



**UNIVERSITÀ DEGLI STUDI DI MILANO**  
**SCUOLA DI SCIENZE DELLA MEDIAZIONE**  
**LINGUISTICA E CULTURALE**

**Corso di Laurea Magistrale in Lingue e Culture per la  
comunicazione e cooperazione internazionale**

**Classe LM-38**

**Costruire senso nell'era dell'IA: comunicazione, leadership e  
co-intelligenza nel progetto Lead-Alliance**

Relatore: Chiar.mo Prof. Di Iorio Luigi

Correlatore: Dott. Manuel Dileo

Tesi di Laurea di: Micol Maria Vittoria Bianchi

Matr. 40996A

Anno Accademico 2024-2025



### *LE TRE LEGGI DELLA ROBOTICA*

1. *Un robot non può recar danno a un essere umano, né permettere che, a causa della propria negligenza, un essere umano patisca danno.*
2. *Un robot deve sempre obbedire agli ordini degli esseri umani, a meno che contrastino con la Prima Legge.*
3. *Un robot deve proteggere la propria esistenza, purché questo non contrasti con la Prima o Seconda Legge.*

Isaac Asimov. *Io, Robot* (1950)



## INDICE

|                                                                                                        |           |
|--------------------------------------------------------------------------------------------------------|-----------|
| <b>INTRODUZIONE.....</b>                                                                               | <b>1</b>  |
| <b>CAPITOLO 1 - COSTRUIRE SENSO NELL'ERA DELL'IA.....</b>                                              | <b>6</b>  |
| 1.1 La leadership come processo comunicativo e di costruzione del senso.....                           | 7         |
| 1.2 Comunicazione uomo-macchina: simulazione emotiva, percezione e fiducia...                          | 10        |
| 1.3 L'intelligenza artificiale come fattore di trasformazione dei processi di senso...                 | 13        |
| 1.4 Co-intelligenza e collaborazione interdisciplinare: una sfida comunicativa e<br>organizzativa..... | 18        |
| 1.5 Dalla teoria alla pratica: Lead-Alliance come laboratorio di co-intelligenza.....                  | 20        |
| <b>CAPITOLO 2: LA SECONDA FASE DI LEAD-ALLIANCE.....</b>                                               | <b>22</b> |
| 2.1 Il passaggio da "hub" a "base operativa".....                                                      | 23        |
| 2.2 Identità di progetto e obiettivi emergenti.....                                                    | 29        |
| 2.3 Struttura organizzativa e governance interna: il ruolo del senato ristretto.....                   | 31        |
| 2.4 La struttura dei team come architettura operativa.....                                             | 33        |
| <b>CAPITOLO 3: COMUNICARE LEAD-ALLIANCE.....</b>                                                       | <b>35</b> |
| 3.1 Comunicazione interna: coordinamento, fiducia e costruzione di pratiche<br>condivise.....          | 37        |
| 3.2 Comunicazione esterna e brand identity: posizionamento, TOV e riconoscibilità.<br>45               |           |
| 3.3 Content strategy e community building: sperimentazioni e output.....                               | 48        |
| 3.3.1 Il sito come infrastruttura narrativa e spazio di legittimazione.....                            | 49        |
| 3.3.2 Il podcast come dispositivo divulgativo e matrice simbolica.....                                 | 54        |
| 3.3.3 La newsletter come protocollo sperimentale e dispositivo metodologico..                          | 57        |
| 3.3.4 Dal whitepaper al libro: formalizzazione e capitalizzazione del sapere....                       | 58        |
| 3.4 Prospettive di sviluppo.....                                                                       | 60        |
| 3.4.1 Tra consolidamento e rischio di dispersione.....                                                 | 61        |
| 3.4.2 Formalizzazione e stabilità: l'ipotesi di integrazione strutturale.....                          | 66        |
| <b>CAPITOLO 4: IL FUTURO DELLE IA NELLE ORGANIZZAZIONI.....</b>                                        | <b>70</b> |
| 4.1 L'efficacia dell'IA nei contesti organizzativi: criteri, limiti e ambiguità.....                   | 72        |
| 4.1.1 I limiti strutturali dell'affidabilità algoritmica: bias, overfitting e drift.....               | 73        |
| 4.1.2 I livelli di efficacia.....                                                                      | 75        |
| 4.1.3 Hard skills e soft skills: una simmetria solo apparente.....                                     | 77        |
| 4.1.4 I rischi metodologici della misurazione in contesti AI-driven.....                               | 78        |
| 4.2 Slash soft-skill assessment: a case study.....                                                     | 80        |
| 4.2.1 Genesi del progetto.....                                                                         | 81        |
| 4.2.2 Architettura del questionario.....                                                               | 82        |
| 4.2.3 Il problema della misurazione delle soft skills nel caso Slash.....                              | 86        |
| 4.2.4 Slash tra governance e co-creazione.....                                                         | 88        |
| <b>CONCLUSIONI.....</b>                                                                                | <b>90</b> |

|                                                                            |            |
|----------------------------------------------------------------------------|------------|
| <b>BIBLIOGRAFIA.....</b>                                                   | <b>94</b>  |
| <b>SITOGRAFIA.....</b>                                                     | <b>107</b> |
| <b>APPENDICE A: Presentazione della seconda fase di Lead-Alliance.....</b> | <b>108</b> |
| <b>ALLEGATO A: Brochure altermAInd Slash.....</b>                          | <b>113</b> |

## INTRODUZIONE

L'integrazione delle tecnologie legate all'intelligenza artificiale (IA) nei contesti organizzativi sta progressivamente spostando la riflessione sulla *leadership* verso un terreno sempre più culturale e strategico. Le questioni in gioco non riguardano più soltanto le capacità computazionali degli algoritmi, ma investono il significato dell'autorità, della responsabilità e della fiducia in contesti nei quali decisioni umane e suggerimenti automatizzati convivono e, talvolta, si sovrappongono. In questo scenario, la *leadership* è chiamata non solo a governare processi e risultati, ma a sostenere la costruzione di senso all'interno di sistemi complessi e tecnologicamente mediati.

In tale quadro, studiare come si formano comunità attorno a progetti che cercano di tenere insieme intelligenza artificiale e *leadership* organizzativa diventa centrale per comprendere alcune traiettorie di cambiamento delle organizzazioni contemporanee. Il presente elaborato si propone pertanto di analizzare la seconda fase del progetto *Lead-Alliance*, un'iniziativa che si colloca a metà tra *hub* multidisciplinare, piattaforma di sperimentazione, rete professionale e manifesto culturale dedicato all'integrazione dell'IA nelle pratiche di *leadership*. *Lead-Alliance* è il risultato di un processo di creazione collettivo nel quale professionisti provenienti da ambiti differenti, tra cui formazione, *high-tech*, consulenza e ingegneria, hanno condiviso l'obiettivo di tradurre orientamenti teorici in azioni concrete, con particolare attenzione alle dimensioni etiche e algoretiche<sup>1</sup>.

Se la prima fase del progetto<sup>2</sup> ha posto le fondamenta concettuali e tematiche attraverso incontri preliminari, *workshop* e prime sperimentazioni, la seconda fase, avviata formalmente a settembre 2025, si concentra sulla traduzione operativa di tali orientamenti. In questa fase il progetto ha lavorato alla costruzione di una presenza pubblica più riconoscibile e strutturata. Tale sviluppo è passato attraverso scelte comunicative ed editoriali precise, tra cui la realizzazione di un sito *web*, l'avvio di una comunicazione continuativa su *LinkedIn*, la redazione di contenuti in forma di *blog* e la

---

<sup>1</sup> Il termine algoretica designa la riflessione etica applicata specificamente agli algoritmi e ai sistemi di intelligenza artificiale, con particolare attenzione alle implicazioni morali dei processi decisionali automatizzati. Per un approfondimento, si veda Benanti (2022).

<sup>2</sup> La prima fase del progetto *Lead-Alliance* è stata oggetto di analisi nella tesi di laurea di Alice De Gennaro, discussa nel dicembre 2025 presso questo stesso dipartimento, dal titolo *Lead-Alliance: un gruppo interdisciplinare per esplorare la leadership nell'era dell'intelligenza artificiale*. Il presente elaborato si pone in continuità con tale lavoro, concentrandosi sulla seconda fase del progetto.

progettazione di ulteriori *output*, quali *newsletter*, prodotti audiovisivi e libri. Sullo sfondo resta inoltre l'ipotesi, ancora in evoluzione, di percorsi formativi digitali orientati sia alla diffusione delle riflessioni emerse sia alla costruzione di una comunità interessata ai temi proposti.

L'approccio adottato in questa tesi è di tipo partecipato e qualitativo. In veste di responsabile della comunicazione interna ed esterna del progetto, ho osservato e contribuito alle decisioni strategiche, prendendo parte alle riunioni di coordinamento e alla definizione dei contenuti. Questa posizione ha offerto il vantaggio di un accesso diretto ai processi di negoziazione di senso che hanno accompagnato lo sviluppo della seconda fase, ma ha anche richiesto un attento rigore metodologico per contenere il rischio di parzialità. Per questa ragione, la ricerca combina l'osservazione partecipante con diverse tecniche di analisi qualitativa con l'obiettivo di mantenere un equilibrio tra esperienza soggettiva e sistematicità nell'interpretazione dei dati.

In coerenza con l'impostazione metodologica adottata e con l'oggetto stesso della ricerca, l'intelligenza artificiale è stata inoltre anche utilizzata come strumento di supporto nella revisione stilistica del testo, nella ricerca bibliografica e nella produzione di alcuni grafici esplicativi presenti nei capitoli. Tale impiego non ha inciso sulle scelte teoriche, sull'impianto argomentativo o sull'interpretazione dei dati, ma si è configurato come pratica di co-intelligenza: un'interazione supervisionata e criticamente orientata, coerente con l'ipotesi di fondo secondo cui l'IA può affiancare e potenziare i processi umani di costruzione di senso senza sostituirli.<sup>3</sup>

Accanto al mio ruolo all'interno del progetto e alle metodologie implementate nel presente lavoro, risulta poi rilevante descrivere brevemente la struttura del nucleo operativo che ha guidato questa fase di sviluppo. Il cosiddetto "senato ristretto", composto da sette persone oltre alla sottoscritta, costituisce il principale centro

---

<sup>3</sup> In linea con quanto dichiarato nel *Decalogo per un utilizzo etico, legittimo e consapevole di strumenti d'Intelligenza Artificiale (AI) in tutte le attività dell'Ateneo* (2024), e in particolare nel documento *1. Linee guida per un uso vantaggioso dell'AI in ambito didattico* (consultabile al link [https://www.unimi.it/sites/default/files/regolamenti/Allegato%201\\_decalogo%20AI.pdf](https://www.unimi.it/sites/default/files/regolamenti/Allegato%201_decalogo%20AI.pdf)), si dichiara che, durante la preparazione di questo lavoro, sono stati utilizzati come strumenti IA quali *Chat-GPT* (5.2) e *Claude Sonnet* (4.6) per: redigere contenuti, parafrasare e riformulare, migliorare lo stile di scrittura e controllare grammatica e ortografia. Inoltre, in alcuni casi è stato utilizzato il tool IA *Perplexity* per la ricerca di fonti bibliografiche accademiche. Infine, si è utilizzato *Canva AI* per la creazione di grafici e diagrammi inseriti nei capitoli. Dopo aver utilizzato questi strumenti, il contenuto è stato attentamente revisionato e modificato secondo necessità. L'autrice si assume dunque la piena responsabilità del proprio elaborato.

decisionale del progetto e riunisce figure con competenze e *background* differenti. La composizione di questo gruppo aiuta a comprendere sia la rapidità con cui le idee sono state elaborate e trasformate in proposte operative, sia le difficoltà legate alla costruzione di una comunicazione coerente e attrattiva. Nei momenti di confronto iniziali è emersa infatti una forte sinergia interdisciplinare, affiancata dalla necessità di tradurre linguaggi specialistici, talvolta molto distanti, in pratiche comunicative accessibili a un pubblico più ampio. È all'interno di queste dinamiche, e dalla mia posizione nel *team*, che si costituiscono le basi empiriche su cui la presente ricerca si fonda.

L'obiettivo del lavoro è duplice. Da un lato, documentare e analizzare in modo critico le scelte metodologiche e comunicative che hanno caratterizzato la seconda fase di *Lead-Alliance*; dall'altro, collocare tale esperienza all'interno della letteratura esistente sulla *leadership* contemporanea, sul *knowledge management* e sull'etica dell'intelligenza artificiale, al fine di trarne indicazioni operative trasferibili ad altri contesti organizzativi. Le domande di ricerca riguardano, in primo luogo, il modo in cui la comunicazione può sostenere coesione, fiducia e pratiche condivise all'interno di un gruppo interdisciplinare impegnato su un tema ad alto impatto come l'integrazione tra IA e *leadership*; in secondo luogo, le strategie narrative e comunicative più efficaci nella costruzione di un'identità di progetto e nella creazione di un pubblico; infine, le modalità attraverso cui le principali questioni etiche e algoretiche connesse all'uso dell'intelligenza artificiale vengono discusse e tradotte in pratiche operative.

Il primo capitolo dell'elaborato propone quindi una cornice teorico-interpretativa che non si configura come una rassegna sistematica della letteratura, ma come una scelta epistemologica coerente con l'oggetto di studio. Costruire senso nell'era dell'intelligenza artificiale implica infatti adottare uno sguardo capace di tenere insieme *leadership*, comunicazione e collaborazione interdisciplinare come fenomeni dinamici. Tale impostazione risponde alla natura stessa dell'IA nei contesti organizzativi, intesa non come strumento neutro, ma come attore in grado di riorganizzare pratiche, linguaggi e processi decisionali.

Il secondo capitolo è dedicato all'analisi del progetto *Lead-Alliance* e ne ricostruisce identità, obiettivi e struttura organizzativa, con particolare attenzione allo sviluppo della seconda fase. Il terzo capitolo approfondisce invece le strategie di comunicazione

interna ed esterna adottate nel corso di tale fase, analizzando i linguaggi utilizzati, le sperimentazioni avviate e i principali *output*, tra cui il sito *web* e i contenuti editoriali.

È qui necessario esplicitare alcuni limiti della presente ricerca. La seconda fase del progetto, iniziata formalmente a settembre 2025, è tuttora in corso e ciò implica che alcune valutazioni potranno essere riviste alla luce di sviluppi successivi. La mia partecipazione diretta al progetto può inoltre introdurre elementi di soggettività nell'analisi, sebbene mitigati dall'adozione di strumenti metodologici coerenti con un approccio qualitativo. Per ragioni di riservatezza, infine, non vengono riportate comunicazioni private o dati sensibili senza previo consenso; nomi e ruoli sono citati esclusivamente nella misura necessaria alla comprensione complessiva del progetto.

Ciononostante, il valore atteso del presente lavoro risiede nella possibilità di contribuire al dibattito accademico sul rapporto tra intelligenza artificiale e *leadership*, offrendo da un lato un caso empirico approfondito e dall'altro indicazioni operative utili a organizzazioni interessate a intraprendere percorsi analoghi, quali pratiche di comunicazione, modelli di *governance* interna e strumenti editoriali. In un contesto in cui l'adozione dell'intelligenza artificiale generativa non rappresenta più un'opzione<sup>4</sup>, ma una sfida culturale e organizzativa attuale, comprendere come costruire fiducia, responsabilità e senso condiviso anche in presenza di tecnologie pervasive diventa una leva centrale per la sostenibilità dei sistemi organizzativi.

L'elaborato si conclude con un capitolo autonomo dedicato a una riflessione metodologica esplorativa sulla valutazione dell'efficacia dell'utilizzo di modelli di intelligenza artificiale nelle organizzazioni. Attraverso il *case study* del *soft skills assessment* di *Slash*, progettato da *Altermaind*, e senza la pretesa di fornire metriche definitive, tale sezione problematizza criteri, limiti e rischi della misurazione, interrogandosi su cosa significhi parlare di efficacia quando sono in gioco dimensioni immateriali quali il senso, la qualità delle decisioni e le dinamiche collaborative. In questo modo, la tesi contribuisce ad aprire una possibile traiettoria di ricerca e di pratica sul rapporto tra tecnologia, *leadership* e costruzione di senso. Questo capitolo, pur non essendo direttamente collegato al caso *Lead-Alliance*, ne rappresenta un complemento

---

<sup>4</sup> Questa lettura è confermata dai dati di adozione rilevati a livello globale: secondo Gartner (2024), la *Generative AI* è oggi la soluzione di intelligenza artificiale più frequentemente implementata nelle organizzazioni.

analitico: mentre *Lead-Alliance* mostra come si costruisce senso attorno all'IA nei processi di *leadership*, il caso di *Slash* offre uno sguardo sulle sfide metodologiche della misurazione e valutazione di tali tecnologie in contesti applicativi. La scelta di includerlo come sezione autonoma risponde alla volontà di ampliare la riflessione oltre il singolo caso empirico, aprendo una traiettoria di ricerca sui criteri di efficacia quando sono in gioco dimensioni qualitative e relazionali.

## CAPITOLO 1 - COSTRUIRE SENSO NELL'ERA DELL'IA

L'importanza della costruzione di un senso condiviso in un'epoca segnata dalla crescente pervasività dell'intelligenza artificiale si configura come un obiettivo centrale per le pratiche di *leadership* e per il funzionamento delle organizzazioni. Prima di entrare nel merito del progetto *Lead-Alliance*, risulta pertanto necessario esplicitare il quadro concettuale adottato in questa ricerca, che nasce dalla volontà di evitare un'analisi meramente descrittiva o tecnologica del fenomeno, privilegiando invece una prospettiva attenta alle sue implicazioni organizzative, relazionali e comunicative.

Nel dibattito contemporaneo, *leadership* e intelligenza artificiale vengono spesso trattate come ambiti distinti, ricondotti a un punto di incontro prevalentemente strumentale, in cui le tecnologie sono chiamate a supportare o rendere più efficienti i processi decisionali dei *leader*. In questo capitolo, al contrario, l'attenzione viene spostata dal piano dell'efficienza a quello del senso. L'intelligenza artificiale non è considerata unicamente come uno strumento operativo, ma come un elemento che interviene nei processi di costruzione del significato, nel coordinamento dell'azione collettiva e nei meccanismi di legittimazione delle decisioni.

In questa prospettiva, la *leadership* viene assunta non come una funzione individuale o una posizione formale, ma bensì come un processo collettivo e situato di costruzione del senso, che si sviluppa nel tempo attraverso pratiche comunicative, dinamiche relazionali e forme di co-creazione. L'introduzione dell'intelligenza artificiale non sostituisce tale processo, ma ne modifica profondamente le condizioni di possibilità, rendendo il lavoro interpretativo al tempo stesso più complesso e più centrale. Le tecnologie digitali, infatti, amplificano la necessità di integrare linguaggi eterogenei, differenze culturali e nuovi modelli di interazione, ponendo al centro temi quali la fiducia, l'inclusione e l'innovazione.

Alla luce di questa impostazione, non si intende offrire una rassegna sistematica dei modelli di *leadership* né una classificazione delle tecnologie di intelligenza artificiale. L'interesse è rivolto piuttosto ai concetti teorici utili a comprendere come la *leadership* si configuri oggi come un processo emergente, relazionale e comunicativo, e come l'IA contribuisca a ridefinire le pratiche di *governance* e di costruzione del senso nei contesti

organizzativi contemporanei. Si tratta dunque di una cornice selettiva e orientata, che non ambisce all'esaustività, ma alla coerenza analitica.

Il presente capitolo si articola pertanto attorno alla lettura della *leadership* come pratica di costruzione del senso condiviso, ponendo particolare attenzione al ruolo della fiducia e della mediazione tra caratteristiche umane e tecnologiche. Tale prospettiva consente di preparare il terreno per l'analisi del caso *Lead-Alliance*, che verrà affrontato non come semplice esempio applicativo, ma come spazio empirico entro cui osservare in modo concreto le dinamiche teoriche qui delineate.

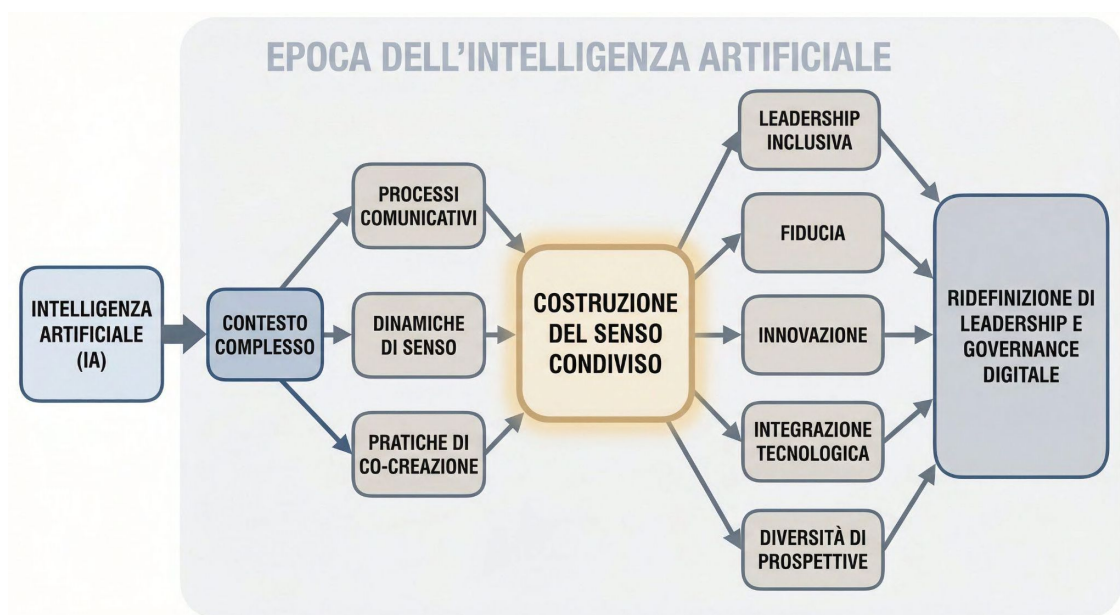


Fig. 1: Modello concettuale per la leadership nell'epoca dell'Intelligenza Artificiale.<sup>5</sup>

## 1.1 La leadership come processo comunicativo e di costruzione del senso

Per spostare il *focus* dalle definizioni astratte della *leadership* alle dinamiche organizzative concrete, è necessario analizzare la *leadership* come l'evoluzione di un processo che prende forma all'interno delle pratiche comunicative e relazionali. In questa prospettiva, essa non precede l'azione del gruppo, ma si costruisce nel corso dell'azione stessa, attraverso tentativi successivi di interpretazione, dialogo e

<sup>5</sup> Il presente diagramma, riassuntivo dell'argomentazione del capitolo, mostra come l'intelligenza artificiale crea un contesto complesso che richiede la costruzione (o ri-costruzione) di senso condiviso. Questo si ottiene tramite la creazione di processi comunicativi, dinamiche di senso e pratiche di co-creazione, portando a risultati come la *leadership* inclusiva e culminando infine nella ridefinizione della *leadership* e della *governance* digitale.

coordinamento. Proprio per questo motivo, nelle organizzazioni contemporanee, la *leadership* tende sempre più a configurarsi come un processo di comunicazione orientato alla costruzione condivisa del senso, che supera completamente l'immagine del *leader* isolato, autocratico e autoreferenziale.

Un riferimento teorico centrale per questa lettura è rappresentato dalla prospettiva del *sensemaking* elaborata da Karl Weick. In *Sensemaking in Organizations*, l'autore sostiene infatti che le organizzazioni non possano essere comprese come strutture oggettive e stabili, ma piuttosto come processi sociali fondati su pratiche continue di attribuzione di significato (Weick, 1995, pp. 1–9). In contesti caratterizzati da ambiguità e incertezza, gli attori organizzativi sono costantemente impegnati nel rendere intelligibili i fatti attraverso narrazioni plausibili che orientano l'azione collettiva. Il senso, in questa prospettiva, non è quindi dato a priori, ma emerge dall'esperienza e viene progressivamente stabilizzato attraverso la comunicazione.

Applicata alla *leadership*, questa visione implica che "guidare" non significhi fornire risposte definitive o imporre una direzione univoca a un gruppo, bensì contribuire alla costruzione di cornici interpretative condivise che rendano possibile l'azione del gruppo stesso. Come evidenziato da Weick, il *sensemaking* è infatti un processo sociale, continuo, orientato alla plausibilità più che all'accuratezza, e fondato sulla negoziazione di significati tra prospettive differenti (Weick, 1995, pp. 17–62). In questa chiave, la *leadership* si esercita proprio all'interno dei processi di *sensemaking*, sostenendo la continuità dell'azione collettiva e rendendo la complessità organizzativa attraversabile.

Questa dimensione processuale emerge con particolare evidenza nei contesti complessi e interdisciplinari, nei quali la pluralità di linguaggi, *background* e competenze rende centrale la dimensione dialogica. In tali scenari, le decisioni richiedono spesso tempi più lunghi, perché la costruzione del consenso e l'emersione di eventuali tensioni implicite non possono essere risolte con un atto unilaterale ma attraverso una mediazione. Tuttavia, è importante sottolineare che proprio questo stesso investimento di tempo consente al contempo di valorizzare la diversità delle prospettive e di integrare punti di vista eterogenei all'interno del sistema organizzativo.

In questo senso, la *leadership* come processo comunicativo e di costruzione del senso può essere intesa come la capacità di integrare visioni differenti e di trasformare le differenze attraverso il dialogo, la negoziazione e l'utilizzo produttivo del conflitto. Ne

deriva un modello di *governance* meno verticale e più collaborativo, in cui la co-intelligenza del gruppo diventa la risorsa principale per tradurre idee e orientamenti in azioni sostenibili. Attivare un processo dialogico di questo tipo, inoltre, spesso incrementa il senso di appartenenza e rafforza la motivazione intrinseca, poiché rende le persone partecipi della costruzione degli obiettivi e delle pratiche.

La comunicazione assume, in questo quadro, un ruolo cruciale come fattore di costruzione dei significati condivisi, degli obiettivi e delle pratiche operative. Ciò risulta particolarmente rilevante proprio nei *team* interdisciplinari, nei quali la sola circolazione delle informazioni non è sufficiente: occorre invece progettare momenti e luoghi di confronto strutturato, in cui il gruppo possa chiarire i concetti, negoziare il lessico comune e ridurre le ambiguità che inevitabilmente emergono quando si integrano linguaggi e competenze differenti. Il compito della *leadership*, in tale direzione, è quello di alimentare questi spazi e sostenerne la continuità nel tempo, evitando che strategie comunicative eccessivamente tecnicistiche o, al contrario, superficiali finiscano per indebolire la coerenza del lavoro condiviso.

