors Influencing the Adoption of Generative AI in Education: A Systematic Review, Proposed Framework and Future Research Agenda. *British Educational Research Journal* (online first).

Wageman, R., Hackman, J. R., & Lehman, E. (2005). Team Diagnostic Survey: Development of an Instrument. *The Journal of Applied Behavioral Science*, 41(4), 373–398. <https://doi.org/10.1177/0021886305281984>

Wang, Z. Z., Shao, Y., Shaikh, O., Fried, D., Neubig, G., & Yang, D. (2025). How Do AI Agents Do Human Work? Comparing AI and Human Workflows Across Diverse Occupations.

Wu, S., Yang, D., Chang, J., Hearst, M. A., & Lo, K. (2025). *Human-AI Collaboration: How AIs Augment Human Teammates*. ACL Tutorial 2025.

Zhang, Guanglu, Leah Chong, Kenneth Kotovsky e Jonathan Cagan. 2023. "Trust in an AI versus a Human Teammate: The Effects of Teammate Identity and Performance on Human-AI Cooperation." *Computers in Human Behavior* 139: 107536.

—

Belbin - Finding Your Authentic Contribution in the Age of AI <https://www.belbin.com/resources/articles-directory/finding-your-authentic-contribution-in-the-age-of-ai>

Government Digital Service (GDS). *Microsoft 365 Copilot Experiment: Cross-Government Findings Report*. <https://www.gov.uk/government/publications/microsoft-365-copilot-experiment-cross-government-findings-report/microsoft-365-copilot-experiment-cross-government-findings-report-html>

Harvard Business Review. *Systematic Approach to Experimenting with Gen AI* <https://hbr.org/2026/01/a-systematic-approach-to-experimenting-with-gen-ai>

UK Department for Science, Innovation and Technology (DSIT). *Landmark government trial shows AI could save civil servants nearly 2 weeks a year* (Press release, 2 giugno 2025) <https://www.gov.uk/government/news/landmark-government-trial-shows-ai-could-save-civil-servants-nearly-2-weeks-a-year>