Un ulteriore elemento centrale nei processi di *leadership* riguarda poi la costruzione della fiducia, sia nelle relazioni tra i membri del gruppo sia nei confronti delle tecnologie. La fiducia rappresenta una variabile strutturale delle relazioni organizzative, capace di ridurre la complessità e di rendere più fluidi i processi decisionali (Covey, 2023, pp. 25–30). In contesti caratterizzati da fiducia, la comunicazione tende a essere più aperta e orientata alla cooperazione; al contrario, in assenza di fiducia, il senso si frammenta e l'azione collettiva si irrigidisce.

La centralità delle relazioni emerge con particolare evidenza anche se si considera il ruolo delle emozioni nei processi di coordinamento collettivo. Daniel Goleman ha mostrato come le competenze emotive incidano in modo significativo sulla qualità della *leadership*, influenzando il clima relazionale e la capacità dei gruppi di affrontare situazioni complesse (Goleman, 2011, pp. 242–245). In *Lavorare con intelligenza emotiva*, l'autore evidenzia proprio come il *leader*, attraverso il suo stile comunicativo, contribuisca a modellare l'ambiente emotivo dell'organizzazione, incidendo direttamente sui livelli di motivazione e coinvolgimento del gruppo (Goleman, 2000, pp. 220–227). Ne deriva che la *leadership* non riguarda soltanto la definizione degli obiettivi, ma

soprattutto la cura delle condizioni relazionali che rendono possibile un lavoro condiviso sul significato dell'azione.

Le stesse definizioni più note del concetto di *leadership* riconoscono, infatti, seppur con accenti diversi, questa dimensione processuale. Per esempio, Northouse definisce la *leadership* come "il processo attraverso il quale un individuo influenza un gruppo di persone per raggiungere un obiettivo comune" (Northouse, 2019, p. 6)<sup>6</sup>. In questa affermazione, ciò che risulta particolarmente rilevante non è tanto l'enfasi sull'influenza del singolo, quanto il riferimento allo svilupparsi di un processo e alla dimensione relazionale dello stesso. L'influenza si costruisce nel tempo attraverso interazioni, scambi comunicativi e forme di riconoscimento reciproco e la *leadership*, in questa prospettiva, non può quindi mai essere separata dal contesto né dalle pratiche interpretative che caratterizzano la vita organizzativa.

In definitiva, la concezione della *leadership* come pratica comunicativa e interpretativa assume una rilevanza particolare nell'era dell'intelligenza artificiale. L'introduzione di sistemi capaci di generare, organizzare e suggerire informazioni modifica le condizioni entro cui il *sensemaking* organizzativo ha luogo: le fonti si moltiplicano, le decisioni diventano più mediate, e le interpretazioni umane si intrecciano con *output* algoritmici non sempre immediatamente comprensibili. In tale scenario, il lavoro di costruzione del senso non viene ridotto, ma tende piuttosto a intensificarsi. La *leadership*, quindi, non consiste nel delegare il giudizio alla tecnologia, ma nel sostenere processi interpretativi consapevoli che tengano insieme competenze umane e supporti tecnici, preservando la responsabilità, la partecipazione e la coerenza del senso.

## **1.2 Comunicazione uomo–macchina: simulazione emotiva, percezione e fiducia**

La comunicazione uomo–macchina è diventata un nodo nevralgico nella riflessione sull'interazione tra esseri umani e sistemi di intelligenza artificiale, soprattutto nei contesti organizzativi interdisciplinari, in cui linguaggi, ruoli e competenze eterogenei devono essere coordinati in tempi rapidi e in condizioni di elevata ambiguità. Accanto allo sviluppo di sistemi sempre più performanti, l'attenzione si sta progressivamente

---

<sup>6</sup> Tale definizione viene qui assunta come riferimento descrittivo di sintesi, utile a evidenziare la dimensione processuale e relazionale della *leadership*, più che come modello teorico prescrittivo.

spostando dalla sola accuratezza tecnica alla qualità relazionale dell'interazione, intesa come capacità del sistema di inserirsi in pratiche comunicative già esistenti senza interromperne il flusso interpretativo.

Negli ultimi anni, tale tendenza è diventata visibile anche al di fuori dei contesti accademici, attraverso pratiche di sperimentazione che mettono in scena l'IA in contesti conversazionali informali e variegati. Per esempio<sup>7</sup>, esperimenti veicolati tramite contenuti audio di tipo *podcast* generati da strumenti di intelligenza artificiale generativa come *Notebook LM*<sup>8</sup> oppure l'interazione con la modalità vocale di vari modelli tra cui *ChatGPT*<sup>9</sup> mostrano come le macchine provino a riprodurre elementi tipici della comunicazione umana: ritmo dialogico, pause, enfasi, persino una forma di "intonazione emotiva". Questo scarto tra simulazione emotiva e percezione umana è centrale nella costruzione della fiducia verso i sistemi di intelligenza artificiale. Da un lato, la capacità di imitare tratti comunicativi percepiti come empatici tende infatti a rendere la macchina più accessibile, aumentando la disponibilità delle persone a interagire con essa. Dall'altro, permane la consapevolezza che tale empatia non è autentica: l'IA non possiede intenzionalità, responsabilità morale né "comprensione" nel senso umano del termine. Ne consegue che la fiducia non dipende esclusivamente dalla qualità comunicativa dell'interfaccia, ma dalla possibilità di attribuire affidabilità, trasparenza e controllabilità al sistema e al contesto organizzativo che lo utilizza.

A livello manageriale, le analisi di settore segnalano poi come l'intelligenza artificiale, e in particolare la *Generative AI*, sia oggi tra le soluzioni più frequentemente implementate nelle organizzazioni (Gartner, 2024). Tale diffusione viene spesso collegata a miglioramenti in termini di produttività, efficienza operativa e rapidità nell'elaborazione delle informazioni. Questo dato non rappresenta soltanto un indicatore tecnologico, ma prevalentemente un segnale culturale: l'IA sta infatti progressivamente diventando parte integrante dei processi decisionali ordinari.

Tuttavia, è anche da sottolineare che l'incremento di produttività non coincide automaticamente con un rafforzamento della qualità interpretativa. Accelerare la

---

<sup>7</sup> Gli esempi qui richiamati non hanno valore di evidenza empirica scientifica, ma svolgono una funzione di osservazione culturale: rende visibile le modalità attraverso cui l'intelligenza artificiale viene progressivamente normalizzata nelle pratiche comunicative quotidiane e i criteri impliciti con cui il pubblico ne valuta l'accettabilità e la credibilità.

<sup>8</sup> Cfr. <https://notebooklm.google.com/>

<sup>9</sup> Cfr. <https://chatgpt.com/it-IT/features/voice/>

produzione di *output* non significa infatti necessariamente approfondire il processo di costruzione del senso. È proprio in questo scarto tra efficienza e significato che la questione della fiducia e della *governance* diventa centrale. In tale contesto, la letteratura sul *trust in automation* e sul *decision-making* supportato da sistemi intelligenti mostra come la fiducia verso sistemi automatizzati costituisca una variabile critica: può abilitare l'efficacia decisionale, ma può anche generare forme di dipendenza o di deresponsabilizzazione se non adeguatamente bilanciata (Lee & See, 2004; Qwaider et al., 2024).

Questo equilibrio si intreccia anche con la letteratura organizzativa sulla fiducia nei *team* e sulla *psychological safety*, intesa come la percezione condivisa che il contesto sia sufficientemente sicuro da consentire l'espressione di dubbi, errori e disaccordi senza timore di ripercussioni (Edmondson, 2012). Nei contesti mediati dall'IA, tale dimensione assume una rilevanza ulteriore: se gli *output* algoritmici vengono percepiti come autorevoli o "oggettivi", diventa essenziale che i membri del gruppo si sentano legittimati a metterli in discussione, interrogarli e persino contraddirli. In assenza di sicurezza psicologica, il rischio non è soltanto l'errore tecnico, ma la normalizzazione acritica della raccomandazione automatizzata.

La qualità della decisione dipende dunque non solo dalla *performance* del sistema, ma dalla possibilità di mantenere aperto uno spazio dialogico in cui il giudizio umano resti attivo e responsabile. Letti in questa chiave, i *report* di settore, quali i già citati di Gartner, non indicano semplicemente una crescita dell'adozione tecnologica, ma rendono evidente la necessità di una *governance* consapevole dell'interazione uomo-macchina, capace di integrare velocità algoritmica e responsabilità interpretativa.

Nei contesti di lavoro collaborativi e interdisciplinari, in questo senso, la fiducia assume dunque ancor più una dimensione relazionale. Essa contribuisce a superare barriere linguistiche, professionali e culturali, sostenendo il senso di appartenenza e il *commitment* dei membri. Quando l'IA entra in questi ambienti, diventa a tutti gli effetti un "terzo attore" nei processi di coordinamento, incidendo non solo sul rapporto individuo-macchina, ma anche sulle dinamiche comunicative tra le persone che operano attraverso la tecnologia.

In aggiunta, un ulteriore elemento di complessità è introdotto, come già esemplificato, dalla crescente diffusione di interfacce conversazionali progettate per

apparire sempre più dialogiche e parzialmente "umanizzate". Se da un lato questa caratteristica facilita l'ingaggio e l'accettazione degli strumenti di IA, dall'altro introduce ambiguità nei processi di attribuzione di responsabilità e legittimità delle decisioni. La simulazione di tratti comunicativi umani può rendere infatti l'interazione indubbiamente più fluida, ma non elimina affatto la necessità di un presidio critico sull'uso della tecnologia.

Proprio per questo motivo, gli approcci orientati alla *human-centered AI* insistono sull'importanza della spiegabilità, della trasparenza e dell'*accountability* dei sistemi intelligenti. In questa prospettiva, la fiducia non deriva solamente dall'illusione di una comunicazione "umana" da parte della macchina, ma dalla possibilità di comprenderne i limiti, interrogare i risultati prodotti e mantenere una chiara attribuzione di responsabilità. *Leadership* e *governance* assumono così un ruolo centrale nel definire le condizioni entro cui l'IA può essere integrata come risorsa organizzativa, senza compromettere l'autonomia decisionale e il senso critico degli attori coinvolti.

In contesti interdisciplinari come *Lead-Alliance*, la fiducia nei sistemi di intelligenza artificiale si configura conseguentemente come un fenomeno organizzativo e culturale, più che puramente tecnologico. Essa incide sulle modalità di coordinamento, sulla costruzione del senso e sulla legittimazione delle decisioni, rendendo necessarie pratiche intenzionali di comunicazione, revisione e controllo umano. È su questo piano che l'interazione uomo-macchina mostra il suo potenziale trasformativo, incidendo non soltanto sull'efficienza operativa, ma sulle condizioni stesse entro cui il senso viene costruito nelle organizzazioni.

### **1.3 L'intelligenza artificiale come fattore di trasformazione dei processi di senso**

In definitiva, l'IA può essere letta come un fattore che trasforma la produzione di senso nelle organizzazioni, spostando l'attenzione dall'automazione verso la necessità di costruire processi interpretativi comuni. La capacità di elaborare grandi quantità di dati e generare risposte, spesso, più che adeguate consente certamente di ridefinire il processo decisionale; tuttavia, il limite strutturale dell'IA nella comprensione di dimensioni analogiche, empatiche e creative interroga criticamente la qualità e il senso delle decisioni prese in contesti in cui esseri umani interagiscono con sistemi artificiali. Si crea così una dicotomia: da un lato l'efficacia e la velocità dell'algoritmo, dall'altro la

necessità di valutazioni qualitative e riflessive, tipiche del giudizio umano, soprattutto quando sono in gioco scelte ad alto valore etico o creativo (Cruciani, 2023, p. 2). In altre parole, l'IA accelera la reazione, ma non garantisce automaticamente la riflessione: ed è proprio qui che le decisioni rischiano di diventare certamente "più rapide" ma anche molto "più povere di senso".

Luciano Floridi aiuta a inquadrare questa trasformazione proponendo di leggere le tecnologie digitali come elementi che ridefiniscono l'infosfera, ossia l'ambiente informazionale entro cui gli esseri umani agiscono e prendono decisioni (Floridi, 2013, p. 51). In questa prospettiva, l'intelligenza artificiale non è un mezzo neutrale: essa partecipa alla produzione, selezione e circolazione delle informazioni, incidendo su ciò che viene reso visibile, rilevante o degno di attenzione. Ne consegue un impatto diretto sulle condizioni stesse del *sensemaking* organizzativo, perché cambia la materia prima da cui gli attori costruiscono interpretazioni e coordinano l'azione.

Se infatti, come sostiene Weick, il senso emerge retrospettivamente dall'azione e viene stabilizzato attraverso narrazioni plausibili (Weick, 1995, pp. 23–25), l'introduzione dell'IA modifica le fonti stesse da cui tali narrazioni attingono. Gli *output* algoritmici – spesso percepiti come più "oggettivi" proprio perché "macchinosi" – entrano nel repertorio interpretativo dell'organizzazione e orientano aspettative, valutazioni e decisioni. Il rischio principale non risiede solo nell'errore tecnico, ma nella tendenza a naturalizzare l'*output*, attribuendogli un'autorità che non sempre è giustificata.

Paolo Benanti affronta questo nodo attraverso il concetto di fiducia algoritmica, mettendo in evidenza come la crescente delega cognitiva ai sistemi di intelligenza artificiale possa generare forme di deresponsabilizzazione (Benanti, 2022, pp. 95–110). Quando l'algoritmo viene percepito come un'istanza neutrale o superiore, il processo interpretativo umano tende a ridursi, lasciando spazio a un'accettazione acritica delle soluzioni proposte. In termini organizzativi, si tratta di uno spostamento del baricentro decisionale che non è soltanto tecnico, ma culturale e comunicativo: cambia chi viene ascoltato, cosa conta come prova, e come si legittima una scelta.

Questo fenomeno può essere letto anche alla luce della riflessione weberiana sulla legittimità. Max Weber descrive infatti la legittimità razionale-legale come fondata su procedure formalizzate e prevedibili (Weber, 1978, pp. 215–220). Nei contesti mediati

dall'IA, questa forma di legittimità tende a intrecciarsi con una legittimazione "procedurale" basata sulla presunta imparzialità del calcolo e sulla formalizzazione algoritmica, che però – a differenza delle procedure burocratiche tradizionali – opera spesso tramite meccanismi opachi e difficilmente contestabili<sup>10</sup>. Floridi descrive questo tipo di influenza come potere infraetico: una forma di autorità che agisce non tanto attraverso norme esplicite, quanto attraverso la configurazione dell'ambiente informazionale e delle possibilità di azione disponibili (Floridi, 2014, pp. 125–140). L'IA, in quanto infrastruttura spesso invisibile, orienta comportamenti e decisioni non sostituendosi direttamente all'agire umano, ma modellando il contesto stesso entro cui il senso viene costruito.

Accanto a fiducia e legittimità, un ulteriore elemento critico riguarda poi i *bias*<sup>11</sup> che emergono nell'interazione tra esseri umani e sistemi di intelligenza artificiale. Da un lato, l'IA può incorporare distorsioni presenti nei dati di addestramento e nelle scelte progettuali: ciò influenza gli *output* e, di conseguenza, le interpretazioni che diventano disponibili e "plausibili" per l'organizzazione. Dall'altro, l'interazione con l'IA attiva *bias* cognitivi negli utenti, che incidono sul modo in cui gli *output* vengono accolti e utilizzati. Fenomeni come l'effetto *ELIZA*<sup>12</sup> – la tendenza ad attribuire alla macchina capacità di comprensione ed empatia superiori a quelle effettive – sono noti fin dalle origini dei *chatbot* (Weizenbaum, 1966; Weizenbaum, 1976). Allo stesso modo, l'*automation bias* può spingere le persone a considerare le indicazioni automatizzate come intrinsecamente più affidabili del giudizio umano, favorendo forme di *over-reliance* (Lee & See, 2004). In queste condizioni, lo spazio del giudizio critico tende a ridursi e il processo interpretativo umano viene indebolito: il rischio non è quindi solo la decisione sbagliata, ma la progressiva perdita di responsabilità

---

<sup>10</sup> In gergo tecnico, per quanto riguarda gli LLM (*Large Language Models*), talvolta si parla di tale opacità con il termine *black box*: nella costruzione di alcune IA, infatti, gli intricati sistemi del loro funzionamento e del loro processo decisionale (spesso simili a reti neurali profonde) sono incomprensibili persino agli sviluppatori, rendendo molto difficile capire come arrivino a certe conclusioni, nonostante forniscano risultati considerabili accurati (Palo Alto Networks, n.d.).

<sup>11</sup> Per un ulteriore approfondimento sul tema dei *bias*, si veda il quarto capitolo.

<sup>12</sup> Il nome deriva dal primo *chatbot* mai inventato, sviluppato da Joseph Weizenbaum nel 1966. La principale funzionalità di *ELIZA* era quella di riuscire a simulare efficacemente i *pattern* delle conversazioni umane tramite il riconoscimento di schemi ripetuti (*pattern matching*). La sua applicazione più nota, chiamata *DOCTOR*, riusciva persino a emulare uno psicoterapeuta, creando l'illusione di comprensione e dando origine al suddetto effetto psicologico, ovvero la tendenza umana ad attribuire personalità alle macchine.

interpretativa, cioè l'abitudine a considerare il senso come "dato" anziché costruito e negoziato.

Per questo, l'introduzione dell'IA impone una rielaborazione delle pratiche di interazione e una maggiore cura dei processi di *sensemaking*: non tanto un *sensemaking* "fatto dalla macchina", ma un *sensemaking* costruito con la macchina e attorno ai suoi *output*, attraverso pratiche collaborative e spesso interdisciplinari. La domanda che ne consegue è però anche concreta: come si traduce questo presidio umano nei sistemi che già oggi mediano decisioni organizzative?

Una risposta significativa viene dagli approcci *human-in-the-loop* (HITL), che descrivono i sistemi in cui l'intervento umano non è solo un gesto simbolico, ma una componente strutturale del *workflow*: la persona supervisiona, corregge, valida o interrompe l'*output* della macchina, preservando *accountability* e possibilità di contestazione (IBM, s.d.).

Una declinazione oggi centrale di questa logica è conseguentemente il *Reinforcement Learning from Human Feedback* (RLHF), ossia un insieme di tecniche in cui un modello apprende preferenze e criteri di qualità a partire da valutazioni umane. In *Deep reinforcement learning from human preferences*, Christiano et al. mostrano come obiettivi complessi possano essere trasmessi al sistema tramite le preferenze espresse da persone su alternative di comportamento, riducendo così la necessità di definire esplicitamente una funzione-ricompensa "perfetta" (Christiano et al., 2017). Allo stesso modo, in ambito linguistico, Ouyang et al. descrivono un processo in cui, dopo una fase supervisionata, il modello viene ulteriormente ottimizzato utilizzando *ranking* di *output* forniti da valutatori umani, con l'obiettivo di riprodurre le risposte preferite dagli utenti in termini di utilità, sicurezza e aderenza all'intento (Ouyang et al., 2022).

Portare il RLHF dentro una riflessione organizzativa consente di esplicitare un punto chiave: la qualità dell'*output* non dipende solo dalla potenza del modello, ma da chi definisce i criteri di "buona risposta", quali dati costituiscono il *feedback*, e quali valori vengono premiati o penalizzati. Ciò non elimina totalmente la possibilità di distorsioni, ma rende visibile che anche l'allineamento è il risultato di scelte umane situate, e dunque storicamente e culturalmente determinate. In altri termini, l'allineamento non è un fatto puramente tecnico: è un processo socio-organizzativo in cui si formalizzano priorità, linguaggi e norme di comportamento che incidono direttamente sulla

legittimazione della decisione e sulla fiducia nel sistema. In questa prospettiva, la *leadership* si colloca non a valle dell'algoritmo, ma a monte dei criteri che ne orientano il comportamento.

Da questo ragionamento deriva quindi la necessità di un bilanciamento costante tra automazione e mantenimento del senso e dei valori organizzativi, e ciò implica anche processi periodici di revisione per identificare ed eliminare errori sistemici e distorsioni ricorrenti, nonché la stipulazione di politiche di *governance* orientate a garantire coerenza decisionale e salvaguardia etica (Cruciani, 2023, p. 2).

Questa fragilità non è però solo teorica, in quanto la tensione tra efficienza algoritmica e responsabilità interpretativa è già ad oggi una realtà vissuta nei contesti organizzativi, e non ancora pienamente risolta. Sul piano empirico, emerge infatti una tensione tipica dei contesti organizzativi contemporanei: se, infatti, molti *manager* riportano benefici in termini di qualità del processo decisionale, permane tuttavia una preoccupazione legata all'opacità e alle implicazioni etiche degli algoritmi. A tal proposito, Qwaider et al. (2024, p. 4) segnalano che il 65% dei *manager* che utilizzano IA indica un netto miglioramento nella presa di decisioni, mentre il 50% esprime preoccupazioni proprio rispetto a possibili derive etiche di queste stesse decisioni. Questo dato chiarisce che la fiducia nel sistema artificiale rimane fragile e non può essere risolta con la sola *performance* tecnica: si richiede comunicazione, trasparenza, capacità di contestualizzazione e presidio critico.

In tale scenario, la *leadership* come pratica di costruzione del senso non può mai limitarsi a "integrare" l'IA nei processi decisionali: deve rendere visibili le distorsioni potenziali, sostenere pratiche riflessive e creare tutte le condizioni affinché gli *output* vengano discussi, tradotti e inseriti in narrazioni condivise. Gli *output* dell'IA non parlano da soli: richiedono sempre spiegazione e mediazione. In assenza di questo lavoro, il rischio è che la tecnologia venga percepita alternativamente come oracolo infallibile o come minaccia, generando resistenze, fraintendimenti o dipendenza cognitiva. È proprio nella comunicazione che tutte queste tensioni possono essere nominate e negoziate, permettendo al gruppo di mantenere un ruolo attivo nella costruzione del senso.

A tal proposito, Ethan Mollick sottolinea come l'IA possa anche funzionare da amplificatore delle capacità umane, ma solo a condizione che vi sia una consapevolezza

chiara dei suoi limiti e una postura critica nel suo uso quotidiano (Mollick, 2025, pp. 45–65). In definitiva, l'IA trasforma il processo decisionale e intensifica la complessità interpretativa e proprio per questo richiede un ripensamento delle pratiche di *leadership* e di *governance*, affinché automazione e responsabilità restino in equilibrio e il senso dell'azione non venga delegato, ma costruito in modo collaborativo.

#### **1.4 Co-intelligenza e collaborazione interdisciplinare: una sfida comunicativa e organizzativa**

L'introduzione dell'intelligenza artificiale nei contesti organizzativi non si limita quindi a ridefinire processi decisionali o modalità operative, ma incide direttamente e in modo significativo sulle forme di collaborazione e sulle dinamiche di coordinamento tra attori con competenze, linguaggi e prospettive differenti. In questo scenario, il concetto di co-intelligenza consente di spostare l'attenzione dall'idea di sostituzione dell'agire umano a quella di interazione tra capacità umane e sistemi artificiali, mettendo in luce tutte le condizioni organizzative e comunicative che rendono tale interazione possibile.

Ethan Mollick utilizza a tal proposito il termine *cointelligence* per descrivere proprio queste forme emergenti di collaborazione uomo–macchina, in cui l'IA agisce come amplificatore delle capacità cognitive e creative, più che come loro sostituto (Mollick, 2025, pp. 1–15). Tuttavia, questa amplificazione non è mai né automatica né garantita: dipende totalmente, come già analizzato, dalla capacità degli attori organizzativi di comprendere il funzionamento dei sistemi, riconoscerne i limiti e integrarli in pratiche di lavoro già esistenti. In altri termini, la co-intelligenza non è una proprietà della tecnologia, ma una qualità che emerge dall'interazione, dal contesto e dalle *routine* comunicative che la sostengono.

Questa dinamica diventa particolarmente complessa nei contesti multidisciplinari, in cui l'IA viene discussa, progettata e utilizzata da soggetti portatori di saperi eterogenei. Professionisti della tecnologia, della formazione, della consulenza, della comunicazione e dell'ingegneria tendono infatti a interpretare l'intelligenza artificiale attraverso categorie spesso molto differenti, attribuendole funzioni, potenzialità e rischi non sempre coincidenti. La sfida principale non riguarda quindi soltanto l'integrazione tecnica dei sistemi, ma la costruzione di un linguaggio comune che consenta collaborazione e coordinamento dell'azione.

Come già riportato, la letteratura sottolinea poi che *team* caratterizzati da *background* differenti producono soluzioni spesso più innovative e ad ampio raggio rispetto a gruppi omogenei, con esiti positivi sia in termini di *problem-solving* sia di costruzione di senso (Bala, 2024, p. 3; Stephens et al., 2020, p. 20). Le diverse discipline favoriscono infatti un approccio complementare che spinge ad affrontare questioni controverse da prospettive non convenzionali e a integrare competenze specialistiche. In questa cornice, la co-intelligenza può essere intesa anche come la sommatoria di diversi strumenti e linguaggi. Tuttavia, solo una gestione efficace di tali elementi può abilitare la creatività del *team*, mentre una gestione carente può di contro favorire collisioni o isolamento.

In tale prospettiva, la collaborazione interdisciplinare non diventa un effetto collaterale dell'adozione dell'intelligenza artificiale, ma una sua condizione di possibilità. È attraverso il confronto tra prospettive diverse che l'IA può essere interrogata criticamente, evitando tanto l'entusiasmo acritico quanto il rifiuto difensivo. L'interazione tra umano e artificiale, infatti, si colloca in uno spazio di tensione tra logiche differenti: l'IA opera per correlazioni probabilistiche e statistiche, mentre gli esseri umani attribuiscono significato attraverso esperienze, valori e la loro cultura di appartenenza. Riconoscere questa stessa tensione, anziché tentare di eliminarla, è una condizione fondamentale per un uso consapevole e responsabile dell'IA nelle organizzazioni.

Proprio per questo la co-intelligenza appare come un problema eminentemente comunicativo. Gli *output* dell'IA devono essere interpretati, contestualizzati e discussi all'interno del gruppo. In questo quadro, la *leadership* assume perciò una funzione cruciale nel sostenere la co-intelligenza come pratica organizzativa. Se la *leadership* è intesa come processo di costruzione del senso, essa si manifesta qui nella capacità di facilitare il dialogo tra competenze diverse, rendere esplicite le assunzioni implicite e mantenere aperto lo spazio dell'interpretazione critica.

A tal proposito, in *Exponential Organizations*, Ismail osserva che le organizzazioni capaci di integrare tecnologie avanzate sviluppano strutture spesso meno gerarchiche e più distribuite, in cui il processo decisionale è quindi il risultato di interazioni multiple più che di comandi centralizzati (Ismail et al., 2015, pp. 47–53). Tuttavia, questa distribuzione dell'intelligenza e dell'iniziativa comporta anche un aumento della

complessità comunicativa: diventa infatti necessario un lavoro costante di allineamento tra obiettivi, valori e interpretazioni, affinché la collaborazione non si frammenti.

In definitiva, la co-intelligenza uomo-uomo e uomo-macchina non elimina affatto il conflitto interpretativo, ma può renderlo produttivo, trasformandolo in una risorsa per l'apprendimento collettivo. In questo quadro, la *leadership* si configura come pratica di accompagnamento e cura dei processi collaborativi, chiamata a sostenere l'emergere di significati condivisi in ambienti segnati dalla presenza crescente dell'intelligenza artificiale.

### **1.5 Dalla teoria alla pratica: Lead-Alliance come laboratorio di co-intelligenza**

Il percorso teorico sviluppato in questo capitolo ha mostrato come la *leadership* possa essere letta come un processo di costruzione del senso, radicato nelle pratiche comunicative e relazionali che attraversano le organizzazioni. In questa prospettiva, l'intelligenza artificiale non si configura come un semplice supporto tecnico, ma come un elemento che trasforma le condizioni entro cui tali processi interpretativi hanno luogo, incidendo sulle modalità di collaborazione, coordinamento e attribuzione di significato all'azione collettiva.

La nozione di co-intelligenza consente di superare una visione oppositiva tra umano e artificiale, mettendo in evidenza come il valore dell'IA emerga dall'interazione con competenze umane, contesti organizzativi e pratiche comunicative condivise. La co-intelligenza non è quindi una proprietà intrinseca della tecnologia, ma una qualità emergente che richiede tempo, negoziazione e un investimento continuo sul piano comunicativo.

In questa prospettiva, gli approcci *human-in-the-loop* e, in particolare, il RLHF possono essere letti come la traduzione tecnica di una tesi organizzativa più ampia: la *leadership* non "subisce" l'IA, ma può contribuire a definirne criteri di funzionamento, soglie di accettabilità e stili comunicativi, attraverso pratiche strutturate di valutazione, revisione e *feedback*. Non si tratta di trasformare tutti i *manager* in ingegneri, ma di rendere visibile che l'allineamento dell'IA richiede attori competenti che sappiano formulare criteri, riconoscere distorsioni e governare i *trade-off* tra efficienza, qualità relazionale e responsabilità.

È proprio su questo sfondo teorico che il progetto *Lead-Alliance* può essere quindi letto non come un caso esemplare in senso astratto, ma come uno spazio empirico in cui osservare concretamente le dinamiche di costruzione del senso nell'interazione tra persone e sistemi di intelligenza artificiale.

## CAPITOLO 2: LA SECONDA FASE DI LEAD-ALLIANCE

Il capitolo precedente ha proposto una cornice teorico-interpretativa fondata sull'idea che *leadership* e comunicazione, nell'era dell'intelligenza artificiale, non possano essere comprese come dimensioni separate: la *leadership* si manifesta sempre più come capacità di orientare processi interpretativi collettivi, mentre l'IA interviene riorganizzando le condizioni entro cui tali processi prendono forma. In questa prospettiva, studiare *Lead-Alliance* significa osservare un contesto concreto in cui la costruzione di senso non è un tema astratto, ma una pratica quotidiana di mediazione tra linguaggi, competenze ed aspettative eterogenee, interne ed esterne al progetto.

A partire da tale impostazione, il presente capitolo ricostruisce la seconda fase di *Lead-Alliance* come momento di svolta: si tratta del passaggio da una configurazione prevalentemente esplorativa e sperimentale a una struttura più definita, operativa e orientata alla produzione di *output* pubblici. Questo passaggio, formalizzato nella riunione del 15 settembre 2025 dell'*Action Committee* (senato ristretto), segna una fase cruciale di consolidamento, in cui il gruppo ha riconosciuto di disporre di una base di *community* sufficientemente matura e ha scelto di rafforzare il proprio posizionamento attraverso una maggiore strutturazione organizzativa e comunicativa.

La transizione a una forma più strutturata ha avuto effetti immediati sulle dinamiche interne: le riunioni si sono progressivamente spostate da uno spazio di confronto prevalentemente discorsivo a un coordinamento più operativo; gli esperimenti con strumenti di intelligenza artificiale hanno lasciato il posto a gruppi di lavoro orientati alla produzione di contenuti e materiali strutturati; è aumentata la necessità di assumere decisioni rapide e coerenti, così come l'esposizione esterna del progetto. In questo senso, la definizione più chiara di ruoli, l'introduzione di pratiche di *governance* partecipativa e l'elaborazione di modelli comunicativi condivisi non rappresentano soltanto un'esigenza organizzativa, ma diventano parte integrante del processo di costruzione di senso e di posizionamento di *Lead-Alliance*.

Nel delineare questa fase, il capitolo adotta un punto di vista interno e situato, coerente con l'impostazione partecipata della ricerca. La mia posizione di responsabile della comunicazione interna ed esterna mi colloca infatti in un ruolo non solo osservativo, ma anche attivo nell'orientamento delle scelte editoriali e nel

coordinamento della coerenza comunicativa del progetto. Tale collocazione consente di restituire con maggiore precisione la natura "in divenire" di *Lead-Alliance*, richiedendo al contempo un'attenzione costante nel distinguere tra ricostruzione analitica dei processi e valutazioni personali.

I paragrafi che seguono approfondiscono tre snodi principali: il passaggio da *hub* di discussione a una maggiore operatività e le conseguenti implicazioni organizzative; l'emergere progressivo di un'identità di progetto più definita e di obiettivi condivisi; infine, la struttura organizzativa e la *governance* interna, con particolare attenzione al ruolo del senato ristretto come centro decisionale formale e come spazio di coordinamento tra i diversi gruppi di lavoro.

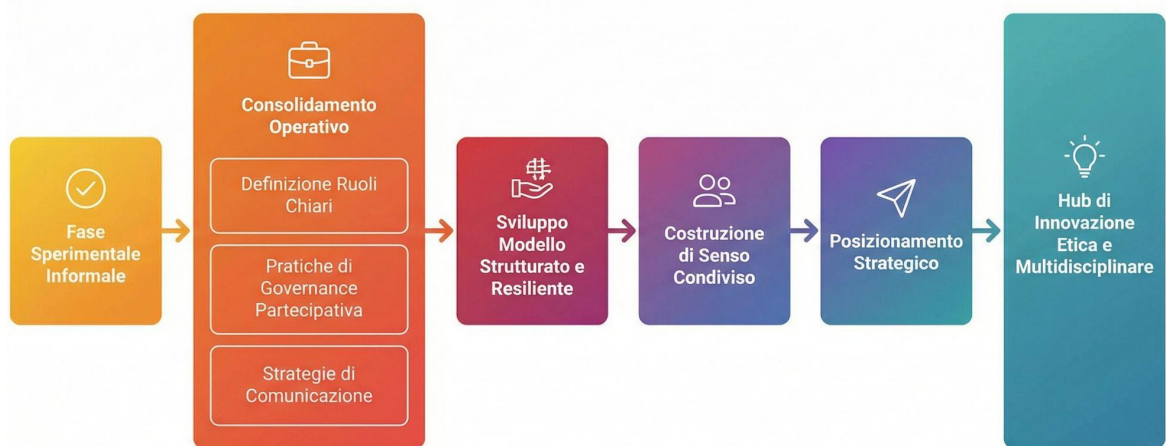


Fig. 2: Evoluzione di *Lead-Alliance*, dalla fase sperimentale all'hub di innovazione.

## 2.1 Il passaggio da "hub" a "base operativa"

*Lead-Alliance* nasce come iniziativa interdisciplinare<sup>13</sup> e volontaria orientata a esplorare il rapporto tra intelligenza artificiale e *leadership*, in un contesto segnato dalla crescente diffusione di sistemi di IA generativa e dal conseguente bisogno di una riflessione critica sulle loro implicazioni organizzative, etiche e culturali. Nel corso del 2025 il progetto si è progressivamente configurato come spazio di confronto tra professionisti provenienti da ambiti differenti, con l'obiettivo di superare una visione

<sup>13</sup> La genesi del progetto, già narrata e analizzata da Alice De Gennaro nella sua tesi, è partita da un'iniziativa personale della Dott.ssa De Gennaro affiancata dal Prof. Luigi Di Iorio.

puramente strumentale dell'IA e di interrogare il suo impatto sui modelli di *leadership* e sui processi decisionali. In questa fase fondativa, l'attenzione si è concentrata soprattutto sulla costruzione di una cornice teorica condivisa e sulla sperimentazione di linguaggi comuni, ponendo le basi per uno sviluppo successivo più strutturato.

Un primo momento di discontinuità rispetto alla fase iniziale può essere individuato nell'incontro del 19 giugno 2025, in cui il progetto, fino ad allora conosciuto come *Captain AI & The Leadership Avengers*, assume ufficialmente il nome di *Lead-Alliance*. Questo passaggio non ha rappresentato una semplice operazione di *rebranding*, ma ha segnato l'emergere di una consapevolezza condivisa: la fase esplorativa non costituiva il punto di arrivo del percorso, bensì il presupposto per la costruzione di un progetto collettivo più stabile. Le riflessioni emerse in quella sede, arricchite dal contributo di esperti esterni quali Luca Ioime<sup>14</sup> e Mattia Pedroncelli<sup>15</sup>, hanno contribuito a chiarire la necessità di dare continuità al lavoro avviato, orientandolo verso forme più stabili di collaborazione e di produzione di conoscenza. Nei mesi successivi, questa consapevolezza si è poi tradotta nei primi tentativi di organizzare il lavoro in modo più stabile, attraverso una *roadmap* preliminare e la costituzione di un nucleo operativo ristretto incaricato di accompagnare l'evoluzione del progetto e coordinare le attività emergenti. Parallelamente, si è delineata una distinzione progressiva tra un cerchio di partecipazione maggiormente coinvolto nella progettazione attiva e un gruppo più ampio di partecipanti interessati a contribuire in modo discontinuo, attraverso *feedback* e momenti di confronto. Questa configurazione, ancora fluida, ha reso visibile l'esigenza di passare da un orizzonte prevalentemente esplorativo a una struttura in grado di sostenere nel tempo la crescita del progetto.

Riassumendo, nella sua fase iniziale, *Lead-Alliance* si è dunque configurato come uno spazio prevalentemente sperimentale, orientato alla condivisione di prospettive teoriche, esperienze professionali e prime sperimentazioni, in assenza di una struttura operativa rigidamente definita. La dimensione discorsiva e la libertà di esplorazione hanno rappresentato una risorsa fondamentale sia per la costruzione di una visione comune, sia per il consolidamento di una prima *community* interessata ai temi proposti. Tuttavia, con il progressivo ampliamento del gruppo e con l'emergere di un nucleo

---

<sup>14</sup> Head of Sales di Boosha AI. (cfr. <https://www.linkedin.com/company/boosha-ai/posts/>)

<sup>15</sup> CEO di XR Solution. (cfr. <https://www.linkedin.com/in/mattia-pedroncelli/>)

stabile di partecipanti, è divenuto evidente come tale impostazione, pur feconda sul piano ideativo, non fosse più sufficiente a garantire l'efficacia del progetto nel medio e lungo periodo.

Il punto di svolta definitivo può essere individuato nella riunione del 15 settembre 2025 del nucleo operativo, denominato da quel momento all'interno del gruppo "senato ristretto", durante la quale i membri maggiormente coinvolti hanno riconosciuto di disporre ormai di una base di relazioni e di interesse tale da rendere necessario e ufficiale un passaggio di fase. In quella sede matura così la decisione di trasformare *Lead-Alliance* da *hub* di riflessione a vera e propria "base operativa", capace di produrre *output* concreti e di rivolgersi a un pubblico più ampio in modo strutturato. Il passaggio dalla fase di esplorazione alla fase operativa rappresenta un momento fondamentale perché segna il passaggio dall'ideazione alla concretezza: così si è resa necessaria la definizione di un nucleo decisionale ristretto e la chiarificazione dei ruoli, con l'obiettivo di rendere più rapido e coerente il processo decisionale senza rinunciare al carattere partecipativo e transdisciplinare del progetto.

In ogni caso, il passaggio alla seconda fase non implica il totale abbandono della dimensione esplorativa, ma piuttosto una sua riformulazione in chiave operativa. A partire da quel momento, le riunioni hanno sicuramente assunto un carattere maggiormente orientato al coordinamento e alla presa di decisioni, riducendo progressivamente lo spazio dedicato alla sola discussione teorica, tuttavia gli esperimenti iniziali con strumenti di intelligenza artificiale non vengono mai abbandonati. Essi sono stati infatti ricondotti dentro gruppi di lavoro tematici, ciascuno incaricato della progettazione e realizzazione di *output* pubblici — contenuti editoriali, iniziative di divulgazione, materiali di approfondimento — e chiamato a rispettare scadenze, vincoli di coerenza e criteri condivisi di qualità. In questa fase emergono così con forza alcune tensioni organizzative tipiche dei gruppi interdisciplinari: la difficoltà di far coesistere coerenza e pragmatismo strategico, l'attrito tra approcci metodologici differenti e la necessità di mediare linguaggi e priorità. Queste tensioni non sono di facile gestione, ma possono diventare una leva generativa, in grado di alimentare soluzioni e visioni realmente interdisciplinari.

Parallelamente alla ridefinizione degli obiettivi, il cambiamento di fase si è tradotto poi nell'adozione di strumenti e pratiche tipiche della progettazione organizzativa,

funzionali a rendere sostenibile il passaggio dall'esplorazione all'operatività. Nei diversi sottogruppi di lavoro sono stati quindi introdotti strumenti di pianificazione e coordinamento finalizzati a rendere visibile l'avanzamento delle attività, a chiarire le responsabilità e a favorire un allineamento operativo tra i membri coinvolti. Tra questi, un ruolo centrale è stato svolto dal diagramma di Gantt<sup>16</sup>, utilizzato come strumento di rappresentazione temporale delle attività. In linea con l'impostazione del *Project Management Body of Knowledge*, il Gantt ha consentito di esplicitare la sequenza delle azioni, le interdipendenze tra compiti e le scadenze, contribuendo a spostare il lavoro del gruppo da una logica prevalentemente discorsiva a una logica orientata all'esecuzione e al rispetto dei tempi (Project Management Institute, 2021). In questo senso, lo strumento non ha avuto una funzione meramente tecnica, ma ha agito come dispositivo cognitivo e comunicativo, facilitando la convergenza tra visione strategica e azione concreta.

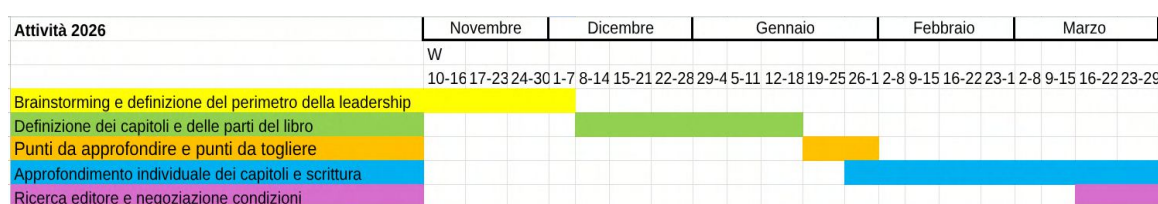


Fig. 2.1.1: Pianificazione Gantt del team Whitepaper.

Accanto agli strumenti di pianificazione temporale, sono state adottate anche tecniche di ideazione strutturata finalizzate allo sviluppo degli *output*. In particolare, durante le fasi di progettazione dei contenuti e delle iniziative, è stato utilizzato in più occasioni il metodo del *lotus blossom*, una tecnica di *brainstorming* che consente di esplorare un tema centrale attraverso una progressiva articolazione in sotto-temi e idee correlate. A differenza di forme di *brainstorming* più libere, il *lotus blossom* favorisce infatti una riflessione sistematica e visuale, permettendo di mantenere una coerenza costante tra la visione generale del progetto e le singole proposte operative. L'adozione

<sup>16</sup> Il diagramma di Gantt, ideato da Henry Gantt nel 1910, costituisce un grafico a barre orizzontali finalizzato alla gestione visiva della pianificazione temporale di un progetto. L'asse verticale presenta l'elenco delle attività, mentre l'asse orizzontale rappresenta la linea temporale. Esso risulta strumentale per la visualizzazione della durata, delle eventuali sovrapposizioni e delle dipendenze tra i compiti, elementi che agevolano il coordinamento del *team* e l'osservanza delle scadenze prefissate.

di questo strumento risulta coerente con l'impostazione multidisciplinare di *Lead-Alliance*, poiché consente di far dialogare approcci e linguaggi differenti senza disperdere il *focus* progettuale, sostenendo al contempo un processo decisionale condiviso.

|                                       |                                                                               |                                                                                  |                                                                         |                                                 |                                       |                              |                            |                                     |
|---------------------------------------|-------------------------------------------------------------------------------|----------------------------------------------------------------------------------|-------------------------------------------------------------------------|-------------------------------------------------|---------------------------------------|------------------------------|----------------------------|-------------------------------------|
| Cina? / punto di vista internazionale | preparazione delle prossime generazioni (aspetto pedagogico) + paletti legali | 5 generazioni                                                                    |                                                                         |                                                 |                                       | puntate generazioni          |                            |                                     |
|                                       | temi                                                                          | interpretazione AI + leadership (non nascondersi/chi ne sa di più) - generazioni | simulazione AI bambino oggi leader di domani (cosa ha imparato con AI?) | modalità                                        | differenziazione pubblico/piattaforma | interviste a esperti         | puntate/stagioni           |                                     |
|                                       | prompt per i leader                                                           | aspetti legali                                                                   | linguaggio spaziale                                                     | shorts (con AI) delle puntate lunghe + commento | AI anchor man (far parlare l'AI)      |                              |                            |                                     |
| text to speech                        | notebook lm                                                                   | audacity                                                                         | temi                                                                    | modalità                                        | puntate/stagioni                      | it director                  | manager                    | alla AI stessa                      |
| copilot                               | programmi utili                                                               |                                                                                  | programmi utili                                                         | PODCAST                                         | interviste                            | formatore (azienda Milena)   | interviste                 | esperto in pedagogia dello sviluppo |
|                                       |                                                                               |                                                                                  | piattaforme                                                             | target                                          | tempistiche                           | invitati degli eventi        | rappresentanti generazioni | persone internazionali              |
|                                       | comunicazione su linkedin                                                     | spotify o youtube?                                                               | allinearsi con altri team (one voice)                                   | target-piattaforma                              |                                       | entro gennaio 26 macro/micro | GANT                       |                                     |
|                                       | piattaforme                                                                   | instagram                                                                        |                                                                         | target                                          |                                       |                              | tempistiche                |                                     |

Fig. 2.1.2: Lotus Blossom nel brainstorming del team Podcast.

L'organizzazione del lavoro in sottogruppi coordinati, ciascuno responsabile di specifici *output*, richiama inoltre modelli già sperimentati in altri contesti di progettazione collaborativa. Un riferimento significativo, emerso anche nel confronto interno al gruppo, è quello delle *junior enterprise*, organizzazioni diffuse in ambito universitario che operano secondo una struttura a *team* con ruoli definiti e responsabilità operative, orientate alla realizzazione di progetti concreti per *stakeholder* esterni. In particolare, il modello di *JECOMM*<sup>17</sup>, la *junior enterprise* dell'Università degli Studi di Milano, rappresenta un esempio di organizzazione in cui la produzione di *output*, il coordinamento tra gruppi di lavoro e la presenza di figure di riferimento consentono di coniugare apprendimento, sperimentazione e operatività. Pur collocandosi in un contesto differente e non esclusivamente universitario, la seconda fase di *Lead-Alliance* mostra alcune analogie strutturali con questo modello: la suddivisione in *team*, la presenza di responsabili per ciascun ambito di lavoro e l'orientamento alla produzione di

<sup>17</sup> Cfr. <https://www.jecomm.it/>

contenuti pubblici costituiscono infatti elementi chiave del passaggio all'operatività. Tale transizione non riduce la dimensione riflessiva del progetto, ma la integra in pratiche organizzative più strutturate, capaci di sostenere nel tempo la costruzione di senso e la traduzione delle idee in azioni.

Un altro elemento rilevante della transizione è poi il riconoscimento della necessità di dotarsi di un'identità organizzativa più definita. Il lavoro di formalizzazione di *mission*, *vision* e valori, analogamente a quanto avviene nelle organizzazioni strutturate, risponde non solo a esigenze di comunicazione esterna, ma anche a un bisogno interno di allineamento: orienta le priorità, rende più leggibili i criteri decisionali e riduce l'ambiguità operativa. In questo quadro, la costruzione del senso non resta confinata alla dimensione discorsiva, ma diventa una vera e propria pratica organizzativa continua che accompagna l'azione.

Questa dinamica di negoziazione del senso si riflette altresì in modo evidente anche nella gestione delle differenze interne. La pluralità di metodologie, culture professionali e aspettative può generare frizioni e dissonanze, soprattutto quando aumenta l'urgenza di decidere e produrre. Proprio per questo, nella fase operativa la comunicazione assume una funzione strutturante: non è solo diffusione all'esterno, ma infrastruttura interna che sostiene coesione, chiarezza e responsabilità reciproca. Le strategie narrative condivise e le pratiche di facilitazione contribuiscono infatti a trasformare la diversità in risorsa, favorendo risonanze tra prospettive differenti e mitigando le dissonanze (Bruno et al., 2008). In questo senso, la *leadership* emerge ancor di più come capacità di costruire senso in modo corale, più che come mera direzione esecutiva.

Parallelamente alla strutturazione interna, la fase operativa richiede conseguentemente un consolidamento delle scelte editoriali e comunicative. La definizione di una strategia di posizionamento, la costruzione di un'identità riconoscibile e l'avvio di una produzione multiformato permettono al progetto di rendersi leggibile e credibile all'esterno. La produzione di contenuti in formati diversi — coerenti tra loro e riconducibili a un impianto valoriale condiviso — aumenta infatti la capacità di coinvolgimento e sostiene lo sviluppo della *community*. In particolare, in contesti formativi e innovativi, questa capacità di tenere insieme identità, apprendimento e traduzione operativa rappresenta un elemento distintivo per sostenere transdisciplinarietà e adattabilità nel tempo (Grivet et al., 2025).

Nel suo complesso, riassumendo, la capacità del progetto di reggere la tensione tra le sue due forze motrici di "*hub*" e "base operativa" — ovvero l'idea di mantenere vivo il laboratorio di sperimentazione, ma rendendolo anche sostenibile e riconoscibile — costituisce uno degli elementi che qualificano *Lead-Alliance* come esperienza in cui la costruzione di senso e l'operatività non si escludono, ma si alimentano costantemente e reciprocamente.

## **2.2 Identità di progetto e obiettivi emergenti**

Il passaggio alla seconda fase ha reso quindi evidente la necessità di chiarire e rendere condivisa l'identità di *Lead-Alliance* non soltanto come cornice simbolica, ma come dispositivo operativo capace di orientare le scelte del gruppo e di sostenerne la coerenza nel tempo. In questo contesto, la definizione progressiva di *mission*, *vision*, *purpose* e valori ha rappresentato un momento cruciale di allineamento interno, finalizzato non a fissare un'identità astratta, bensì a rendere esplicite e negoziabili le priorità emergenti all'interno di un progetto caratterizzato da elevata complessità e pluralità di approcci.

La *mission* di *Lead-Alliance* può essere ricondotta alla volontà di tradurre riflessioni teoriche e sperimentazioni sull'intelligenza artificiale e la *leadership* in azioni concrete, attraverso il dialogo interdisciplinare, la sperimentazione e la condivisione sistematica di conoscenze. Essa non delimita un ambito rigido di attività, ma definisce una modalità operativa fondata sull'interazione continua tra pensiero e pratica.

La *vision* del progetto si orienta invece verso la costruzione di un futuro in cui intelligenza artificiale e *leadership* possano coesistere in modo sinergico, contribuendo sia alla trasformazione delle organizzazioni sia allo sviluppo della *leadership* a livello individuale. Tale *vision*, per altro, non assume i tratti di una previsione deterministica, ma si configura come un orizzonte interpretativo che guida le scelte del progetto, mantenendo aperta la possibilità di adattamento e riformulazione in relazione ai contesti e alle esperienze maturate.

Il *purpose* di *Lead-Alliance* può essere individuato nel contributo alla messa in discussione dello *status quo* che separa tradizionalmente tecnologia e dimensione umana. Il progetto si propone infatti di promuovere una lettura critica e responsabile dell'adozione dell'IA, interrogandone non solo l'efficienza tecnica, ma anche le

implicazioni etiche, culturali e relazionali. Il *purpose* agisce in questo senso come principio orientativo di fondo, che conferisce senso e direzione alle attività senza tradursi in un obiettivo immediatamente misurabile.

I valori del progetto si articolano di conseguenza attorno ad alcuni assi centrali: la multidisciplinarietà come risorsa, la centralità della dimensione umana nei processi di innovazione, l'attenzione alle implicazioni etiche dell'intelligenza artificiale e la valorizzazione del confronto come strumento di apprendimento collettivo. Tali valori non vengono assunti come enunciazioni astratte, ma come criteri pratici che orientano le modalità di lavoro e le scelte organizzative della seconda fase.

Come già affermato, l'identità di *Lead-Alliance* può essere interpretata come l'esito di un processo di *sensemaking* collettivo, più che come il risultato di una definizione strategica a priori. L'identità non precede quindi l'azione, ma si costruisce retrospettivamente attraverso le pratiche, le interazioni e le narrazioni che consentono agli attori coinvolti di dare senso a ciò che fanno e al contesto in cui operano. *Lead-Alliance* si configura così come un progetto a natura ibrida, che tiene insieme le caratteristiche di *hub* di confronto, laboratorio sperimentale e manifesto culturale. Le diverse traiettorie professionali, i linguaggi disciplinari e le esperienze pregresse dei membri non costituiscono quindi un semplice sfondo individuale, ma diventano una vera e propria infrastruttura del senso, attraverso cui l'identità del progetto prende forma come configurazione dinamica.

La chiarificazione dell'identità è andata poi di pari passo con l'emergere di obiettivi progressivamente più definiti. Nella seconda fase, *Lead-Alliance* ha infatti iniziato a delineare direttrici di sviluppo più chiare, pur mantenendo un margine di flessibilità coerente con la propria natura sperimentale. Tra queste si collocano il posizionamento come punto di riferimento culturale sul tema del rapporto tra intelligenza artificiale e *leadership*, la costruzione di una *community* di pratica autorevole e la funzione di laboratorio creativo orientato alla traduzione delle riflessioni teoriche in soluzioni e strumenti applicabili.

In questo quadro, l'attenzione al pubblico esterno assume un ruolo rilevante nella costruzione dell'identità del progetto. La scelta di produrre *output* pubblici e di strutturare una presenza comunicativa riconoscibile risponde alla volontà di dialogare non solo con gli "addetti ai lavori", ma anche con professionisti, organizzazioni e

istituzioni interessate a comprendere e sperimentare l'integrazione tra IA e *leadership*. La strategia comunicativa multicanale sostiene perciò questo processo di posizionamento, consentendo di articolare l'identità del progetto attraverso linguaggi e formati differenti, pur mantenendo una coerenza narrativa di fondo.

Nel suo complesso, la seconda fase mostra come l'identità di *Lead-Alliance* non sia il risultato di una definizione astratta o puramente dichiarativa, ma il prodotto di un processo di negoziazione continua tra *vision*, pratiche operative e aspettative interne ed esterne. L'identità del progetto emerge conseguentemente come un dispositivo dinamico, capace di tenere insieme innovazione tecnologica, *leadership* etica e inclusività, e di orientarne lo sviluppo nel tempo.

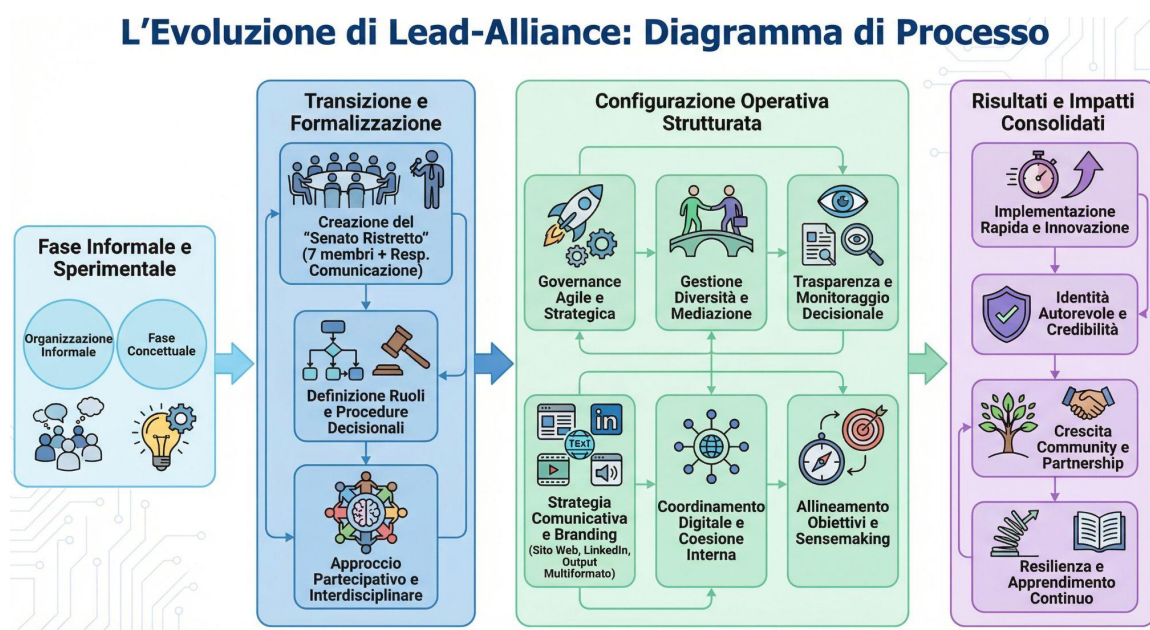


Fig. 2.2: Il processo di evoluzione di *Lead-Alliance*.

### 2.3 Struttura organizzativa e governance interna: il ruolo del senato ristretto

La struttura organizzativa e la governance interna di *Lead-Alliance* si fondano su un approccio partecipativo e interdisciplinare, in cui il nucleo operativo denominato "senato ristretto" svolge la funzione di fulcro decisionale e di coordinamento. Composto da sette membri provenienti da ambiti professionali differenti e affiancato dalla figura della responsabile della comunicazione (ovvero la sottoscritta), il senato ristretto

rappresenta una soluzione organizzativa che consente di integrare rapidamente contributi eterogenei, favorendo l'implementazione operativa delle idee in tempi contenuti, senza sacrificare la pluralità metodologica e culturale che caratterizza il progetto. Proprio questa stessa configurazione ha reso possibile il superamento della fase prevalentemente sperimentale e ha accompagnato *Lead-Alliance* nel suo passaggio a piattaforma operativa.

In una prospettiva coerente con l'impianto teorico adottato, il senato ristretto può essere interpretato non soltanto come organo decisionale formale, ma come dispositivo di *sensemaking*. Esso costituisce uno spazio in cui la complessità generata dall'incontro tra competenze provenienti da settori diversi viene resa intelligibile e trasformata in orientamenti condivisi all'azione. Le decisioni non derivano da un processo lineare o gerarchico, ma emergono attraverso pratiche di confronto periodico, discussione e mediazione, che consentono di mantenere un equilibrio tra rapidità operativa e coerenza strategica.

Un elemento rilevante emerso nella fase operativa riguarda infatti l'impatto della *governance* partecipativa sulla percezione di autoefficacia dei membri. La trasparenza dei processi decisionali, la cadenza regolare delle riunioni e l'uso condiviso di strumenti digitali di coordinamento (come per esempio *Notion*<sup>18</sup> e *Google Drive*<sup>19</sup>) hanno contribuito a rafforzare il senso di responsabilità individuale e collettiva. I membri hanno percepito con maggiore chiarezza il proprio ruolo e il valore del proprio contributo, con effetti positivi sulla motivazione, sulla coesione del gruppo e sul senso di appartenenza al progetto.

L'utilizzo sistematico di strumenti di comunicazione digitale ha ulteriormente sostenuto queste dinamiche, riducendo le incomprensioni e favorendo uno stile relazionale basato su rispetto, fiducia e *feedback* tempestivi. Ogni partecipante ha potuto sperimentare uno spazio di espressione riconosciuto, all'interno di una cornice organizzativa chiara, in cui tempi, responsabilità e obiettivi risultavano costantemente aggiornati e condivisi.

Nel bilanciamento continuo tra urgenza operativa e orientamento strategico, il senato ristretto ha dato spazio alla tensione tra pragmatismo e idealità, intesa come dialogo

---

<sup>18</sup> Cfr. <https://www.notion.com/>

<sup>19</sup> Cfr. <https://drive.google.com/drive/my-drive?hl=it>

costante tra l'esigenza di agire in modo efficiente e quella di mantenere coerenza con i valori fondativi del progetto. Le divergenze emerse rispetto alle priorità operative e agli obiettivi di medio-lungo periodo hanno reso necessario un lavoro costante di riallineamento, attuato attraverso strategie di revisione progressiva del piano operativo e un approccio *step by step*, capace di garantire flessibilità senza perdita di direzione.

Nel suo complesso, il passaggio dalla dimensione sperimentale a quella operativa ha consentito a *Lead-Alliance* di consolidare una posizione peculiare, in cui il progetto si configura simultaneamente come laboratorio di innovazione e come piattaforma organizzata. La *governance* del senato ristretto emerge così come uno dei fattori centrali del successo del progetto, non solo per la sua efficacia decisionale, ma per la sua capacità di sostenere processi di intelligenza collettiva, inclusione e *leadership* partecipativa. In questo senso, la struttura di *governance* adottata non rappresenta un semplice *asset* organizzativo, ma un modello operativo potenzialmente trasferibile ad altri contesti orientati alla sperimentazione e all'innovazione collaborativa.

## **2.4 La struttura dei team come architettura operativa**

Oltre alla suddivisione tra gruppo generale e senato ristretto, la suddivisione in *team* tematici — *podcast*, *newsletter*, *whitepaper* (successivamente evoluto in progetto libro) ed eventi — rappresenta una scelta deliberata di articolazione funzionale dell'ecosistema editoriale.

La configurazione adottata presenta tratti riconducibili a quella che Mintzberg (1983) definisce struttura semplice o organizzazione imprenditoriale: un nucleo strategico ristretto (il senato ristretto, che include i responsabili dei vari *team*) che orienta le decisioni e una periferia operativa flessibile, caratterizzata da coordinamento diretto più che da formalizzazione procedurale. In contesti emergenti e a bassa formalizzazione, la centralità del vertice non implica rigidità, bensì rapidità decisionale e adattabilità. Allo stesso tempo, la suddivisione in *team* tematici introduce elementi di specializzazione funzionale, pur mantenendo confini permeabili. Questa combinazione tra nucleo strategico e moduli operativi riflette una tensione tipica delle organizzazioni ibride: la ricerca di equilibrio tra controllo e autonomia, tra visione unitaria e iniziativa distribuita.

A differenza del senato ristretto, organo di coordinamento e indirizzo strategico, i *team* operativi nascono con una finalità prevalentemente produttiva. L'obiettivo non era

soltanto distribuire il carico di lavoro, ma presidiare formati differenti e complementari. Tale impostazione riflette una logica di copertura multilivello, nella quale ogni *format* risponde a una diversa modalità di attivazione del pubblico.

Gli obiettivi dei singoli *team*, inizialmente impliciti e guidati da convergenza valoriale più che da pianificazione formalizzata, sono stati progressivamente esplicitati nella fase di costruzione del sito *web*. Come si analizzerà nel terzo capitolo, questa transizione dall'intuizione alla formalizzazione segnala un primo processo di maturazione organizzativa: ciò che nasce come energia condivisa richiede, nel momento in cui si espone pubblicamente, una codifica più chiara.

La struttura adottata non presenta tuttavia confini rigidi. I membri possono transitare tra i gruppi in base a disponibilità e competenze, configurando un modello ibrido tra coordinamento e flessibilità. Questa permeabilità risponde alla natura volontaria e professionale del gruppo, composto da individui con impegni lavorativi esterni che incidono sulla continuità operativa.

In tale assetto, non tutte le iniziative hanno però mostrato la stessa sostenibilità. Il *team* eventi, inizialmente concepito come leva di visibilità attraverso *webinar* e incontri in presenza, è stato progressivamente assorbito dagli altri gruppi. La sua dissoluzione non è avvenuta per decisione formale, ma per progressiva inattività dovuta a difficoltà di coordinamento. Tuttavia, l'episodio non è stato interpretato come fallimento, bensì come alleggerimento necessario e ricalibrazione rispetto alle risorse disponibili. Questo passaggio evidenzia un apprendimento organizzativo: la necessità di allineare ambizione progettuale e capacità esecutiva in una fase ancora embrionale.

### CAPITOLO 3: COMUNICARE LEAD-ALLIANCE

Ogni progetto attraversa una fase in cui non esiste ancora come struttura definita, ma esiste come possibilità. *Lead-Alliance*, nei mesi compresi tra l'avvio della seconda fase e la costruzione operativa delle prime iniziative, si è trovata precisamente in questo spazio intermedio: non più idea informale, non ancora organizzazione formalizzata. In questa zona di incubazione, la comunicazione non ha rappresentato un'attività successiva o accessoria, ma il luogo stesso in cui l'identità ha cominciato a prendere forma. Questa dinamica richiama nuovamente il concetto di *sensemaking*, secondo cui le organizzazioni possono produrre una nuova realtà attraverso l'azione comunicativa (Weick, 1995).

Costruire un sito *web*, immaginare un *podcast*, progettare una *newsletter* o un *whitepaper* non ha significato semplicemente scegliere canali di diffusione, bensì interrogarsi su cosa *Lead-Alliance* stesse diventando. *Community* aperta? *Network* professionale? Laboratorio di ricerca? *Brand* editoriale? Le risposte non sono ancora del tutto univoche, e proprio questa ambivalenza ha trasformato la comunicazione in uno spazio di negoziazione identitaria. Albert e Whetten (1985) hanno mostrato infatti come l'identità organizzativa si definisca attraverso tre domande fondamentali: chi siamo, cosa ci rende distinti, cosa permane nel tempo. Nel caso di *Lead-Alliance*, queste domande non trovano ancora una risposta stabile, e sono proprio i processi comunicativi — la scelta delle parole, dei formati, dei toni — a fungere da terreno di negoziazione.

La metafora della navicella e dell'equipaggio, adottata come immaginario condiviso del progetto, non nasce poi come elemento decorativo. Seguendo la prospettiva di Gareth Morgan (1998), le metafore organizzative non sono infatti semplici abbellimenti retorici: strutturano la percezione, orientano l'azione e producono effetti reali sull'identità collettiva. In questo caso, la metafora rappresenta una condizione reale — un gruppo interprofessionale che esplora un territorio in evoluzione, senza mappe definitive, ma con una rotta da tracciare — e al tempo stesso impone una disciplina narrativa: ogni contenuto prodotto deve risultare coerente con quel immaginario, pena la perdita di riconoscibilità. L'intelligenza artificiale, in questa prospettiva, non è semplicemente oggetto di analisi; è l'ambiente stesso in cui il progetto si muove. Comunicare diventa quindi un esercizio di orientamento collettivo.

Parallelamente, emerge una tensione strutturale. *Lead-Alliance* aspira a evolvere verso un modello ibrido, con logiche *freemium* e una possibile traiettoria imprenditoriale in ambito formativo. Tuttavia, nella fase attuale, opera come *network* volontario composto da professionisti già impegnati nei rispettivi contesti lavorativi. Questa condizione genera una frizione costante tra ambizione strategica e sostenibilità operativa — una dinamica ben documentata nella letteratura sui progetti collaborativi a partecipazione volontaria, in cui l'assenza di vincoli formali tende a produrre intermittenza operativa (Cnaan & Goldberg-Glen, 1991). Eventi pensati in grande scala vengono così ridimensionati; *podcast* progettati con forte carica narrativa rallentano di fronte ai vincoli tecnici e temporali; infine, la costruzione del sito diventa prerequisito simbolico e funzionale per ogni ulteriore passo, pur subendo rallentamenti legati ai tempi e agli impegni individuali.

In tale scenario, la comunicazione assume una duplice funzione. Internamente, è strumento di coordinamento poiché definisce la struttura dei *team*, le priorità, le scadenze e i modelli decisionali. Esternamente, è laboratorio pubblico di co-creazione, in cui l'uso dichiarato dell'intelligenza artificiale nei processi editoriali diventa scelta etica e identitaria. Per esempio, nel caso dei progetti che si stanno costruendo, dichiarare che una *newsletter* è redatta con il supporto dell'IA non costituisce mera trasparenza tecnica, ma una presa di posizione culturale che si inserisce nel dibattito più ampio sul rapporto tra tecnologia e responsabilità editoriale.

Questo capitolo analizza la comunicazione di *Lead-Alliance* come pratica di *leadership* in atto. Non si tratta di descrivere una strategia già consolidata, ma di osservare il processo attraverso cui una visione si traduce in infrastruttura digitale, in contenuti, in ritualità operative e in apprendimento organizzativo. Attraverso l'esame delle dinamiche interne, della costruzione del sito *web* e della *content strategy* si intende mostrare come la comunicazione diventi il dispositivo attraverso cui un progetto ancora fluido tenta di stabilizzare la propria traiettoria. In questo senso, il capitolo adotta una prospettiva processuale: non fotografa un'organizzazione già formata, ma segue il movimento attraverso cui essa prende forma.

L'analisi che segue racconta quindi scelte riuscite e tentativi interrotti, accelerazioni e stalli, entusiasmo e vincoli. In un ecosistema digitale saturo di discorsi sull'intelligenza artificiale, la differenza non sta nella quantità di contenuti prodotti, ma nella coerenza

tra visione dichiarata, *governance* adottata e capacità di esecuzione. La coerenza, infatti, è la condizione necessaria affinché la strategia non resti pura intenzione, ma si traduca in azione organizzativa riconoscibile (Mintzberg, 1994).

### **3.1 Comunicazione interna: coordinamento, fiducia e costruzione di pratiche condivise**

L'inizio della seconda fase di *Lead-Alliance* si colloca in uno spazio organizzativo intermedio a metà tra sperimentazione condivisa e struttura non ancora formalizzata. In questa condizione, la comunicazione interna non si limita a trasmettere informazioni operative, bensì contribuisce direttamente alla configurazione del progetto. In assenza di un assetto giuridico definito o di risorse dedicate a tempo pieno, il coordinamento si fonda prevalentemente sulla qualità delle interazioni e sulla capacità di allineare visioni. Questo è un esempio di ciò che Wenger (1998) descrive come processo costitutivo delle comunità di pratica, in cui la partecipazione condivisa precede e produce la struttura formale.

Fin dalle prime riunioni operative è emersa quindi la necessità di rendere governabile la complessità. Il gruppo, composto da professionisti con competenze e impegni differenti, rischiava di restare in una dimensione puramente ideativa. La scelta di suddividere il lavoro in *team* tematici e legati ad *output* — che inizialmente erano quattro: eventi, *newsletter*, *podcast* e *whitepaper* — nasce da un'esigenza pragmatica: trasformare l'entusiasmo collettivo in micro-cantieri operativi. Questa impostazione riflette da un lato l'esperienza presente in altre iniziative simili a cui ci si è ispirati — quali le già citate *junior enterprise* — dall'altro un principio consolidato nella letteratura organizzativa, secondo cui la differenziazione in sottogruppi consente di ridurre i costi di coordinamento e aumentare il senso individuale di responsabilità (Mintzberg, 1979).

Tuttavia, la divisione non è mai stata concepita come rigida. I membri possono spostarsi tra i *team* in base a disponibilità e competenze, mantenendo una configurazione fluida coerente con la natura volontaria dell'iniziativa. Questa elasticità riflette un modello di *governance* distribuita: la *leadership* è diffusa nell'operatività quotidiana, ma esiste un nucleo decisionale ristretto che garantisce coerenza complessiva. Tale assetto è coerente con quanto descritto da Gronn (2002), secondo cui la *leadership* distribuita non implica assenza di centri di responsabilità, ma una

ridefinizione del rapporto tra autonomia locale e coordinamento generale. Le decisioni vengono generalmente discusse in modo aperto, con un orientamento più visionario che strettamente strategico; quando emergono divergenze, si procede attraverso votazione e negoziazione, fino a una sintesi affidata alla responsabilità congiunta del coordinamento generale.

Questo equilibrio tra partecipazione diffusa e responsabilità ultima evita tanto la frammentazione quanto l'accentramento decisionale. In una fase di incubazione, la chiarezza decisionale non può dipendere esclusivamente dal consenso espresso in modo informale, ma necessita di un punto di stabilizzazione — in linea con ciò che Laloux (2014) definisce come tensione costitutiva delle organizzazioni auto-organizzate, in cui libertà e struttura non si escludono ma si integrano in forme ibride e adattive.

La comunicazione interna svolge, in questo quadro, una funzione di progressiva strutturazione. La definizione di scadenze, l'introduzione di strumenti di pianificazione come i diagrammi di Gantt e la ritualità degli incontri periodici rappresentano tentativi di tradurre la visione in temporalità concreta. La letteratura sui progetti collaborativi a partecipazione volontaria mostra come la distanza tra ideazione e implementazione costituisca la criticità più ricorrente: l'entusiasmo genera proposte ambiziose, ma la realtà operativa richiede selezione, priorità e ridimensionamento (Deci & Ryan, 2000). La motivazione intrinseca, se non sostenuta da meccanismi minimi di *accountability*, tende infatti a concentrarsi sulla fase ideativa, simbolicamente gratificante, a scapito della fase esecutiva, più ripetitiva e meno visibile.

Un caso emblematico è rappresentato dal *team* Eventi. L'idea iniziale prevedeva *webinar* di ampia portata e incontri in presenza con ospiti di rilievo, in linea con la visione espansiva del progetto. Tuttavia, la complessità organizzativa richiesta e la difficoltà di stabilire una ritualità stabile di incontro hanno progressivamente rallentato il lavoro, fino alla decisione di sciogliere il gruppo e redistribuire le energie su iniziative più sostenibili. L'episodio non costituisce un fallimento isolato, ma un caso di apprendimento organizzativo a ciclo singolo: il gruppo ha riconosciuto lo scarto tra obiettivi dichiarati e capacità operative, correggendo la rotta senza tuttavia rimettere in discussione le assunzioni di fondo.

Parallelamente, la pratica sistematica di utilizzare l'intelligenza artificiale nelle fasi preliminari di scrittura e revisione introduce un elemento distintivo nella dinamica

interna. I testi vengono solitamente sottoposti a una prima elaborazione tramite IA prima di essere discussi collettivamente. Questa modalità non sostituisce il confronto umano, ma lo prepara: l'IA funziona come dispositivo di pre-strutturazione del pensiero, che abbassa la soglia di ingresso alla discussione e aumenta la qualità del materiale condiviso. Il risultato è una forma di co-creazione che coinvolge autore, gruppo e strumento tecnologico in una sequenza iterativa, ovvero un modello di collaborazione uomo-macchina in cui le rispettive capacità si integrano anziché sostituirsi. La comunicazione interna diventa così laboratorio concreto di quella *leadership* aumentata che il progetto stesso si propone di esplorare a livello teorico.

In questo senso, la comunicazione non è semplicemente uno strumento organizzativo. È il dispositivo attraverso cui *Lead-Alliance* tenta di trasformare una visione condivisa in struttura emergente. Mantenere un clima armonico e visionario senza perdere ancoraggio operativo non è dunque solo una sfida gestionale, ma una condizione costitutiva del progetto stesso.

La stabilizzazione progressiva del progetto non avviene soltanto attraverso strumenti di pianificazione, ma anche mediante la costruzione di una ritualità condivisa. In assenza di una struttura gerarchica formalizzata o di una presenza quotidiana in uno spazio fisico comune, la periodicità degli incontri diventa il principale dispositivo di coesione. Questa dinamica richiama il concetto di confine simbolico elaborato da Wenger (1998): in una comunità di pratica distribuita, sono le pratiche ricorrenti, e non le regole formali, a produrre appartenenza e identità collettiva.

Accanto alle riunioni operative del senato ristretto, si è poi consolidata la pratica di incontri plenari con l'intero gruppo, generalmente con cadenza bi o trimestrale. Questi momenti non hanno esclusivamente una funzione informativa, ma assumono un valore simbolico: rappresentano il luogo in cui la visione viene rinnovata, le traiettorie vengono riallineate e le tensioni latenti trovano uno spazio di verbalizzazione controllata. In un contesto armonico-consensuale, il rischio non è il conflitto esplicito, bensì la dispersione silenziosa. La ritualità periodica contrasta proprio questa deriva, riaffermando l'esistenza di un orizzonte comune — ciò che Weick (1995, pp. 90-100) definisce come funzione di *sensegiving*: la capacità di un gruppo di ridefinire collettivamente il significato della propria azione in momenti di incertezza.

Il clima che caratterizza tali incontri è prevalentemente visionario. Le riunioni non si limitano alla verifica dello stato di avanzamento, ma tendono ad aprire scenari, generare nuove idee, immaginare sviluppi futuri. Questo orientamento alimenta la motivazione e consolida l'identità del gruppo, ma richiede un successivo lavoro di selezione e concretizzazione nei momenti più ristretti del coordinamento. Si produce così una dinamica oscillatoria tra espansione immaginativa e restringimento operativo, che è una condizione tipica dei sistemi organizzativi ad apprendimento attivo, in cui la generazione di nuove ipotesi deve essere bilanciata dalla capacità di valutarle e selezionarle (Argyris e Schön, 1978).

La ritualità svolge inoltre una funzione di legittimazione reciproca. Ogni membro, pur provenendo da ambiti professionali differenti, trova in questi spazi un riconoscimento del proprio contributo. La condivisione pubblica degli avanzamenti, anche minimi, rafforza la percezione di appartenenza e riduce la sensazione di isolamento che può emergere nei progetti distribuiti. In tal senso, la comunicazione interna non si limita a coordinare attività, ma produce capitale simbolico: rende visibile l'impegno e lo inserisce in una narrazione collettiva (Bourdieu, 1991, pp. 163-200). Il riconoscimento reciproco non è dunque un effetto collaterale della comunicazione, ma una sua funzione costitutiva.

Un ulteriore elemento rilevante è la progressiva consapevolezza della necessità di strutturare il tempo. Se nella fase iniziale la motivazione fungeva da motore sufficiente, con il passare dei mesi emerge la necessità di strumenti più rigorosi di pianificazione. L'introduzione dei diagrammi di Gantt in alcuni *team*, in particolare in quello che inizialmente era il *team whitepaper*, rappresenta un passaggio verso una maggiore maturità organizzativa. La ritualità non resta così confinata alla dimensione simbolica, ma si integra con una scansione temporale più definita — riducendo il rischio di quella procrastinazione strutturale che nei contesti volontari tende a manifestarsi non come negligenza, ma come effetto sistemico dell'assenza di scadenze vincolanti (Cnaan & Goldberg-Glen, 1991).

In questa combinazione di incontri periodici, clima armonico e progressiva strutturazione temporale si delinea una forma embrionale di *governance* adattiva. Non si tratta di un modello codificato, ma di un equilibrio dinamico tra visione condivisa e necessità operative: una struttura organica capace di adattarsi alle condizioni ambientali

senza irrigidirsi in forme burocratiche (Burns & Stalker, 1961). La comunicazione interna agisce così come infrastruttura invisibile che permette al progetto di non dissolversi nella dispersione delle agende individuali. La coesione non deriva quindi dalla formalizzazione giuridica, bensì dalla ripetizione di pratiche: è la ripetizione che genera riconoscibilità, e la riconoscibilità è ciò che consolida l'identità.

Accanto alla ritualità e alla progressiva strutturazione temporale, emergono due questioni che incidono in modo determinante sulla sostenibilità del progetto nel medio periodo. La prima riguarda, come già sottolineato, la natura volontaria dell'impegno; la seconda il clima armonico e consensuale che caratterizza le dinamiche interne. Entrambe rappresentano condizioni tipiche dei progetti collaborativi in fase di incubazione e, come tali, meritano un'analisi che vada oltre la semplice registrazione del dato empirico.

*Lead-Alliance* si fonda su una motivazione intrinseca molto elevata. I membri partecipano perché condividono una visione. Questo produce un livello di entusiasmo iniziale significativo e una disponibilità a investire tempo e competenze in modo generoso. Tuttavia, proprio l'assenza di obblighi formali introduce una variabile critica: la priorità effettiva del progetto dipende dalla compatibilità con gli impegni professionali individuali. Secondo la teoria dell'autodeterminazione, la motivazione intrinseca è una risorsa potente ma fragile, che tende a erodersi in assenza di strutture minime di supporto e riconoscimento (Deci & Ryan, 2000). La motivazione, se non sostenuta da una struttura minima di *accountability*, rischia dunque di trasformarsi in intermittenza operativa.

La dissoluzione del *team* eventi può essere letta anche in questa chiave. Non si è trattato esclusivamente di una difficoltà organizzativa, ma di uno scarto tra ambizione progettuale e capacità reale di allocare tempo continuativo. L'esperienza conferma quanto documentato nella letteratura sui contesti volontari: l'energia iniziale tende a concentrarsi sulla fase ideativa, mentre la fase esecutiva, più ripetitiva e meno simbolicamente gratificante, richiede meccanismi di responsabilizzazione più espliciti (Cnaan & Goldberg-Glen, 1991). La progressiva introduzione di strumenti di pianificazione e la definizione più chiara dei ruoli rappresentano tentativi di rispondere a questa criticità senza snaturare la dimensione partecipativa.

Parallelamente, il clima armonico-consensuale costituisce una forza e, al tempo stesso, una potenziale fragilità. L'assenza di conflitti espliciti facilita la coesione e riduce l'attrito decisionale; al tempo stesso, può attenuare quella frizione critica necessaria per selezionare con rigore le priorità. Janis (1982) ha descritto questo fenomeno come *groupthink*: la tendenza dei gruppi coesi a privilegiare la convergenza rispetto alla valutazione critica delle alternative, con il rischio di produrre decisioni subottimali. In un ambiente prevalentemente visionario, la produzione di idee tende a superare la capacità di implementazione; se non adeguatamente filtrata, questa espansione immaginativa rischia di generare sovraccarico progettuale.

La cultura del consenso, in particolare, può indurre una tendenza implicita a preservare l'armonia anche a scapito della chiarezza. Il ricorso alla votazione e alla negoziazione rappresenta uno strumento di riequilibrio, ma non elimina del tutto il rischio di eccessiva convergenza. In questo senso, la *leadership* del senato ristretto assume una funzione di contenimento: non solo coordinare, ma talvolta selezionare e ridimensionare. Questa funzione corrisponde a ciò che Heifetz (1994) definisce come *adaptive leadership*: la capacità di mantenere un livello produttivo di tensione all'interno del gruppo, evitando tanto il conflitto distruttivo quanto la convergenza acritica.

Le due tensioni menzionate (la motivazione volontaria e l'armonia consensuale) non devono però essere interpretate come limiti strutturali, bensì come condizioni tipiche dei progetti collaborativi in fase di incubazione. La loro consapevolezza costituisce già un elemento di maturazione organizzativa. La comunicazione interna, lungi dall'essere neutra, diventa lo spazio in cui tali tensioni vengono riconosciute, rielaborate e progressivamente integrate in pratiche più sostenibili. Nel momento in cui un progetto aspira a evolvere verso un modello ibrido con logiche *freemium* e, in prospettiva, verso una possibile configurazione imprenditoriale, la gestione di queste tensioni diventa cruciale: la motivazione deve trasformarsi in continuità, e l'armonia deve convivere con una capacità selettiva più marcata. In tale transizione, la comunicazione interna non è solo supporto organizzativo, ma laboratorio di apprendimento istituzionale.

Nel momento in cui un progetto aspira a evolvere verso un modello ibrido con logiche *freemium* e, in prospettiva, verso una possibile configurazione imprenditoriale, la gestione di queste tensioni diventa cruciale. La motivazione deve trasformarsi in

continuità, e l'armonia deve convivere con una capacità selettiva più marcata. In tale transizione, la comunicazione interna non è solo supporto organizzativo, ma laboratorio di apprendimento istituzionale.

Se la ritualità consolida l'identità attraverso la ripetizione degli incontri, la costruzione del sito *web* introduce un elemento ulteriore: la materializzazione visibile del progetto. Prima della pubblicazione effettiva, il sito assume già una funzione catalizzatrice interna, diventando il punto attorno a cui convergono discussioni, revisioni, scelte lessicali, definizioni di missione e priorità narrative. Il sito non è uno strumento neutro: è un artefatto che orienta l'azione organizzativa e, nel processo della sua costruzione, produce effetti reali sulla struttura del gruppo (Orlikowski, 2007).

La necessità di riempire il sito di contenuti — come la descrizione della missione, del gruppo e delle prime iniziative — costringe infatti il *team* a raccontare in maniera esplicita ciò che fino a quel momento era rimasto un pensiero condiviso implicito. Scrivere la sezione "Chi siamo" o definire la *mission* non significa semplicemente redigere un testo; implica stabilire confini identitari. È in questo passaggio che emergono interrogativi più profondi: siamo principalmente piattaforma, manifesto, *hub* o laboratorio? Qual è il nostro livello di formalizzazione? Qual è l'orizzonte temporale che intendiamo dichiarare pubblicamente? Queste domande rimandano direttamente alle tre dimensioni dell'identità organizzativa individuate da Albert e Whetten (1985): centralità, distintività e continuità temporale.

Il sito, dunque, non nasce solo come strumento di comunicazione esterna, ma come dispositivo di auto-chiarificazione. Ogni parola inserita diventa un impegno simbolico. La scelta di presentarsi come piattaforma e, al tempo stesso, come manifesto per l'integrazione tra intelligenza artificiale e *leadership* richiede coerenza tra visione e capacità operativa. L'atto di scrivere obbliga a selezionare, sintetizzare, talvolta rinunciare, producendo quello che Weick (1995, pp. 24-47) definisce come effetto di *commitment*: l'azione comunicativa vincola retroattivamente il gruppo a una traiettoria, rendendo alcune scelte più difficili da revocare.

Inoltre, l'adozione di un immaginario fantascientifico coerente — la metafora della navicella, della *crew*, della *home* come base — svolge inizialmente una funzione interna. Fornisce infatti un linguaggio condiviso che facilita l'identificazione e rafforza il senso di appartenenza. Seguendo Morgan (1998), non si tratta quindi di un espediente

estetico, ma di una grammatica simbolica attraverso cui il gruppo si riconosce e orienta la propria azione. La metafora consente di rappresentare la fase di esplorazione senza nascondere la dimensione sperimentale del progetto; al tempo stesso, obbliga a una disciplina narrativa: mantenere coerenza tra immaginario e contenuti diventa parte integrante della *governance* comunicativa.

La costruzione tecnica del sito, affidata a una piattaforma come *WordPress*<sup>20</sup>, introduce poi una dimensione concreta di apprendimento organizzativo. La necessità di comprendere strumenti di *editing*, organizzare sezioni, definire *call to action* e integrare moduli come la *newsletter* tramite *Mailchimp*<sup>21</sup> rende evidente che la comunicazione non è solo visione, ma infrastruttura. Il ritardo nella pubblicazione non è semplicemente un rallentamento operativo; è l'effetto di una fase in cui il progetto sta ancora definendo il proprio perimetro.

In questo senso, il sito *web*, che verrà analizzato meglio nelle prossime sezioni, funziona come primo banco di prova della trasformazione da comunità informale a entità riconoscibile. Prima ancora che generare visibilità esterna, produce un effetto di consolidamento interno: rende tangibile ciò che fino a quel momento era stato prevalentemente discorsivo. La comunicazione, dunque, non si limita a raccontare *Lead-Alliance*; contribuisce a farla esistere in una forma stabile confermando, sul piano pratico, la prospettiva enattiva di Weick (1995) secondo cui le organizzazioni non precedono la comunicazione, ma emergono attraverso di essa.

Infine, per quanto concerne il tema della fiducia, in questa fase di identificazione è emersa nel tempo la consapevolezza della necessità di strumenti di ascolto strutturato, come l'introduzione di un questionario anonimo di riallineamento interno<sup>22</sup>, volto a prevenire possibili fenomeni di *disengagement* silenzioso. Il tema dell'erosione progressiva dell'*engagement* è stato più volte richiamato anche nel dibattito recente sull'innovazione organizzativa. Per esempio, durante la conferenza *Tracciare la rotta del cambiamento: AI, nuove strategie e competenze per il futuro del lavoro* degli Osservatori Digital Innovation del Politecnico di Milano (13 maggio 2025), è stato evidenziato come molto spesso una cultura organizzativa fragile e la mancata

---

<sup>20</sup> Cfr. <https://it.wordpress.org/>

<sup>21</sup> Cfr. <https://mailchimp.com/it/>

<sup>22</sup> Tale questionario verrà analizzato nel dettaglio nell'ultima sezione del presente capitolo.

valorizzazione delle competenze possano incidere negativamente sul coinvolgimento, aumentando il rischio di *turnover* e disaffezione<sup>23</sup>. A tal proposito, proprio gli strumenti di *assessment* vengono indicati come leve per l'*engagement* interno, poiché favoriscono fiducia, riconoscimento, *purpose* e valorizzazione delle competenze.

In questa prospettiva, il questionario di *Lead-Alliance* non rappresenta un semplice strumento operativo, ma un dispositivo culturale di ascolto attivo, volto a intercettare segnali deboli prima che si traducano in abbandono silenzioso.

### **3.2 Comunicazione esterna e brand identity: posizionamento, TOV e riconoscibilità**

Se la comunicazione interna ha contribuito alla stabilizzazione organizzativa del progetto, la comunicazione esterna introduce una dimensione ulteriore: la costruzione di una identità riconoscibile nello spazio pubblico. In questa fase, *Lead-Alliance* non si configura ancora come entità imprenditoriale strutturata, né come *think tank* formalizzato, ma come laboratorio culturale che interroga criticamente il rapporto tra *leadership* e intelligenza artificiale. Tale auto-definizione orienta in modo decisivo le scelte di posizionamento.

Secondo la teoria del posizionamento, ciò che distingue un'organizzazione non è soltanto ciò che produce, ma il posto che occupa nella mente del pubblico di riferimento (Ries, A., & Trout, J., 1981). In un ecosistema comunicativo saturo di contenuti sull'intelligenza artificiale, spesso polarizzati tra entusiasmo tecnocratico e allarmismo etico, *Lead-Alliance* sceglie di collocarsi in una zona intermedia: non come fornitore di soluzioni immediate, ma come spazio di riflessione applicata. Il posizionamento come "laboratorio culturale" consente di valorizzare il capitale teorico già accumulato (tra cui la progettualità editoriale, le ricerche e le tesi sviluppate sul progetto) e, al contempo, di mantenere apertura sperimentale.

In questa fase embrionale si avvia anche la costruzione della reputazione organizzativa. Fombrun (1996) definisce la reputazione come una rappresentazione collettiva costruita nel tempo attraverso la coerenza tra identità dichiarata,

---

<sup>23</sup> Cfr.

<https://www.osservatori.net/convegno/hr-innovation-practice/hr-innovation-practice-osservatorio-convegno-risultati-ricerca/>

comportamenti osservabili e comunicazione pubblica. In tale prospettiva, per esempio, la scelta di rendere visibili i processi di evoluzione del progetto anche sul sito *web*, prima ancora dei risultati, non è soltanto un atto comunicativo, ma un investimento reputazionale: l'organizzazione si espone e si vincola alla propria narrazione.

La *brand identity*, in questo senso, non è costruita attorno a un prodotto, ma attorno a un orientamento valoriale. Le teorie sull'identità di marca evidenziano come un *brand* non coincida con un logo o con un insieme di canali, bensì con un sistema coerente di significati, simboli e promesse implicite (Aaker, D. A., 1996, Kapferer, J.-N., 2012). Nel caso di *Lead-Alliance*, l'identità si struttura attorno a tre elementi: la centralità del contenuto teorico, la dimensione narrativa e l'attenzione etica.

Il contenuto teorico rappresenta quindi il principale capitale simbolico del progetto. L'esistenza di elaborazioni strutturate, la progettazione di un libro collettivo e la produzione di riflessioni approfondite conferiscono legittimità culturale. Tuttavia, ciò che rende riconoscibile l'iniziativa non è soltanto la qualità del contenuto, ma la modalità con cui esso viene raccontato. L'adozione di un immaginario fantascientifico non si configura in questo senso come elemento ludico superficiale, bensì come scelta strategica di differenziazione narrativa. In un panorama comunicativo dominato da linguaggi tecnici o *corporate*, l'immaginario esplorativo consente di tradurre complessità teoriche in metafore accessibili, mantenendo al contempo una coerenza simbolica.

Il *tone of voice* si colloca in equilibrio tra autorevolezza e accessibilità. Non assume il registro freddamente tecnocratico tipico di molte comunicazioni sull'IA, ma evita anche l'eccessiva semplificazione divulgativa. La scelta di dichiarare esplicitamente l'utilizzo dell'intelligenza artificiale nei processi di scrittura, ad esempio nella *newsletter* sperimentale, rappresenta poi una presa di posizione identitaria significativa. In un contesto in cui l'uso dell'IA nella produzione di contenuti è spesso implicito o non dichiarato, la trasparenza diventa elemento distintivo e rafforza la coerenza tra valori professati e pratiche adottate.

La riconoscibilità, in questa fase, non deriva ancora da metriche quantitative consolidate, quali numero di *follower*, tasso di conversione o *engagement* strutturato, bensì dalla coerenza narrativa e dalla chiarezza del posizionamento. La costruzione del sito *web* e della presenza *LinkedIn* non risponde prioritariamente a una logica di

conversione commerciale, ma alla necessità di stabilire un perimetro identitario visibile. Rendere pubblico il progetto significa cristallizzarne i contorni e assumere una responsabilità simbolica verso l'esterno.

In questo senso, la comunicazione esterna non è mera amplificazione. È un processo di stabilizzazione attraverso l'esposizione. L'identità, dichiarata pubblicamente, smette di essere ipotesi interna e diventa impegno. Proprio per questo, la progettazione del sito non è partita da un'architettura tecnica, ma anzitutto da una domanda identitaria: “Come vogliamo apparire?”. Questa inversione di priorità — dalla tecnica alla narrazione — segnala come la comunicazione esterna venga concepita come costruzione simbolica prima che infrastruttura digitale. Solo successivamente l'architettura tecnica si è adattata alla visione narrativa, confermando che l'identità precede lo strumento.

In aggiunta, la già citata scelta di dichiarare esplicitamente l'uso dell'IA nella produzione e revisione di alcuni contenuti, in particolare la *newsletter*, non viene trattata esclusivamente come opzione etica, ma come elemento di posizionamento differenziante. In un contesto in cui l'uso dell'intelligenza artificiale nei processi comunicativi è diffuso ma raramente esplicitato, la trasparenza diventa atto identitario. Il tema della *trustworthy AI communication*, emerso anche nel dibattito promosso dagli Osservatori Digital Innovation del Politecnico di Milano nel contesto della conferenza *Trust and governance in AI: building responsible artificial intelligence for Europe* (2 luglio 2025)<sup>24</sup>, si inserisce nel più ampio quadro europeo di regolazione e *governance* dell'IA, che pone al centro *accountability*, tracciabilità e chiarezza nell'uso dei sistemi intelligenti. In questo senso, la trasparenza non è solo coerenza valoriale, ma costruzione di fiducia.

La costruzione dell'identità non si è però limitata alla definizione teorica del posizionamento, ma ha trovato una prima verifica pubblica in occasione della presentazione del progetto all'interno di una lezione di Organizzazione Aziendale presso l'Università degli Studi di Milano, il 4 dicembre 2025, che ho svolto personalmente insieme ad Alice De Gennaro e sotto la guida del prof. Luigi Di Iorio.

---

<sup>24</sup> Cfr.

<https://eng.osservatori.net/conference/european-digital-tech-watch-eng/convegno-dei-risultati-di-ricerca-osservatorio-european-digital-tech-watch/>

L'interazione con gli studenti ha rappresentato un banco di prova significativo: l'interesse manifestato e i suggerimenti emersi hanno confermato la percezione di *Lead-Alliance* non come semplice iniziativa divulgativa, ma come possibile modello organizzativo evolutivo. Le proposte di miglioramento fornite dagli studenti — dalla costruzione di un modello di *leadership data-driven*, all'ipotesi di un'architettura *freemium* (applicazione gratuita con abbonamento), fino alla configurazione di un'offerta formativa strutturata rivolta a PMI e formatori — hanno evidenziato come il progetto venga spontaneamente interpretato in chiave imprenditoriale e sistemica. Questo episodio suggerisce quindi che la *brand identity* non è percepita esclusivamente come progetto culturale, ma come piattaforma potenzialmente scalabile. L'identità dichiarata come laboratorio culturale si confronta così con una tensione evolutiva verso modelli più strutturati, senza tuttavia perdere la propria matrice sperimentale.

### **3.3 Content strategy e community building: sperimentazioni e output**

Se nel paragrafo precedente la comunicazione esterna è stata analizzata soprattutto come costruzione identitaria e posizionamento pubblico, in questa sezione l'attenzione si sposta sul livello operativo: la produzione di contenuti e la costruzione progressiva della *community* come forma di sperimentazione organizzativa. In un progetto che si può definire al momento come laboratorio culturale, la *content strategy* non rappresenta un apparato accessorio, ma un dispositivo di ricerca applicata. I contenuti non vengono concepiti soltanto come strumenti di visibilità, bensì come occasioni per testare linguaggi, formati e modalità di interazione, mettendo in relazione produzione simbolica e apprendimento del gruppo.

La seconda fase di *Lead-Alliance* si caratterizza per la scelta di articolare il lavoro in nuclei operativi dedicati a *format* differenti. Tale impostazione risponde a una logica di presidio dell'ecosistema editoriale: *podcast*, *newsletter* e progetto editoriale (inizialmente *whitepaper* e successivamente ipotizzato come libro) vengono pensati come canali complementari, capaci di attivare livelli diversi di profondità e accessibilità. In questa fase, la costruzione dei contenuti non procede secondo una linearità tipica delle organizzazioni già strutturate, ma secondo un andamento iterativo. Alcune iniziative avanzano rapidamente, altre rallentano, alcune vengono ricalibrate in base alla sostenibilità. Questa dinamica non costituisce un'anomalia, bensì, come già

evidenziato, un tratto coerente con un contesto embrionale e volontario, composto da professionisti con disponibilità di tempo limitata e con vincoli lavorativi esterni.

All'interno di tale scenario, la strategia di *community building* non viene intesa come crescita puramente quantitativa, ma soprattutto come attivazione progressiva di interazioni. L'obiettivo non è soltanto essere visibili, ma costruire un pubblico che partecipi, commenti, suggerisca e contribuisca alla definizione del progetto, coerentemente con l'impostazione di co-creazione che *Lead-Alliance* rivendica anche sul piano della *governance*. Per questo motivo, la lettura degli *output* prodotti in questa fase deve essere condotta non soltanto in termini di performance o risultati, ma come osservazione di processi: ogni *format* attivato diventa un esperimento che mette alla prova la capacità del gruppo di coordinarsi, mantenere coerenza narrativa e trasformare la visione in presenza pubblica.

I paragrafi che seguono analizzano dunque i principali cantieri comunicativi della seconda fase, descrivendone logica, obiettivi e criticità emergenti. L'attenzione sarà rivolta sia agli aspetti di progettazione editoriale sia alle implicazioni organizzative: ciò che cambia, ciò che rallenta, ciò che si consolida e ciò che, in modo significativo, viene rimandato o riformulato.

### **3.3.1 Il sito come infrastruttura narrativa e spazio di legittimazione**

Prima dell'attivazione dei singoli *format* editoriali, *Lead-Alliance* ha ritenuto necessario dotarsi di un'infrastruttura digitale capace di rendere visibile e coerente l'identità del progetto. Il sito web (<https://www.leadalliance-ai.com/>) non nasce come semplice canale comunicativo, ma come dispositivo di formalizzazione organizzativa. In una fase in cui il progetto non è ancora configurato come entità imprenditoriale strutturata, la costruzione di uno spazio digitale pubblico rappresenta un atto performativo: dichiarare l'esistenza significa iniziare a stabilizzarla.



*Fig. 3.3.1: Header e Menu del sito web.*

*FONTE: sito ufficiale di Lead-Alliance.*

L'identità organizzativa non è soltanto un dato interno, ma un processo che si consolida attraverso l'esposizione pubblica e la messa in scena coerente dei propri elementi distintivi (Goffman, 1959). Il sito, in questa prospettiva, agisce come palcoscenico istituzionale: rende visibili ruoli, missione, valori e struttura decisionale, contribuendo alla costruzione di una realtà organizzativa riconoscibile.

L'architettura del sito riflette tale intenzione. Le sezioni principali — *Home Base*, *Missione*, *Crew*, *Risorse* e iscrizione alla *newsletter* — rispondono a una logica prevalentemente funzionale, ma sono attraversate da una coerenza simbolica che richiama l'immaginario esplorativo sviluppato nel progetto. Per esempio, la scelta consapevole di denominare l'*homepage* come *Home Base* segnala l'intenzione di configurare il sito non come semplice vetrina, bensì come base operativa da cui si diramano le iniziative editoriali.

Dal punto di vista decisionale, la *governance* del sito è concentrata in un nucleo ristretto composto dalla responsabile della comunicazione, ovvero la sottoscritta (che propone e struttura i contenuti), dal referente tecnico (che ne valida la fattibilità operativa) e dal fondatore del progetto (Luigi Di Iorio, che fornisce l'approvazione finale). Questa triade decisionale evidenzia una gerarchia funzionale chiara ma non conflittuale, coerente con la struttura del senato ristretto già descritta nel capitolo precedente. L'assenza di conflitti simbolici rilevanti nella definizione dei contenuti finali proposti suggerisce un elevato allineamento identitario tra i membri centrali del progetto.

Il sito è stato sviluppato in ambiente di *test* e successivamente pubblicato sul *web* la sera del 10 Marzo 2026 durante un momento condiviso, concepito quasi come un evento di attivazione collettiva, del senato ristretto, in concomitanza con l'apertura della pagina *LinkedIn* ufficiale, in formato aziendale<sup>25</sup>. Nei giorni precedenti al lancio, si è deliberatamente costruita una forma di attesa interna nel gruppo, condividendo *teaser* visivi e messaggi di anticipazione: le mascotte del progetto — due figure supereroiche e un piccolo *robot*, inizialmente coperti da un velo — sono state presentate come simboli di un'identità pronta a emergere pubblicamente.



*Fig. 3.3.1: Schermate di preparazione al lancio del sito web.*

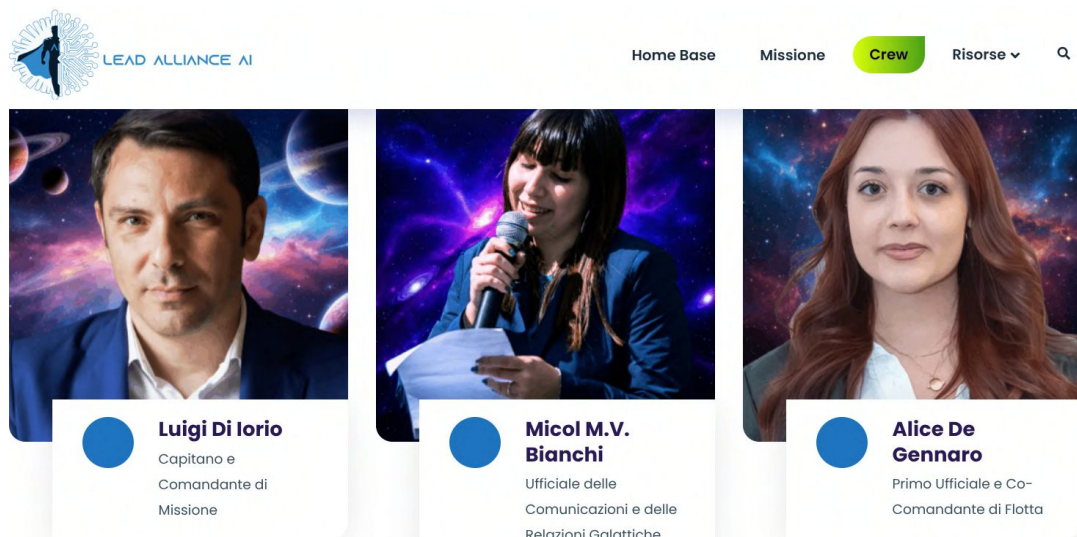
*FONTE: sito ufficiale e LinkedIn di Lead-Alliance.*

<sup>25</sup> Su suddetta pagina sono stati anche pubblicati, nei giorni precedenti alla pubblicazione del sito, *post* che a loro volta sono serviti a creare aspettativa e *hype* nei confronti del progetto.

Il ricorso a questa strategia di creazione di aspettativa e attesa non rispondeva a una logica promozionale in senso stretto, bensì a un'esigenza di coesione: trasformare un passaggio tecnico in un momento narrativo condiviso. Il disvelamento delle *mascotte* ha assunto così la funzione di rito di passaggio, segnando il momento in cui *Lead-Alliance* cessava di essere esclusivamente spazio progettuale interno per assumere una forma riconoscibile e dichiarata all'esterno.

Questo passaggio evidenzia inoltre come la messa *online* non venga percepita come mero completamento tecnico, ma come soglia simbolica: il passaggio da progetto interno a presenza pubblica. Anche le questioni di conformità normativa (formalizzazione delle *policy* tramite consulenza legale esterna) hanno contribuito a rendere più esplicita la struttura organizzativa, obbligando il gruppo a tematizzare responsabilità, investimenti e gestione dei dati.

In termini funzionali, il sito svolge tre ruoli principali. In primo luogo, rappresenta una vetrina istituzionale, offrendo un perimetro identitario chiaro a interlocutori esterni. In secondo luogo, agisce come spazio di legittimazione, rendendo visibile la composizione della *crew* e la missione del progetto. In terzo luogo, attraverso la sezione *Diario di bordo*, assume una funzione narrativa, documentando l'evoluzione delle attività e trasformando il processo in contenuto.



*Fig. 3.3.1: Alcuni componenti del senato ristretto.*

*FONTE: sito ufficiale di Lead-Alliance.*

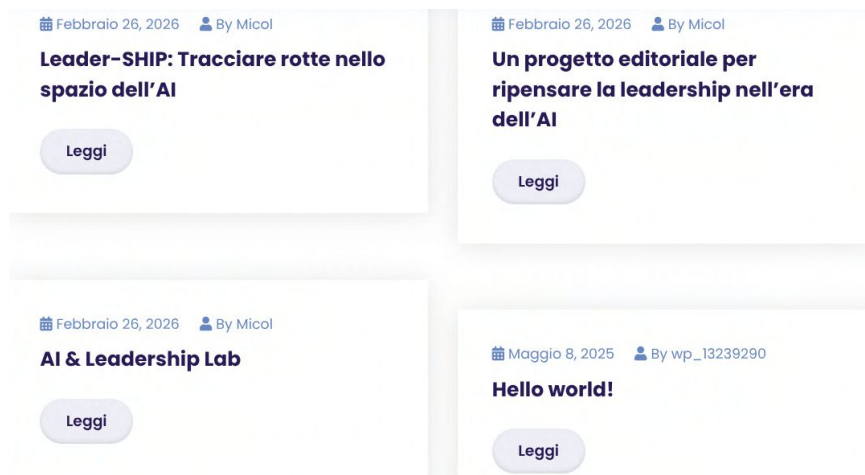


Fig. 3.3.1: I primi articoli del Diario di bordo di Lead-Alliance.

FONTE: sito ufficiale di Lead-Alliance.

La costruzione del sito ha inoltre prodotto un effetto trasformativo interno: la necessità di rendere pubblici missione, visione e valori ha imposto un livello più elevato di formalizzazione concettuale. Ciò conferma come le infrastrutture digitali non siano meri contenitori neutri, ma dispositivi che influenzano la configurazione organizzativa stessa.

Strettamente connessa al sito è la scelta di presidiare esclusivamente *LinkedIn* come canale *social*. L'esclusione di altre piattaforme non deriva da una limitazione operativa, ma da una decisione strategica: il *target* di riferimento è composto principalmente da professionisti, *manager*, consulenti e organizzazioni. *LinkedIn* viene quindi concepito come estensione relazionale del sito, spazio di distribuzione dei contenuti e canale di legittimazione professionale coerente con il posizionamento del progetto.



Fig. 3.3.1: Lead-Alliance su LinkedIn.

FONTE: *LinkedIn ufficiale di Lead-Alliance.*

La creazione congiunta di sito e pagina *LinkedIn* rende evidente un ulteriore passaggio evolutivo: la comunicazione non può più essere episodica o spontanea, ma richiede una pianificazione strutturata. La necessità di definire un piano editoriale — sia per il *blog* interno sia per *LinkedIn* — emerge come impegno organizzativo futuro, non soltanto comunicativo. La continuità dei contenuti diventa infatti condizione di credibilità.

### 3.3.2 Il podcast come dispositivo divulgativo e matrice simbolica

Il *podcast*, dal nome *Leader-SHIP*, rappresenta la prima idea di *format* strutturato attraverso cui *Lead-Alliance* traduce la propria identità in produzione concreta. Nato da un *brainstorming* spontaneo, il progetto prende forma attorno a una definizione che ne chiarisce fin da subito l'orientamento: non divulgazione tecnica, ma navigazione concettuale. L'obiettivo non è spiegare l'intelligenza artificiale in termini operativi, bensì attraversarla come ambiente culturale, mettendola in relazione con *leadership*, responsabilità e trasformazioni generazionali.

Il nome stesso — *Leader-SHIP* — attiva una stratificazione simbolica. Dal termine *leadership* emerge l'idea della *ship*, inizialmente evocata come nave o un galeone e successivamente evoluta in navicella, coerentemente con un immaginario più tecnologico e in linea con il tema dell'IA. La metafora non nasce come elemento decorativo, ma come dispositivo interpretativo: se la *leadership* tradizionale presuppone stabilità e controllo, l'immaginario dell'esplorazione suggerisce incertezza, rotta variabile e necessità di orientamento continuo. Il lessico della navigazione — rotte, avamposti, equipaggio — diventa così un modo per raccontare il cambiamento non come minaccia, ma come territorio da esplorare.

La struttura operativa del *podcast* riflette questa impostazione. Coordinato dalla sottoscritta, il *team* funziona secondo una *leadership* condivisa: le decisioni emergono dal confronto collettivo e vengono successivamente sintetizzate, senza che si produca una distinzione marcata tra *leadership* formale e *leadership* percepita. Questa dinamica rispecchia, a livello micro, la concezione di *leadership* distribuita che il progetto stesso intende promuovere.

Sul piano editoriale, il formato privilegia episodi di durata contenuta (circa 15–20 minuti), con la presenza di almeno un ospite per puntata e una predilezione per la singola *take*. Tale scelta segnala un orientamento alla spontaneità e alla conversazione autentica. La selezione degli ospiti avviene principalmente attraverso la rete personale, con criterio esplicito di coerenza tematica. In questa fase iniziale, la fiducia relazionale precede l'ampliamento sistematico del *network*, configurando il *podcast* come spazio di divulgazione qualificata più che come prodotto mediatico ad ampia diffusione.

La programmazione della prima stagione — articolata in episodi dedicati al cambiamento generazionale, ai regolamenti europei, al tema educativo, alla nascita delle IA e al confronto intergenerazionale — evidenzia un tentativo di coprire ambiti differenti dello stesso fenomeno.

Tuttavia, la distanza tra pianificazione e produzione effettiva mette in luce una tensione tipica delle organizzazioni emergenti: il carico di lavoro e le difficoltà tecniche rallentano l'esecuzione rispetto alla progettazione iniziale. Tale discrepanza non segnala incoerenza strategica, ma un problema di sostenibilità operativa, già osservato nella dissoluzione del *team* eventi. Anche in questo caso, emerge la necessità di calibrare ambizione e capacità esecutiva.

Un elemento particolarmente significativo è l'ipotesi di parlare attraverso l'IA, ad esempio utilizzando strumenti come i *podcast* di *NotebookLM* in un episodio della prima stagione, dedicata prevalentemente al tema del cambiamento generazionale nelle aziende e di come l'IA possa incidere sulla ridefinizione del senso organizzativo.

L'idea nasce come esperimento e provocazione: non limitarsi a discutere dell'IA, ma integrarla direttamente nel processo comunicativo. In tal modo, il *podcast* si configura non solo come spazio di divulgazione, ma come laboratorio applicato che mette alla prova, in tempo reale, il rapporto tra voce umana e mediazione algoritmica.

Anche la proposta di un logo raffigurante Saturno con un anello frammentato si inserisce in questa logica simbolica. L'anello spezzato non rappresenta soltanto un segno grafico, ma la rottura di un ordine consolidato: l'intelligenza artificiale come forza che interrompe equilibri precedenti e costringe a ridefinire coordinate e confini. L'estetica fantascientifica non è quindi un semplice elemento creativo, ma una traduzione visiva della tesi implicita del progetto: la *leadership* contemporanea opera in un contesto di discontinuità sistemica.

Al momento della stesura del presente lavoro, nessun episodio risulta ancora pubblicato. La fase attuale è di pre-produzione: definizione del *vademecum* per gli ospiti (in corso di elaborazione), formalizzazione degli inviti, prove tecniche di registrazione a distanza e pianificazione del montaggio tramite strumenti quali *Audacity*<sup>26</sup>. La pubblicazione è prevista su piattaforme come *Spotify*, con integrazione di *plug-in* dedicato all'interno del sito e successiva promozione su *LinkedIn*. La scelta organizzativa prevede la registrazione preliminare dell'intera stagione e la pubblicazione successiva con cadenza regolare, in modo da garantire continuità e affidabilità nel tempo.

La distanza tra pianificazione e *output* pubblicato non è tuttavia riconducibile a indecisione strategica, ma a carico di lavoro e complessità tecnica, elementi tipici dei progetti emergenti a gestione volontaria. In tal senso, il riferimento a modelli già sperimentati — come il podcast *reJeneration*<sup>27</sup> promosso da *JECOMM*, al quale ho personalmente collaborato come *speaker*, o il progetto *IA Spiegata Semplice*<sup>28</sup> per quanto riguarda la strutturazione editoriale integrata — rappresenta non un tentativo imitativo, ma un apprendimento organizzativo per analogia. *Lead-Alliance*, tuttavia, si differenzia da potenziali *competitors* per la verticalizzazione tematica su *leadership* e IA, evitando una trattazione esclusivamente tecnologica e privilegiando una prospettiva relazionale e organizzativa.

Al di là della funzione divulgativa, il *podcast* persegue inoltre due obiettivi impliciti ma centrali: costruire una *community* dialogica e attrarre collaboratori. Il *format* conversazionale favorisce l'identificazione e l'interazione, mentre la scelta di coinvolgere ospiti eterogenei amplia progressivamente il perimetro relazionale. In questo senso, il *podcast* non è soltanto un canale di comunicazione, ma un dispositivo di attivazione: attraverso la voce, il progetto rende visibile la propria postura culturale e invita altri soggetti a salire a bordo dell'equipaggio.

---

<sup>26</sup> Cfr. <https://www.audacityteam.org/>

<sup>27</sup> Cfr. <https://open.spotify.com/show/580VYJFsw995wryo2OTOp>

<sup>28</sup> Cfr. <https://www.iaspiegatasemplice.it/>

### 3.3.3 La newsletter come protocollo sperimentale e dispositivo metodologico

A differenza del *podcast*, che privilegia la dimensione conversazionale e relazionale, la *newsletter AI & Leadership Lab* nasce come dispositivo metodologico esplicito. Non si configura primariamente come strumento di aggiornamento, bensì come esperimento dichiarato di osservazione del rapporto tra intelligenza artificiale e *leadership*.

Fin dalla dichiarazione iniziale, la *newsletter* assume una postura radicalmente trasparente: il testo è generato da un'intelligenza artificiale e tale condizione non viene nascosta, ma esplicitata come scelta metodologica. In un contesto in cui l'uso dell'IA nella produzione di contenuti professionali è diffuso ma raramente dichiarato, la trasparenza diventa non soltanto opzione etica, ma presa di posizione epistemologica. Rendere visibile il processo significa spostare l'attenzione dal prodotto finale al meccanismo che lo genera.

L'obiettivo primario non è costruire una *mailing list* né consolidare un'abitudine di lettura, ma testare come l'IA scrive di *leadership*. La domanda di fondo non riguarda la brillantezza del sistema, bensì la sua utilità pratica: un'intelligenza artificiale può contribuire in modo significativo alla riflessione di chi guida persone e prende decisioni? La *newsletter* si presenta quindi come unità sperimentale autonoma: ogni numero nasce in una chat dedicata, con *prompt*<sup>29</sup> espliciti e dichiarati, e con un intervento umano limitato e osservativo.

La struttura è volutamente fluida e ancora in fase di definizione. La cadenza di pubblicazione non è rigidamente stabilita, anche perché in parte dipende dall'avanzamento degli altri *team*. Tuttavia, l'impianto tende a combinare approfondimento tematico e *call to action*, invitando i lettori a fornire *feedback* critico. In questo senso, la *newsletter* assume una doppia funzione: autonoma come laboratorio di scrittura algoritmica e complementare rispetto agli altri *format*, fungendo anche da canale di aggiornamento per iniziative come il *podcast*.

Il processo di revisione è collettivo. Il testo generato dall'IA viene discusso dal *team*, commentato e analizzato, ma non riscritto in modo sostanziale. L'errore, la semplificazione e il *bias* non sono corretti per mascherarne l'origine, bensì considerati

---

<sup>29</sup> È importante sottolineare, inoltre, che, in fase di addestramento dell'IA, è stato stabilito che, al termine delle 12 *newsletter* previste nella sperimentazione, la macchina produrrà un *report* su come è andato il progetto, cosa ha funzionato e cosa no. Questo documento verrà poi messo a confronto con il *report* prodotto dalle persone all'interno del *team*.

parte del materiale di ricerca. Questa scelta distingue il progetto da una semplice pratica di *co-writing* umano–macchina: qui l'IA non è assistente invisibile, ma oggetto di osservazione.

L'infrastruttura tecnica è già predisposta: *landing page* e modulo di iscrizione risultano integrati nel sito, sebbene non ancora attivi in attesa della pubblicazione ufficiale e della definizione delle procedure di gestione della *privacy* e dei dati personali. Tale attenzione alla dimensione normativa si inserisce coerentemente nel quadro più ampio di *trustworthy AI communication* discusso precedentemente, ribadendo che sperimentazione e responsabilità non sono in contraddizione.

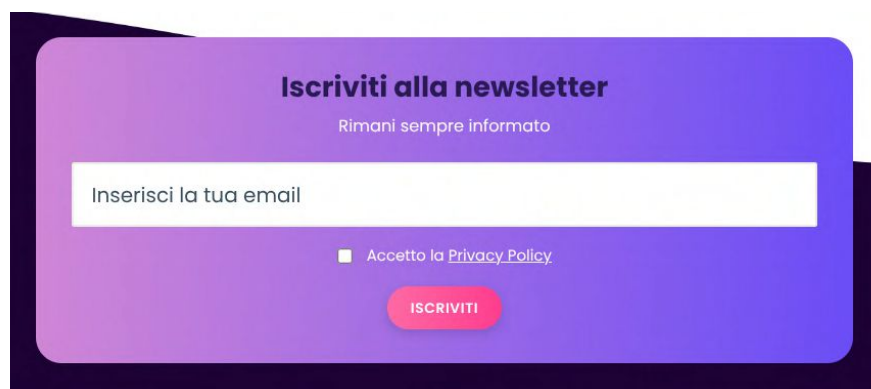
The image shows a digital form for newsletter registration. It has a purple-to-blue gradient background. At the top, the text 'Iscriviti alla newsletter' is in bold white font, followed by 'Rimani sempre informato' in a smaller white font. Below this is a white rectangular input field with the placeholder text 'Inserisci la tua email'. Under the input field is a checkbox followed by the text 'Accetto la Privacy Policy', where 'Privacy Policy' is a blue link. At the bottom center is a red rounded rectangular button with the white text 'ISCRIVITI'.

Fig. 3.3.3: Modulo di iscrizione alla newsletter.

FONTE: sito ufficiale di Lead-Alliance.

Se il *podcast* attiva la dimensione relazionale, la *newsletter* costruisce una comunità epistemica. Chi si iscrive non è soltanto destinatario di contenuti, ma osservatore critico chiamato a valutare pertinenza, applicabilità e limiti del testo generato. Il *feedback* dei lettori diventa parte integrante del protocollo, configurando la *newsletter* come laboratorio aperto più che come prodotto editoriale.

In questa prospettiva, *AI & Leadership Lab* non si limita a comunicare sull'intelligenza artificiale: la mette alla prova pubblicamente. La trasparenza radicale non è un elemento retorico, ma una modalità di costruzione di fiducia fondata sulla visibilità del processo.

### 3.3.4 Dal *whitepaper* al libro: formalizzazione e capitalizzazione del sapere

Il progetto in questione nasce originariamente con l'idea di produrre un *whitepaper*, inteso come documento strutturato di sintesi teorica e operativa sul rapporto tra intelligenza artificiale e *leadership*. Nel corso della progettazione, tuttavia, questa impostazione si è progressivamente evoluta verso l'ipotesi di un libro, segnando un passaggio rilevante nella traiettoria del progetto.

Il passaggio non è, infatti, solo una formalità. Il *whitepaper* rispondeva a una logica di posizionamento e legittimazione iniziale: un documento agile, autorevole, capace di sintetizzare visione, *framework* e primi *output*. L'ipotesi del libro introduce invece una diversa ambizione: non soltanto dichiarare una posizione, ma contribuire in modo sistematico al dibattito scientifico e manageriale.

L'obiettivo è una pubblicazione concepita anche come prodotto editoriale vendibile, capace di trasformare il capitale simbolico accumulato in un *asset* riconoscibile e trasferibile. Essa, inoltre, non si configura come manuale tecnico né come saggio accademico tradizionale, bensì come mappa per orientarsi nel cambiamento che l'intelligenza artificiale produce nei contesti di *leadership*.

Il libro prende infatti forma a partire da una domanda fondativa: come cambia la *leadership* quando l'intelligenza artificiale entra realmente nel lavoro delle persone e delle organizzazioni? L'obiettivo non è tuttavia offrire soluzioni prescrittive, ma accompagnare i lettori — *leader*, *manager*, consulenti, formatori e *team* — in un percorso progressivo di consapevolezza e integrazione.

La struttura prevista si articola in cinque parti e un epilogo, organizzati secondo una logica formativa che procede dal cambiamento di prospettiva fino alle applicazioni concrete:

1. *Cambiare sguardo*: ridefinizione dell'*Augmented Leadership* e decisioni *data-driven*,
2. *Il fattore umano*: un *focus* sulla centralità del fattore umano,
3. *Etica, consapevolezza e limiti*,
4. *Il nuovo leader*: nuove competenze del *leader* di oggi e del futuro,
5. *Dalla teoria alla realtà*: applicazioni operative e modelli di *governance*,

6. *Tenere la barra dritta*: riflessione conclusiva sull'identità del *leader* in un mondo “aumentato”.

Questa impostazione evidenzia un tratto distintivo del progetto: la volontà di tenere insieme riflessione teorica, dimensione etica e applicazione concreta. Il libro viene concepito come esperienza formativa strutturata, non come raccolta episodica di contributi.

Sebbene l'impianto sia unitario, la redazione coinvolge tuttavia membri diversi del gruppo, ciascuno prevalentemente sulle aree di propria competenza. Si configura quindi una dinamica ibrida: coerenza narrativa centrale e contributo interdisciplinare distribuito. L'unitarietà dell'opera non deriva dall'omogeneità delle voci, ma dalla convergenza su una domanda comune.

Dal punto di vista organizzativo, il passaggio dal confronto laboratoriale alla progettazione di un'opera editoriale rappresenta un momento di maturazione. Le discussioni, le sperimentazioni e le riflessioni sviluppate vengono progressivamente trasformate in una struttura argomentativa stabile. Il libro diventa così dispositivo di formalizzazione del percorso di ricerca.

In questo senso, il progetto editoriale non è solo *output* comunicativo, ma strumento di consolidamento identitario. Se il sito rende visibile il progetto e il *podcast* ne esplora le rotte conversazionali, il libro ne sistematizza la visione e ne fissa le coordinate.

### **3.4 Prospettive di sviluppo**

Dopo aver analizzato la struttura organizzativa e le sperimentazioni editoriali di *Lead-Alliance*, si rende necessario interrogarsi sulle sue prospettive evolutive. Ogni progetto in fase embrionale attraversa infatti una soglia critica: il passaggio dall'entusiasmo fondativo alla sostenibilità nel tempo. Nel caso di *Lead-Alliance*, tale soglia non coincide soltanto con la produzione di contenuti o con la pubblicazione del sito, ma con una questione più profonda: quale forma organizzativa può garantire continuità a un laboratorio culturale nato in modo volontario e interdisciplinare?

Le tensioni che emergono in questa fase non sono conflittuali, bensì strutturali. Da un lato, il progetto aspira a diventare un *brand* editoriale riconoscibile e, in prospettiva, una realtà formativa capace di offrire modelli applicativi sull'*Augmented Leadership*.

Dall'altro, la natura non formalizzata dell'organizzazione e la pluralità di impegni professionali dei membri impongono limiti operativi che rendono evidente il rischio di dispersione.

In questa sezione verranno analizzate le principali traiettorie di sviluppo possibili, considerando sia le dinamiche interne di coinvolgimento e governance, sia le opzioni strategiche future — dalla configurazione come *community* stabile fino all'ipotesi di una struttura più formalizzata, eventualmente connessa a realtà già esistenti.

L'obiettivo non è formulare previsioni, ma evidenziare le tensioni evolutive che caratterizzano i progetti collaborativi ad alta intensità valoriale quando si confrontano con la necessità di stabilizzazione organizzativa.

### **3.4.1 Tra consolidamento e rischio di dispersione**

A distanza di alcuni mesi dall'avvio della seconda fase, *Lead-Alliance* si presenta come un'organizzazione in via di consolidamento selettivo. Il nucleo attivo — in particolare i membri impegnati nei *team* editoriali e nel senato ristretto — appare più coeso rispetto alla fase iniziale. Al contempo, si osserva un progressivo allontanamento di alcuni partecipanti meno coinvolti operativamente.

Questa dinamica non è anomala nei progetti volontari o semi-strutturati, ma segnala una tensione evolutiva: il passaggio da entusiasmo fondativo a sostenibilità nel tempo. Il rischio principale non è il conflitto né la divergenza valoriale, bensì la dispersione. La pluralità di impegni professionali dei membri, unita alla natura non ancora formalizzata del progetto, espone l'organizzazione al rischio di rallentamenti e interruzioni cicliche.

In questo contesto si colloca la decisione di introdurre un breve questionario anonimo di riallineamento, volto a misurare il *sentiment* interno rispetto al carico di lavoro, al coinvolgimento e alla percezione del progetto. Lo strumento non è stato concepito soltanto come rilevazione descrittiva, ma come dispositivo preventivo: intercettare eventuali segnali di disaffezione prima che si traducano in abbandono silenzioso.

Il fenomeno del *disengagement* progressivo, talvolta descritto nei report HR contemporanei come *Great Detachment*, evidenzia come la perdita di coinvolgimento possa manifestarsi in modo graduale e non conflittuale, attraverso una riduzione dell'energia investita piuttosto che tramite un'esplicita rottura. In questo senso, il

questionario assume una funzione triplice: diagnostica, preventiva e simbolica. Diagnostica, perché consente di raccogliere dati sul clima interno; preventiva, perché mira a intervenire prima che la dispersione diventi strutturale; simbolica, perché comunica attenzione e cura verso il gruppo.

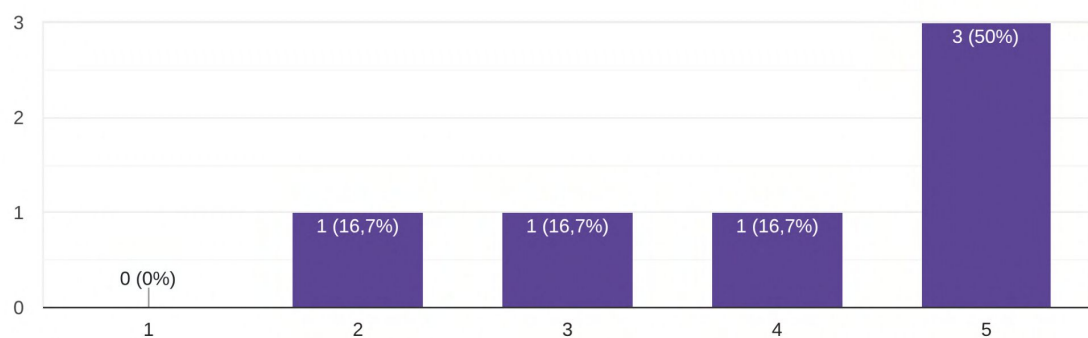
Il questionario, somministrato agli attuali 13 membri attivi del progetto, ha raccolto solamente 6 risposte (46% del totale)<sup>30</sup>. Pur trattandosi di un campione numericamente contenuto, la rilevazione assume valore esplorativo e qualitativo, coerente con la natura partecipativa della ricerca. La non-risposta, in questo senso, costituisce essa stessa un indicatore di differenziazione nel livello di *engagement*.

Dall'analisi dei dati emergono tre elementi principali.

In primo luogo, il livello di coinvolgimento dichiarato risulta mediamente medio-alto, con una concentrazione delle risposte nelle fasce superiori della scala. I membri che partecipano attivamente ai *team* editoriali e alle decisioni strategiche percepiscono il proprio contributo come significativo e si dichiarano motivati a proseguire la collaborazione. Questo dato conferma l'esistenza di un nucleo coeso, caratterizzato da un senso di appartenenza e da una motivazione intrinseca legata sia alla crescita personale sia alla coerenza valoriale del progetto.

Quanto ti senti coinvolto/a nelle decisioni e negli obiettivi strategici di Lead-Alliance?

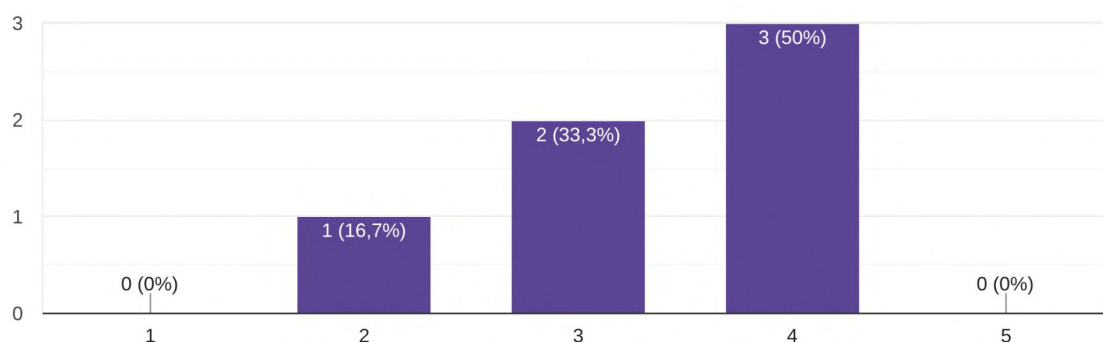
6 risposte



<sup>30</sup> I dati sono stati raccolti tramite *Google Form* (<https://forms.gle/WBFtiJnHdgAkvcWS7>), pertanto i grafici presentati in questa sezione sono stati generati automaticamente in base alle risposte date dai partecipanti.

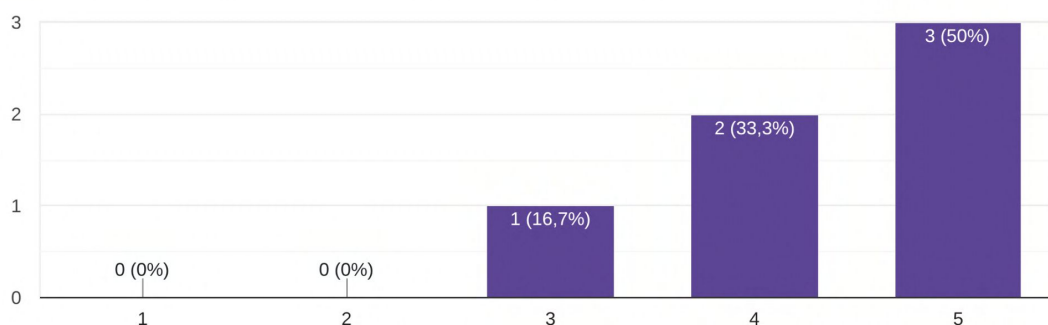
Quanto senti che il tuo contributo stia facendo la differenza per il progetto?

6 risposte



Qual è il tuo livello di motivazione attuale nel collaborare con noi?

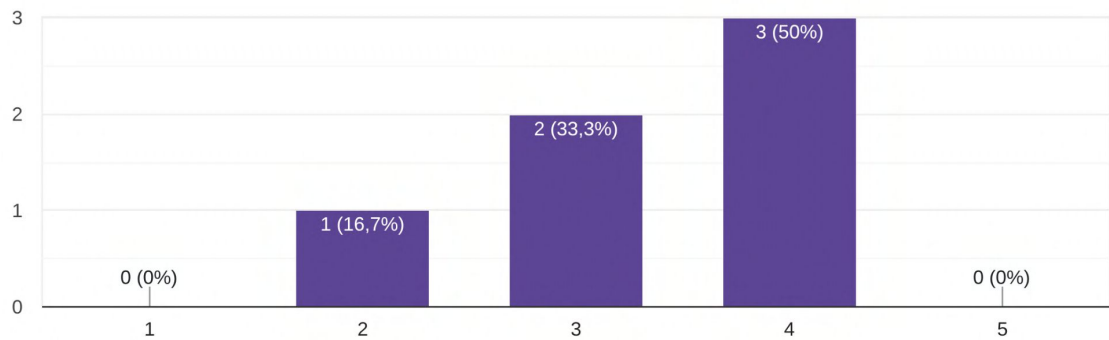
6 risposte



In secondo luogo, emerge una percezione complessivamente positiva del clima umano e della collaborazione interna. Le valutazioni relative alla chiarezza comunicativa e al supporto ricevuto si collocano prevalentemente nelle fasce medio-alte, segnalando un ambiente relazionale collaborativo e non conflittuale. Il rischio, dunque, non appare legato a tensioni interpersonali o a divergenze valoriali, bensì a fattori strutturali.

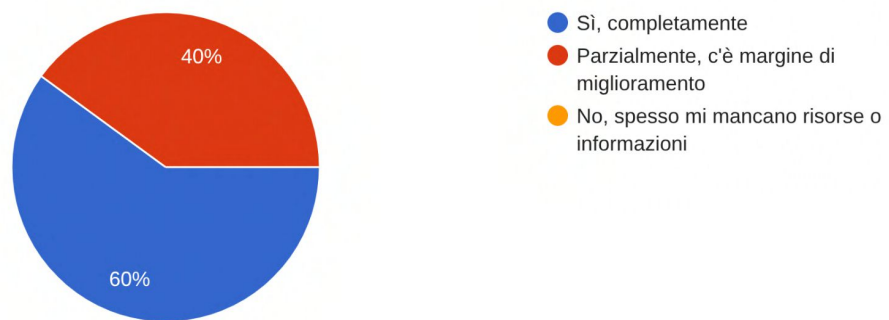
Quanto ritieni chiara la comunicazione interna tra i diversi team e con il 'Senato'?

6 risposte



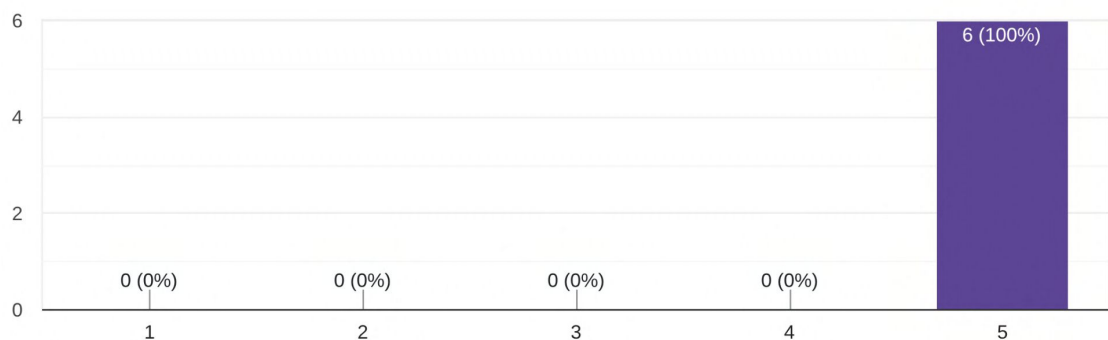
Senti di avere tutto il supporto (strumenti, informazioni, mentoring) necessario per svolgere i tuoi task in modo efficace?

5 risposte



Come valuti il clima umano e la collaborazione all'interno del tuo team specifico?

6 risposte

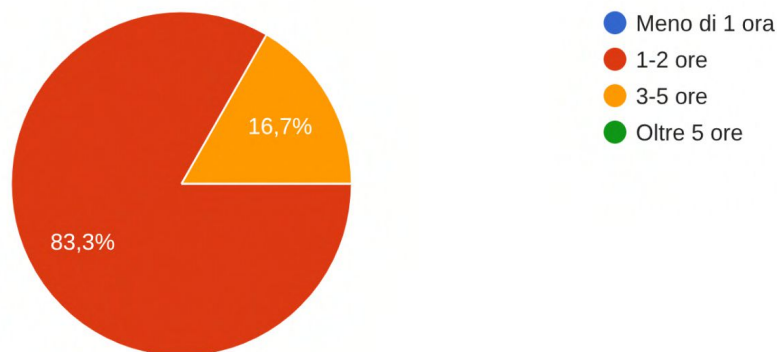


È proprio su questo piano che si manifesta la terza evidenza: la sostenibilità nel tempo. Alcune risposte evidenziano come il principale elemento critico non sia il

disaccordo, ma la dispersione. La pluralità di impegni professionali, l'assenza di vincoli formali e la natura volontaria dell'iniziativa rendono il progetto vulnerabile a rallentamenti progressivi. Il dato più significativo, in questo senso, non è la presenza di insoddisfazione esplicita, bensì la possibile riduzione graduale dell'energia investita.

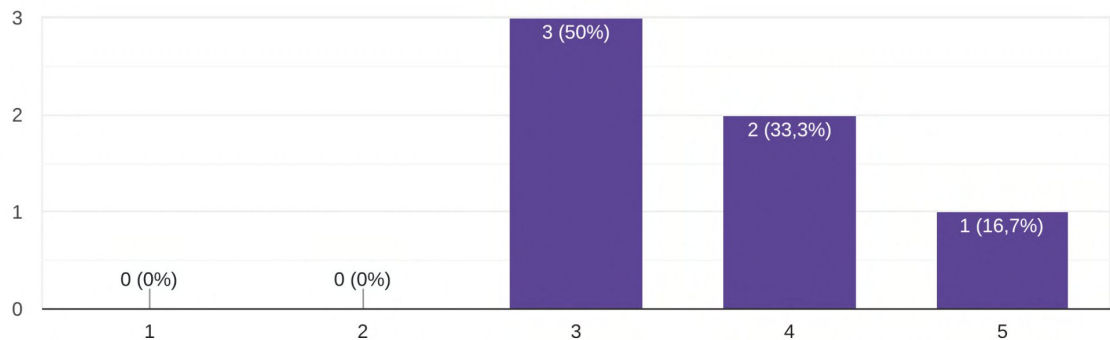
Quanto tempo riesci a dedicare mediamente al progetto a settimana?

6 risposte



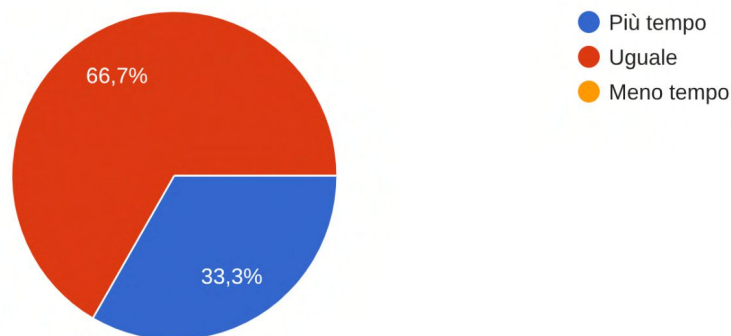
Il carico di lavoro attuale è sostenibile per te?

6 risposte



In futuro, vorresti dedicare più o meno tempo al progetto?

6 risposte



Il quadro complessivo restituisce quindi un'organizzazione che si sta consolidando intorno a un *core team* motivato e collaborativo, ma che deve affrontare una sfida tipica delle realtà emergenti: trasformare l'entusiasmo fondativo in struttura sostenibile.

In questa prospettiva, il questionario non rappresenta soltanto uno strumento di monitoraggio, ma un primo gesto di governance. L'atto stesso di chiedere *feedback* formalizza implicitamente un passaggio: da gruppo informale a organizzazione che riflette su se stessa.

La prospettiva di sviluppo di *Lead-Alliance* dipende quindi non soltanto dalla qualità dei contenuti prodotti, ma dalla capacità di mantenere un equilibrio tra visione e sostenibilità. La solidità del nucleo attivo costituisce un elemento di stabilizzazione; tuttavia, l'assenza di una formalizzazione giuridica e di un impegno temporale definito rende il progetto ancora esposto a oscillazioni. La questione evolutiva diventa allora strategica: *Lead-Alliance* rimarrà un laboratorio fluido, alimentato dalla spinta motivazionale dei suoi membri più attivi, oppure intraprenderà un percorso di maggiore formalizzazione, capace di distribuire responsabilità, tempi e carichi in modo più stabile?

### 3.4.2 Formalizzazione e stabilità: l'ipotesi di integrazione strutturale

Tra le prospettive evolutive emerse nel confronto interno, assume particolare rilievo l'ipotesi di un'integrazione formale con *OpenHS*, organizzazione già strutturata e composta in larga parte dagli stessi membri attivi nel senato ristretto. In questo scenario,

l'eventuale inserimento di *Lead-Alliance* sotto una struttura giuridicamente definita non verrebbe percepito come perdita di autonomia, ma come acquisizione di stabilità. La sovrapposizione significativa tra i membri delle due realtà riduce infatti il rischio di eterodirezione e rende l'integrazione più simile a un consolidamento organizzativo che a una subordinazione.

La questione non riguarda tanto il controllo decisionale, quanto la sostenibilità infrastrutturale. L'assenza di una forma giuridica autonoma espone *Lead-Alliance* a fragilità operative: gestione delle *policy*, eventuale monetizzazione dei contenuti, investimenti economici, distribuzione delle responsabilità legali. L'integrazione con una struttura già esistente potrebbe offrire un contenitore normativo e amministrativo in grado di sostenere l'evoluzione del progetto senza snaturarne la matrice culturale.

Tale passaggio segnerebbe tuttavia una trasformazione identitaria rilevante. La dimensione volontaria e laboratoriale, che oggi costituisce uno degli elementi distintivi del progetto, dovrebbe confrontarsi con una maggiore formalizzazione di ruoli, tempi e *commitment*. La stabilità organizzativa richiede infatti una definizione più esplicita delle responsabilità individuali e degli investimenti di risorse. In questa prospettiva, l'integrazione con *OpenHS* non appare come alternativa alla visione originaria, ma come possibile fase successiva di maturazione. La domanda cruciale non è se formalizzare, ma eventualmente quando e a quali condizioni.

In generale, se si adotta una prospettiva realistica sui prossimi dodici mesi, il modello più sostenibile appare quello di una *community* stabile non ancora orientata alla monetizzazione. L'attuale configurazione organizzativa, fondata su impegno volontario e distribuzione fluida dei ruoli, consente al progetto di mantenere elasticità e sperimentazione senza introdurre pressioni economiche premature.

L'aspirazione evolutiva, tuttavia, è dichiaratamente più ambiziosa: trasformare *Lead-Alliance* in una *startup* formativa capace di offrire percorsi, contenuti strutturati e servizi legati all'intersezione tra *leadership* e intelligenza artificiale. Questa proiezione introduce una tensione temporale tra ciò che è sostenibile oggi e ciò che è desiderabile domani.

La letteratura sull'evoluzione delle organizzazioni emergenti evidenzia come le realtà nascenti attraversino fasi successive di consolidamento prima di giungere a una piena formalizzazione strutturale. Greiner (1972), nella sua teoria delle fasi di crescita

organizzativa, descrive infatti come i progetti inizialmente fondati sull'entusiasmo e sulla creatività collettiva siano destinati a confrontarsi con momenti di crisi evolutiva legati alla *leadership*, alla definizione dei ruoli e al controllo. Analogamente, Stinchcombe (1965) ha evidenziato la cosiddetta *liability of newness*, ossia la vulnerabilità intrinseca delle organizzazioni giovani, esposte a fragilità operative e carenza di legittimazione istituzionale. In una prospettiva istituzionale, Tolbert e Zucker (1996) mostrano infine come i processi di formalizzazione non siano immediati, ma si sviluppino progressivamente attraverso fasi di sperimentazione, diffusione e sedimentazione. In questo quadro teorico, anticipare la monetizzazione in assenza di processi consolidati potrebbe generare tensioni premature; al contrario, una fase di consolidamento comunitario può rappresentare un investimento in stabilità futura.

Nel caso di *Lead-Alliance*, la scelta strategica appare orientata a una gradualità intenzionale: consolidare identità, rete e reputazione prima di attivare modelli economici strutturati. In questa fase, la funzione principale del progetto rimane la produzione di capitale relazionale, simbolico e culturale. La trasformazione in *startup* formativa rappresenta dunque non un'immediata necessità operativa, ma un orizzonte di sviluppo subordinato alla stabilizzazione organizzativa. In termini di *governance*, ciò implica un equilibrio delicato: preservare la dimensione laboratoriale e volontaria che alimenta creatività e senso di appartenenza, mentre si costruiscono progressivamente le condizioni infrastrutturali per un eventuale salto di scala.

In questa prospettiva, l'eventuale integrazione con *OpenHS* non costituirebbe il punto di arrivo dell'evoluzione di *Lead-Alliance*, bensì una fase di stabilizzazione intermedia. L'obiettivo di lungo periodo rimane infatti la costruzione di una realtà autonoma, capace di configurarsi come *startup* formativa specializzata nell'intersezione tra *leadership* e intelligenza artificiale.

La scelta di non attivare immediatamente un modello di monetizzazione non equivale dunque a una rinuncia strategica, ma a una sospensione consapevole. Consolidare comunità, identità e processi prima di strutturare un'offerta economica consente di ridurre il rischio di formalizzazioni premature e di preservare coerenza culturale.

Nel breve termine, la sostenibilità si fonda sulla stabilità del nucleo attivo e sulla progressiva definizione di ruoli e responsabilità. Nel medio-lungo periodo, la trasformazione in entità autonoma richiederà invece un salto di scala: definizione di

modello di *business*, assetto giuridico indipendente, struttura operativa dedicata e distribuzione formale delle responsabilità economiche e legali.

La traiettoria evolutiva di *Lead-Alliance* può quindi essere descritta come tripartita: fase laboratoriale, fase di consolidamento infrastrutturale, eventuale fase imprenditoriale autonoma. La questione strategica non riguarda la direzione, già delineata, ma il ritmo e le condizioni di attivazione del passaggio successivo.

In generale, le prospettive di sviluppo di *Lead-Alliance* non possono essere lette esclusivamente come scelte organizzative, ma come esercizio concreto di *leadership* in contesti tecnologicamente mediati. La costruzione di una *community*, la riflessione sulla formalizzazione giuridica, l'eventuale integrazione infrastrutturale e l'aspirazione a una futura autonomia imprenditoriale rappresentano momenti differenti di uno stesso processo: apprendere a governare l'incertezza senza irrigidire prematuramente la struttura.

In questo senso, *Lead-Alliance* si configura non soltanto come oggetto di analisi, ma come laboratorio vivente di co-evoluzione tra visione, tecnologia e *governance*. Le decisioni assunte — dalla gestione del rischio di dispersione alla scelta di gradualità nella monetizzazione — mostrano come la *leadership* nell'era dell'intelligenza artificiale non consista nell'accelerare indiscriminatamente, ma nel modulare tempi, responsabilità e formalizzazione in funzione della sostenibilità.

Il passaggio alla dimensione pubblica, la misurazione del *sentiment* interno e la riflessione sulle traiettorie organizzative segnano quindi una maturazione progressiva: da iniziativa sperimentale a realtà consapevole delle proprie dinamiche strutturali. Proprio in questa tensione tra laboratorio e istituzionalizzazione si colloca il contributo più rilevante dell'esperienza analizzata: dimostrare che la *governance* dell'IA nelle organizzazioni non è soltanto questione di strumenti, ma di scelte collettive, ritmo evolutivo e responsabilità condivisa.

## CAPITOLO 4: IL FUTURO DELLE IA NELLE ORGANIZZAZIONI

Negli ultimi anni l'intelligenza artificiale è entrata nelle organizzazioni con una promessa tanto seducente quanto ambigua: aumentare l'efficienza, accelerare i processi decisionali, ridurre i costi e potenziare la capacità di governare complessità crescenti. Tuttavia, questa rappresentazione rischia di semplificare un fenomeno profondamente trasformativo sul piano culturale, relazionale e strategico. Come già analizzato nei capitoli precedenti, l'adozione dell'IA non coincide con un semplice *upgrade* tecnologico, ma ridisegna modelli di lavoro, responsabilità e pratiche di coordinamento.

L'esperienza di *Lead-Alliance*, osservata attraverso un approccio partecipato e qualitativo, ha mostrato come un gruppo interdisciplinare tenti di conciliare innovazione e prudenza, sperimentazione e centralità umana. Quel caso costituisce lo sfondo interpretativo da cui questo capitolo muove, spostando però l'attenzione su un problema diverso: come valutare l'efficacia dell'IA nei contesti organizzativi?

In questa prospettiva, il "futuro" delle IA nelle organizzazioni non viene inteso come previsione tecnologica o come anticipazione di scenari, ma come traiettoria metodologica e culturale. Il modo in cui l'intelligenza artificiale verrà integrata, governata e valutata nei prossimi anni dipenderà infatti dai criteri attraverso cui ne definiamo oggi l'efficacia. Il futuro non è dunque soltanto questione di potenza computazionale o di innovazione tecnica, ma di cornici interpretative, di scelte organizzative e di responsabilità decisionali.

Ogni promessa di miglioramento implica infatti un criterio di valutazione. Sostenere che un sistema decisionale sia "più efficace" significa di conseguenza definire cosa costituisca una buona decisione. Se la valutazione di *output* quantitativi può avvalersi di indicatori relativamente stabili, il discorso si complica quando l'IA interviene in ambiti che coinvolgono dinamiche collaborative, interpretazione del contesto, costruzione della fiducia e qualità del giudizio. In questi casi, l'efficacia non si esaurisce nella velocità o nell'accuratezza statistica, ma riguarda la coerenza tra strumenti tecnologici, obiettivi organizzativi e valori condivisi.

Questa tensione emerge quando l'attenzione si sposta dalle *hard skills* alle *soft skills*: competenze quali creatività, pensiero critico, comunicazione e capacità di gestire l'ambiguità. Proprio queste competenze vengono oggi indicate come decisive per

lavorare efficacemente con sistemi di IA, soprattutto in modelli organizzativi ibridi dove la decisione risulta dall'interazione tra suggerimenti algoritmici e valutazione umana. Tuttavia, la loro misurazione pone problemi di validità, interpretazione e uso organizzativo dei risultati.

L'obiettivo di questo capitolo è proporre una riflessione metodologica esplorativa su criteri, limiti e ambiguità della valutazione dell'efficacia dell'IA. Non si intende costruire un sistema definitivo di metriche, ma problematizzare l'idea stessa di misurazione, mostrando come il concetto di efficacia cambi a seconda degli scopi, del contesto e delle dimensioni osservate. L'efficacia viene qui intesa non come mero incremento di *performance*, ma come equilibrio tra automazione e responsabilità, tra supporto algoritmico e capacità critica umana.

Per rendere questa riflessione più concreta, il capitolo include un *case study* dedicato a *Slash*, una piattaforma sviluppata da *Altermaind* che propone un percorso di *assessment* e formazione per l'adozione dell'IA. L'attenzione si concentra sul *soft skills assessment* di *Slash*, progettato per mappare competenze strategiche in ambienti caratterizzati dall'integrazione uomo-macchina. Il caso rappresenta un terreno privilegiato per osservare le sfide metodologiche della misurazione di dimensioni immateriali e i rischi di un uso rigido dei risultati nell'ambito delle Risorse Umane (HR).

La scelta di includere *Slash* risponde a una logica di complementarità analitica. Se *Lead-Alliance* mostra come si costruisce senso attorno all'IA nei processi di *leadership*, *Slash* permette di osservare cosa accade quando si tenta di tradurre la complessità dell'integrazione uomo-macchina in strumenti di misurazione e decisioni operative.

La struttura del capitolo segue due movimenti principali. Nella prima parte vengono discussi criteri e limiti dell'efficacia dell'IA, articolandola su tre livelli (tecnico, organizzativo, culturale) e mettendo in evidenza i rischi della misurazione: riduzionismo, opacità, *bias* e *over-reliance*. Nella seconda parte, il *case study* di *Slash* consente di osservare come un progetto di *assessment* tenti di bilanciare l'esigenza di sintesi e complessità del reale, proponendo strumenti che mirano a supportare — mai a sostituire — la valutazione umana.

#### 4.1 L'efficacia dell'IA nei contesti organizzativi: criteri, limiti e ambiguità

Nel dibattito contemporaneo sull'intelligenza artificiale applicata alle organizzazioni, il termine "efficacia" viene frequentemente utilizzato come sinonimo di miglioramento della *performance*. Tuttavia, in ambito manageriale, l'efficacia possiede un significato teorico più preciso e non sovrapponibile alla semplice efficienza operativa.

Drucker distingue in modo netto tra efficienza ed efficacia: la prima consiste nel "fare bene le cose", mentre la seconda nel "fare la cosa giusta" (Drucker, 1966). L'efficacia manageriale riguarda quindi il grado di allineamento tra risultati ottenuti e obiettivi strategici. Non si tratta di massimizzare una prestazione isolata, ma di orientare le scelte verso ciò che è realmente rilevante per l'organizzazione. L'efficienza può essere quindi misurata in termini di ottimizzazione dei processi; l'efficacia implica invece una valutazione di priorità, direzione e coerenza con lo scopo.

Applicata all'intelligenza artificiale, questa distinzione assume un rilievo particolare. Un sistema di IA può essere altamente efficiente sul piano tecnico — riducendo tempi di elaborazione o aumentando l'accuratezza predittiva — senza essere necessariamente efficace in senso manageriale. Se l'algoritmo ottimizza variabili che non corrispondono agli obiettivi strategici, oppure se produce effetti collaterali che compromettono fiducia o responsabilità, l'aumento di efficienza non coincide con un incremento di efficacia.

Trist e Bamforth (1951), nella ricerca *Some social and psychological consequences of the longwall method of coal-getting*, documentano l'esempio di come l'introduzione di una tecnologia finalizzata a massimizzare l'efficienza estrattiva (nello specifico del carbone) abbia inizialmente aumentato la produttività, ma anche generato tensioni sociali, calo della coesione e ridotta *performance* complessiva. L'ottimizzazione tecnica non garantisce quindi automaticamente un miglioramento sistemico. Nel caso dell'IA, il rischio è analogo: concentrarsi sulle metriche computazionali può oscurare l'impatto sui processi organizzativi.

Inoltre, le organizzazioni, come già analizzato, producono senso attraverso processi interpretativi collettivi: le decisioni non emergono da calcoli ottimali, ma dall'elaborazione condivisa di segnali ambigui (Weick, 1995, pp. 49–91). In questo quadro, un sistema di IA tecnicamente performante ma percepito come opaco o incomprensibile dai suoi utilizzatori rischia di essere marginalizzato o aggirato nei

processi reali di *decision-making*, risultando di fatto inefficace sul piano organizzativo, indipendentemente dalla sua accuratezza computazionale.

A ciò si aggiunge una dimensione normativa. Floridi (2019, pp. 185–193) sottolinea come la *governance* dell'IA richieda la traduzione di principi etici in pratiche concrete. L'efficacia non coincide perciò con il raggiungimento di un risultato, ma con la modalità attraverso cui tale risultato viene perseguito. Trasparenza, responsabilità ed equità diventano quindi parte integrante della valutazione.

L'efficacia dell'IA può in definitiva essere compresa come il grado di allineamento tra prestazione tecnica, obiettivi strategici e coerenza normativa. Non è, di conseguenza, una proprietà intrinseca dell'algoritmo, ma una relazione dinamica tra tecnologia, scopi organizzativi e sistema di valori.

#### 4.1.1 I limiti strutturali dell'affidabilità algoritmica: bias, overfitting e drift

L'idea che l'intelligenza artificiale costituisca un sistema neutro e oggettivo rappresenta una delle narrazioni più persistenti nel dibattito pubblico. Tale rappresentazione si fonda spesso sulla errata associazione tra "algoritmo" e "matematica"<sup>31</sup>, presupponendo imparzialità. Tuttavia, l'affidabilità di un modello non dipende esclusivamente dalla sua architettura formale, ma dalla qualità dei dati, dal contesto di applicazione e dal monitoraggio nel tempo.

In questo senso, tre fenomeni risultano particolarmente rilevanti: il *bias*, l'*overfitting* e il *drift*.

Il *bias*, già citato nel primo capitolo, non è un errore accidentale, ma una distorsione sistematica nei risultati. Come evidenziato da Mehrabi et al. (2021), i sistemi di *machine learning* spesso assorbono e riproducono gli squilibri presenti nei dati di addestramento. Se i dati storici riflettono decisioni discriminatorie o distribuzioni non rappresentative, l'algoritmo perpetua tali squilibri sotto una parvenza di neutralità tecnica. A tal proposito, anche gli studi di Barocas e Selbst (2016, pp. 671–732)

---

<sup>31</sup> L'associazione tra questi due termini richiama l'eredità di Muḥammad ibn Mūsā al-Khwārizmī (ca. 780–850), matematico persiano e autore del *Al-Kitāb al-mukhtaṣar fī ḥisāb al-jabr wa-l-muqābala* (da cui deriva il termine algebra). Il suo nome, latinizzato nel XII secolo come Algorismus indicava originariamente regole aritmetiche per calcoli con cifre indo-arabe. Solo nel XX secolo, con Turing e von Neumann, ottenne l'attuale significato che indica una sequenza finita di istruzioni deterministiche per i *computer*. Anche questo passaggio storico sottolinea perciò che gli algoritmi non sono entità neutre, ma costruzioni umane soggette a bias culturali e contestuali.

mostrano come l'uso di *big data* possa produrre effetti discriminatori anche senza intenzionalità esplicita. Per esempio, in ambito organizzativo, un sistema di selezione automatica addestrato su dati sbilanciati tenderà a favorire profili simili a quelli precedentemente selezionati, rafforzando dinamiche di esclusione.

L'*overfitting* si verifica invece quando il modello apprende in modo eccessivamente aderente ai dati di addestramento, memorizzandone non solo le regolarità ma anche il "rumore statistico" (Goodfellow et al., 2016). Come conseguenza, il sistema mostra prestazioni eccellenti nei *test* interni, ma perde capacità di generalizzazione su dati nuovi. In contesti valutativi, un modello affidabile sul campione iniziale può quindi non mantenere la stessa capacità discriminante su popolazioni diverse. L'apparente precisione genera perciò un'illusione di affidabilità che si dissolve in scenari non previsti (Hastie et al., 2009)<sup>32</sup>.

Il *drift* riguarda poi il deterioramento progressivo delle prestazioni dovuto al mutamento del contesto. Nello specifico, Gama et al. (2014) definiscono il *concept drift* come la variazione nel tempo della relazione tra variabili *input* e *output*. In altre parole: ciò che era valido ieri potrebbe non esserlo più oggi. Un modello non aggiornato rischia così di basarsi su logiche obsolete, producendo decisioni progressivamente meno accurate (Lu et al., 2018, pp. 2346–2363). La sua affidabilità è conseguentemente temporanea e condizionata. In particolare, nei sistemi per valutare competenze, il *drift* può alterare completamente la comparabilità nel tempo, mettendo in crisi l'idea stessa di "progresso" misurato<sup>33</sup>.

---

<sup>32</sup> All'estremo opposto rispetto all'*overfitting* si colloca il fenomeno dell'*underfitting*, che si verifica quando il modello risulta eccessivamente semplice rispetto alla complessità del problema e non riesce a catturare le strutture rilevanti presenti nei dati (Goodfellow et al., 2016). In questo caso, l'errore non deriva da una memorizzazione eccessiva del rumore, bensì da una capacità rappresentativa insufficiente: il sistema non apprende né le regolarità profonde né le eccezioni significative, mostrando prestazioni scadenti sia sui dati di addestramento sia su quelli nuovi. Se l'*overfitting* produce un'illusione di precisione, l'*underfitting* genera invece una falsa percezione di stabilità, mascherando l'incapacità del modello di discriminare in modo accurato. In ambito organizzativo, ciò può tradursi in sistemi valutativi eccessivamente generici, incapaci di cogliere differenze sostanziali tra profili o contesti differenti (Hastie et al., 2009).

<sup>33</sup> Una dinamica correlata, ma distinta, è quella del *catastrophic forgetting*, fenomeno tipico dei sistemi di apprendimento neurale che si manifesta quando un modello, aggiornato su nuovi dati, perde improvvisamente la capacità di eseguire correttamente compiti appresi in precedenza (McCloskey & Cohen, 1989). Nel cosiddetto *continual learning*, l'acquisizione di nuove informazioni può interferire con le rappresentazioni precedenti, sovrascrivendole parzialmente o totalmente (Goodfellow et al., 2013). Se il *drift* riguarda il cambiamento del mondo rispetto al modello, il *catastrophic forgetting* riguarda invece il cambiamento del modello rispetto al proprio passato. Nei contesti organizzativi, questo implica che un sistema aggiornato per adattarsi a nuove competenze o metriche potrebbe perdere coerenza con le

Questi tre fenomeni mostrano che l'efficacia di un sistema di IA non è una proprietà intrinseca e definitiva, ma una condizione situata, dinamica e reversibile. Il problema risiede nella natura stessa dell'apprendimento automatico, che dipende da dati sempre parziali e da un mondo che cambia più rapidamente dei modelli.

Da una prospettiva di *governance*, la responsabilità non è delegabile: monitoraggio, aggiornamento e contestualizzazione degli *output* restano compiti organizzativi e umani (Floridi, 2019), rendendo la supervisione umana una condizione operativa strutturale piuttosto che una scelta discrezionale.

#### 4.1.2 I livelli di efficacia

L'efficacia dell'intelligenza artificiale nei contesti organizzativi si distribuisce su tre dimensioni interdipendenti: una tecnica, una organizzativa e una culturale-normativa.

Sul piano tecnico, l'efficacia riguarda la "robustezza" del modello, ossia: qualità e rappresentatività dei dati, capacità di generalizzazione, stabilità delle prestazioni nel tempo. Come mostrato nella sezione precedente, fenomeni quali *bias*, *overfitting* e *drift* evidenziano la fragilità strutturale dell'affidabilità tecnica. A questo livello, l'efficacia coincide quindi con la coerenza tra prestazione del sistema e compito per cui è stato progettato. Tuttavia, l'ottimizzazione tecnica non garantisce mai in modo automatico il buon funzionamento del sistema organizzativo nel suo complesso (Trist & Bamforth, 1951).

Su un piano organizzativo, l'efficacia concerne invece l'integrazione dell'IA nei processi decisionali esistenti. L'introduzione di uno strumento algoritmico ridefinisce ruoli, responsabilità e flussi informativi; un modello tecnicamente accurato può risultare inefficace se non è chiaro chi debba validarne l'*output*, in quali circostanze esso sia vincolante e come venga gestito l'errore. La domanda diventa allora: il sistema supporta decisioni strategicamente rilevanti per l'organizzazione, o si limita a ottimizzare procedure isolate?

Infine, vi è una dimensione culturale e normativa. L'introduzione dell'IA modifica il rapporto con l'autorità, con l'errore e con la responsabilità professionale. Un sistema percepito come opaco o imposto può risultare formalmente corretto ma sostanzialmente

---

valutazioni storiche, compromettendo la continuità decisionale e la tracciabilità dei criteri adottati nel tempo.

inefficace, perché non legittimato all'interno delle dinamiche organizzative. In questa prospettiva, trasparenza, *accountability* ed equità non costituiscono vincoli esterni, bensì condizioni interne dell'efficacia stessa. La dimensione normativa trova poi oggi una traduzione esplicita anche nel quadro regolatorio europeo. Ad esempio, il *Regolamento (UE) 2024/1689* sull'intelligenza artificiale (*AI Act*) classifica come *high-risk*<sup>34</sup> i sistemi impiegati in ambito occupazionale, inclusi quelli utilizzati per selezione, valutazione e gestione del personale. Per tali sistemi sono previsti obblighi stringenti in termini di documentazione, tracciabilità, gestione del rischio e supervisione umana. Ciò implica che l'efficacia di uno strumento non possa essere valutata esclusivamente in base alla *performance* tecnica, ma debba includere condizioni di controllabilità, contestabilità e responsabilità organizzativa. In questo senso, la regolazione europea non introduce un limite esterno all'innovazione, ma istituzionalizza la necessità di un *human-in-the-loop* strutturato, trasformando la supervisione da scelta opzionale a requisito sistemico<sup>35</sup>.

Esplicitare questi tre piani significa infine riconoscere che l'efficacia dell'IA è un equilibrio dinamico tra tecnica, integrazione organizzativa e maturazione culturale. La fragilità di uno solo di questi livelli può compromettere la tenuta complessiva del sistema intero. Questa tensione riflette poi una tematica già citata nel primo capitolo: mentre infatti l'intelligenza artificiale opera su correlazioni probabilistiche e *pattern* statistici, le organizzazioni operano su significati condivisi, interpretazioni e attribuzioni di senso. L'efficacia deve quindi emergere dall'incontro — non sempre semplice e lineare — tra una logica computazionale orientata alla previsione e una logica organizzativa orientata alla comprensione.

È a partire da qui che si pone perciò un problema metodologico fondamentale: attraverso quali criteri e strumenti l'efficacia può essere osservata senza ridurla alla sola *performance* computazionale?

---

<sup>34</sup> La sezione dell'*AI Act* (entrato in vigore nel mese di agosto 2024 con applicazione progressiva) che cita gli *high-risk* per sistemi HR è contenuta nell'Allegato III, punto 4.

<sup>35</sup> In questo senso, ciò che nel dibattito teorico viene definito *human-in-the-loop* assume perciò una configurazione giuridica concreta.

#### 4.1.3 Hard skills e soft skills: una simmetria solo apparente

Se l'efficacia dell'IA è multilivello, allora anche le competenze necessarie per interagire con essa non possono essere ridotte alla sola dimensione tecnica.

Una delle ambiguità più rilevanti nel dibattito sull'efficacia dell'IA nei contesti organizzativi riguarda infatti la presunta simmetria tra competenze tecniche e competenze trasversali. Se le prime sembrano prestarsi a metriche oggettive e relativamente stabili, le seconde vengono sempre più frequentemente sottoposte agli stessi strumenti di misurazione, come se condividessero la medesima natura epistemologica. Questa assimilazione, tuttavia, è tutt'altro che neutra.

Le *hard skills* si riferiscono generalmente a competenze tecniche specifiche, verificabili attraverso compiti strutturati, *test* standardizzati o indicatori di *performance* chiaramente definiti. In questi casi, il criterio di correttezza è relativamente esplicito: un calcolo è corretto o errato, un codice funziona o non funziona, una procedura è eseguita secondo *standard* o no. La misurazione si fonda su parametri relativamente stabili e su una relazione più lineare tra competenza, prestazione e risultato.

Le *soft skills*, al contrario, riguardano dimensioni relazionali, cognitive e adattive che si manifestano in modo situato. Creatività, pensiero critico, comunicazione, capacità di collaborazione o gestione dell'ambiguità non si esauriscono in una prestazione isolata, ma emergono dall'interazione tra individuo, compito e contesto. Esse non sono semplicemente proprietà interne del soggetto, bensì configurazioni dinamiche che prendono forma in situazioni concrete.

In ambito psicologico, strumenti come il modello dei *Big Five* (Costa & McCrae, 1992) consentono di descrivere tratti relativamente stabili della personalità offrendo una base comparativa utile per l'analisi individuale. Tuttavia, è necessario distinguere tra tratti disposizionali e competenze operative. La letteratura sulle competenze organizzative distingue infatti tra caratteristiche disposizionali relativamente stabili e competenze osservabili in azione, intese come combinazioni di conoscenze, abilità e comportamenti manifestati in contesti specifici (Boyatzis, 1982; Spencer & Spencer, 1993). Confondere questi livelli implica un errore metodologico nella progettazione degli strumenti di *assessment*: misurare un tratto non equivale a valutare una competenza situata. Eppure questa confusione è sistematicamente presente in molti

strumenti di *assessment* supportati da IA, che trattano *output* comportamentali osservabili come *proxy* affidabili di disposizioni profonde, o viceversa. Un tratto stabile non equivale infatti automaticamente a una competenza agita in un contesto organizzativo specifico, e misurare una predisposizione non significa quindi misurare la qualità di una decisione in situazione, né la capacità di negoziare significati in un *team* interdisciplinare.

La tendenza a trattare le *soft skills* come variabili pienamente quantificabili rischia così di trasformare un processo dinamico in una fotografia statica. Uno *score* può infatti offrire un'indicazione sintetica, ma non esaurisce la complessità della *performance* situata. In questo senso, la misurazione non è mai una semplice registrazione neutrale di una realtà data, bensì un'operazione di riduzione e selezione: decide cosa conta, cosa viene reso visibile e cosa rimane escluso.

Quando tali strumenti di misurazione vengono impiegati in ambienti caratterizzati dall'integrazione uomo-macchina, la questione si fa ancora più delicata. Se l'efficacia dell'IA dipende anche dalla capacità degli attori di interpretare e contestualizzare gli *output* algoritmici, allora le competenze trasversali non sono soltanto oggetto di valutazione, ma diventano parte integrante del funzionamento del sistema stesso. È su questa ambiguità — tra misurazione, interpretazione e decisione — che si innesta un altro importante problema metodologico.

#### **4.1.4 I rischi metodologici della misurazione in contesti AI-driven**

Se la misurazione delle *soft skills* comporta già di per sé una riduzione della complessità, l'integrazione di strumenti di intelligenza artificiale nei processi valutativi introduce un ulteriore livello di criticità. Il rischio metodologico più rilevante non risiede esclusivamente nell'accuratezza tecnica dell'algoritmo, ma nel momento in cui la misurazione viene trasformata in criterio decisionale rigido.

Nei contesti organizzativi, i sistemi di *assessment* supportati da IA producono frequentemente *output* sintetici sotto forma di punteggi, classificazioni o profili comparativi. Tali rappresentazioni, proprio in virtù della loro forma numerica e formalizzata, possono essere percepite come oggettive e neutrali. Tuttavia, come evidenziato negli studi sull'*automation bias* (Lee & See, 2004), gli individui tendono ad attribuire ai sistemi automatizzati un grado di affidabilità superiore rispetto al giudizio

umano, riducendo la propensione alla verifica critica. In questo modo, il risultato algoritmico rischia di acquisire un'autorità epistemica che eccede le reali possibilità dello strumento.

Il problema non riguarda qui soltanto l'eventuale errore tecnico, ma la naturalizzazione del punteggio fornito dall'IA. Quando un punteggio relativo a competenze trasversali viene utilizzato come criterio decisionale vincolante — ad esempio nei processi di selezione, promozione o allocazione di opportunità formative — la misurazione smette di essere uno strumento conoscitivo e diventa un dispositivo normativo. Non descrive più soltanto una *performance*, ma stabilisce soglie di accesso, definisce *standard* impliciti di adeguatezza e contribuisce a modellare le traiettorie professionali degli individui.

In questo passaggio si manifesta un rischio strutturale: la riduzione di competenze situate e dinamiche a indicatori sintetici può generare effetti di esclusione non intenzionali, soprattutto se i criteri di costruzione dello strumento non sono pienamente trasparenti o contestualizzati.

A questo si aggiunge un ulteriore fenomeno, meno immediato ma altrettanto rilevante: la progressiva delega cognitiva. Se l'algoritmo fornisce valutazioni sintetiche e apparentemente coerenti, l'organizzazione può ridurre gradualmente lo spazio del giudizio interpretativo, confondendo l'efficacia con la standardizzazione e la decisione con l'applicazione meccanica di un criterio numerico. In tal modo, la complessità della valutazione viene sostituita dalla conformità a uno *score*.

Riconoscere questi rischi non significa negare il valore degli strumenti di *assessment* supportati da IA, ma chiarire una condizione fondamentale: nessuna misurazione esaurisce la realtà che intende rappresentare. In questo senso, la supervisione umana non costituisce soltanto un requisito procedurale, ma il presidio attraverso cui l'organizzazione mantiene aperta la possibilità di interpretazione, contestazione e revisione del risultato algoritmico. Senza questo spazio, l'efficacia rischia di ridursi a mera conformità statistica, perdendo la dimensione critica e responsabile che dovrebbe caratterizzare ogni decisione organizzativa.

È in questo quadro di tensioni irrisolte — tra misurazione e interpretazione, tra sintesi algoritmica e giudizio situato — che il *case study* di *Slash* offre un terreno di osservazione particolarmente fertile. Analizzarlo significa interrogarsi non tanto su se

sia possibile misurare le competenze in contesti *AI-driven*, ma su come farlo senza trasformare la valutazione in automatismo decisionale.

#### **4.2 Slash soft-skill assessment: a case study**

Il caso di *Slash* consente quindi di osservare in modo situato come le tensioni metodologiche fin qui discusse vengano affrontate in un contesto applicativo concreto.

Sviluppata da *Altermaind*, la piattaforma si propone di supportare le organizzazioni nell'adozione dell'intelligenza artificiale attraverso percorsi di *assessment* e formazione, con particolare attenzione alla valutazione delle competenze trasversali rilevanti in ambienti di lavoro caratterizzati dall'integrazione uomo-macchina.

A differenza di una prospettiva puramente tecnica, l'interesse analitico di questo caso non risiede nella sofisticazione dell'algoritmo in sé, ma nel modo in cui lo strumento viene collocato all'interno di un processo decisionale più ampio. *Slash* non si presenta infatti come dispositivo sostitutivo del giudizio umano, bensì come supporto alla riflessione organizzativa, producendo *output* sintetici che richiedono interpretazione e contestualizzazione.

L'analisi del caso si fonda su una combinazione di fonti, ovvero la documentazione messa a disposizione dall'organizzazione e materiali descrittivi del progetto. A questi si aggiunge poi un colloquio semi-strutturato condotto presso la sede di *Altermaind* con la Dott.ssa Viola Marinelli, una dei responsabili del progetto, che ha consentito di integrare la lettura documentale con una comprensione più articolata delle scelte progettuali sottese allo strumento. Il colloquio, svolto in forma non registrata con il consenso dell'interlocutrice e nel rispetto dei criteri comunicativi indicati dall'organizzazione, è stato ricostruito attraverso appunti e note successive alla conversazione. I contenuti riportati sono stati inoltre sottoposti a validazione da parte di *Altermaind*.

Comprendere *Slash* richiede anzitutto di ricostruirne la genesi: le ragioni che hanno portato alla sua ideazione, il problema organizzativo che intende risolvere e la logica progettuale che ne orienta l'architettura. È da qui che emerge con maggiore chiarezza il posizionamento dello strumento come interlocutore di supporto al processo decisionale.

#### 4.2.1 Genesi del progetto

*Slash* si inserisce all'interno delle attività di *Altermaind*, una *AI & technology company* focalizzata sullo sviluppo di soluzioni per l'integrazione dell'intelligenza artificiale nei processi organizzativi. Il progetto nasce da una constatazione empirica maturata nel lavoro con le aziende: l'interesse verso l'adozione dell'IA cresce rapidamente, ma tale interesse non è sempre accompagnato da un adeguato livello di competenze diffuse tra i collaboratori.

Il problema non riguarda soltanto la conoscenza tecnica degli strumenti, ma la capacità di integrarli in modo critico, responsabile e produttivo nei flussi decisionali e collaborativi. In altre parole, molte organizzazioni desiderano implementare soluzioni IA, ma non sempre dispongono di persone preparate a utilizzarle, interpretarle e governarle in modo efficace. Da qui la domanda generativa che ha orientato la nascita di *Slash*: cosa accade quando si introduce l'IA in azienda senza aver prima misurato e sviluppato le competenze necessarie per utilizzarla?

*Slash* risponde a questa tensione attraverso una logica sequenziale: *assessment* prima, *training* poi. L'idea non è proporre immediatamente percorsi formativi standardizzati, ma costruire una mappatura preliminare delle competenze *soft* e *hard* rilevanti per lavorare in contesti *AI-driven*. Solo a partire da questa "fotografia" iniziale diventa possibile progettare interventi mirati, differenziati per ruolo e per livello di maturità.

Il *target* primario della piattaforma sono le funzioni HR, intese come snodo strategico dei processi di formazione e sviluppo organizzativo. L'*assessment* può essere infatti utilizzato in fase di selezione, ma la sua logica dichiarata è principalmente evolutiva, cioè offrire alle aziende uno strumento per comprendere il proprio livello di *AI readiness* e accompagnare nel tempo la crescita delle competenze interne.

In questa prospettiva, *Slash* non si propone come sostituto del giudizio umano, bensì come dispositivo di supporto alla valutazione. L'*assessment* produce indicatori sintetici e *insight* interpretativi, ma ogni decisione resta in capo agli attori organizzativi. Come emerso dall'interlocuzione con la responsabile del progetto, l'obiettivo non è infatti automatizzare il giudizio, bensì affiancare gli HR con uno strumento capace di rendere visibili *pattern* e *gap* che altrimenti resterebbero impliciti. *Slash* può essere inteso, in

una metafora organizzativa, come un "collega" che fornisce dati e prospettive, senza sostituirsi alla responsabilità decisionale.

Letto alla luce della riflessione teorica sviluppata nella sezione precedente, il progetto rappresenta perciò un tentativo di tradurre operativamente il principio del *human-in-the-loop*: l'intelligenza artificiale viene integrata come supporto analitico, ma non assume una funzione deliberativa autonoma. Questa impostazione costituisce il presupposto su cui si costruisce l'architettura metodologica dello strumento.

#### 4.2.2 Architettura del questionario

L'architettura del *soft skills assessment* di Slash combina due piani di valutazione distinti ma complementari: da un lato un *assessment* di competenze trasversali ritenute rilevanti quando si lavora in ambienti organizzativi mediati dall'IA; dall'altro una valutazione delle *hard skills* legate alla conoscenza pratica e all'uso operativo degli strumenti di intelligenza artificiale. Questa scelta risponde a un presupposto preciso emerso nella genesi del progetto: l'efficacia dell'adozione dell'IA non dipende solo dalla capacità tecnica di utilizzare un *tool*, ma anche dalla capacità di interpretarne gli *output*, integrare il supporto algoritmico nel ragionamento e mantenere attivo il presidio critico umano (*human-in-the-loop*). In questo senso, *Slash* non valuta semplicemente il "saper usare l'IA", ma tenta di intercettare anche le condizioni cognitive e relazionali che rendono tale uso produttivo e governabile nel tempo.

Sul piano metodologico, questa impostazione evita una riduzione frequente negli strumenti di *AI readiness*, ovvero trattare l'adozione dell'IA come un problema esclusivamente tecnico. L'*assessment* delle *hard skills* consente infatti una misurazione più diretta e verificabile della familiarità operativa, mentre la componente sulle *soft skills* mira a offrire una lettura orientativa delle risorse personali e comportamentali che influenzano la qualità dell'interazione uomo-macchina, soprattutto in situazioni di ambiguità decisionale e di collaborazione. Tale distinzione è rilevante anche rispetto alla letteratura sulle competenze, che separa le caratteristiche disposizionali relativamente stabili dalle competenze agite e sviluppabili, intese come repertori di comportamenti osservabili e contestuali (Boyatzis, 1982; Spencer & Spencer, 1993). L'*assessment* di *Slash* si colloca quindi nella zona metodologicamente più delicata: trasformare

dimensioni qualitative in indicatori sintetici senza confondere tratti, competenze e *performance* situate.

La componente relativa alle *soft skills* si articola in 24 domande situazionali, costruite per esplorare il modo in cui l'individuo tende ad agire di fronte a scenari concreti. La scelta di adottare *item* basati su situazioni — anziché affermazioni astratte su di sé — si colloca nella tradizione dei *situational judgment tests* (SJT), strumenti ampiamente utilizzati in ambito organizzativo per valutare comportamenti potenziali in contesti realistici (McDaniel et al., 2001; Lievens & Sackett, 2012). In questa impostazione, l'attenzione si sposta dalla dichiarazione di tratti (per esempio "sono creativo", "sono collaborativo", ...) alla simulazione di decisioni e reazioni in circostanze specifiche.

Un elemento progettuale rilevante riguarda poi la natura degli scenari proposti. Le situazioni non sono collocate esplicitamente in contesti aziendali formali, ma possono riferirsi a episodi di vita quotidiana o interazioni generali. Tale scelta risponde all'esigenza di ridurre l'effetto della *social desirability*, ossia la tendenza dei rispondenti a fornire risposte percepite come socialmente più accettabili o professionalmente vantaggiose (Crowne & Marlowe, 1960). In contesti HR, dove l'*assessment* può essere associato a processi di valutazione o sviluppo, il rischio che l'individuo cerchi di "mostrarsi nel modo migliore possibile" è particolarmente elevato. La costruzione di scenari meno direttamente connessi al ruolo professionale mira quindi a favorire risposte più spontanee e meno strategicamente orientate all'*impression management*.

Le opzioni di risposta sono formulate in modo tale da evitare una gerarchia esplicita tra alternative "migliori" o "peggiori". In molti casi, le risposte sono tutte valutative, cioè esprimono modalità differenti di azione piuttosto che contrapposizioni tra comportamento corretto ed errore evidente. Questa impostazione intende ridurre l'effetto di risposte normative e incoraggiare una scelta più autentica, pur nella consapevolezza che nessun formato elimina completamente le distorsioni motivazionali.

La struttura sintetica dello strumento — circa dieci minuti di compilazione — risponde inoltre a una precisa esigenza organizzativa: garantire fruibilità in contesti lavorativi caratterizzati da tempi limitati e soglie di attenzione ridotte. La letteratura sui questionari brevi evidenzia come la riduzione del numero di *item*, se progettata con attenzione, possa infatti mantenere livelli accettabili di affidabilità pur aumentando la

partecipazione e la qualità della risposta (Rammstedt & John, 2007; Rammstedt & Beierlein, 2014). In questo senso, la brevità non è solo un vincolo operativo, ma una scelta strategica coerente con l'uso aziendale dello strumento.

Oltre alla natura situazionale degli *item*, un ulteriore elemento rilevante dell'architettura riguarda il formato delle risposte. Il questionario combina modalità differenti — risposta singola, risposta multipla e utilizzo di differenziali semantici — con l'obiettivo di rendere la compilazione varia e cognitivamente attiva, pur mantenendo una struttura sintetica.

In particolare, l'impiego di differenziali semantici merita una breve precisazione metodologica. Il *semantic differential*, introdotto da Osgood, Suci e Tannenbaum (1957), consiste in una scala bipolare collocata tra due poli concettualmente opposti (ad esempio: "analitico" vs "intuitivo"), lungo la quale il rispondente posiziona la propria tendenza comportamentale. Questo formato consente di cogliere gradazioni di orientamento piuttosto che scelte dicotomiche, offrendo una rappresentazione più sfumata delle inclinazioni individuali. In un contesto come quello di *Slash*, tale impostazione permette di evitare giudizi impliciti di valore e di restituire una misura continua del comportamento percepito.

Le risposte multiple non gerarchiche costituiscono poi un altro elemento distintivo. In questi casi, le opzioni non sono costruite come alternative corrette o scorrette, ma come modalità diverse di affrontare una situazione. Il sistema di *scoring* opera quindi su ponderazioni interne che non sono immediatamente trasparenti al compilatore, riducendo la possibilità di orientare strategicamente la risposta verso quella ritenuta "migliore". Questa scelta si colloca in una tensione tipica degli strumenti di *assessment* organizzativo: bilanciare comprensibilità per l'utente e robustezza metodologica.

La logica complessiva della misurazione non mira perciò a produrre un giudizio definitivo, ma una "fotografia" del momento attuale. Il risultato viene restituito attraverso una rappresentazione grafica *radar*, accompagnata da *insight* interpretativi che descrivono il significato dei punteggi ottenuti. È importante sottolineare che tale restituzione non esprime una valutazione normativa, bensì un'indicazione orientativa sul livello di sviluppo relativo delle diverse competenze. In questo senso, la piattaforma enfatizza la dimensione evolutiva del processo: la misurazione rappresenta un punto di

partenza per la consapevolezza e il miglioramento e non una classificazione statica dell'individuo.

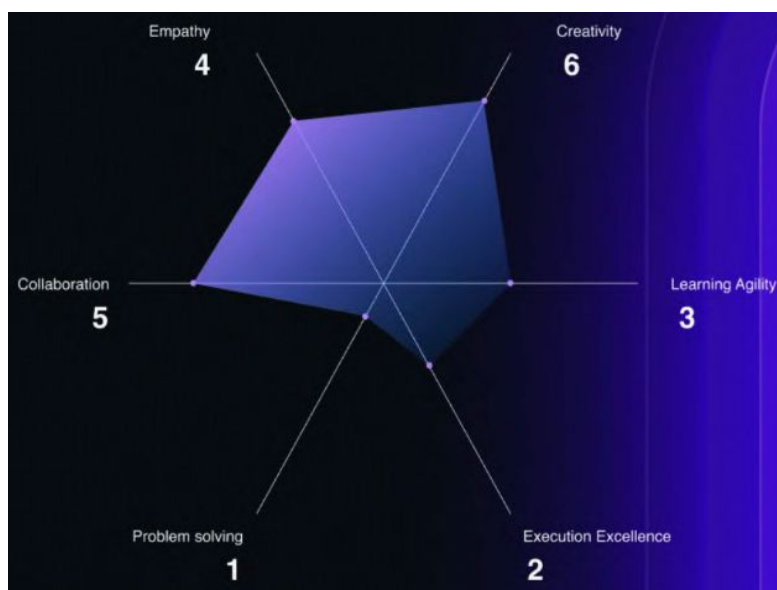


Fig. 4.2.2: Esempio di valutazione del soft skills assessment di Slash.

FONTE: AM Slash: Brochure generale [Allegato A].

È opportuno precisare che l'*assessment* di *Slash* si trova attualmente in fase di validazione. Questa fase rappresenta una tappa metodologicamente necessaria nel ciclo di sviluppo di uno strumento psicometrico. La validazione non consiste infatti in un singolo *test*, bensì in un processo articolato volto a verificare coerenza interna, stabilità nel tempo e corrispondenza con costrutti teorici affini.

Un primo livello riguarda l'analisi della stabilità dei punteggi nel tempo (*test-retest reliability*), attraverso la ripetizione del questionario a distanza di un intervallo definito. La previsione di una ri-somministrazione a un anno consente di osservare se i punteggi rimangano coerenti in assenza di interventi formativi significativi, o se evolvano in presenza di percorsi di sviluppo mirati. Tale passaggio è particolarmente rilevante in strumenti che dichiarano una vocazione evolutiva, poiché permette di distinguere tra fluttuazioni casuali e cambiamenti strutturati.

Un secondo livello concerne la validità convergente e discriminante. In questa prospettiva, *Slash* prevede il confronto con scale validate e ampiamente riconosciute nella letteratura psicologica, tra cui ad esempio il già citato modello dei *Big Five* (Costa & McCrae, 1992). Il riferimento al *Big Five* non implica una sovrapposizione diretta tra

tratti di personalità e *soft skills* operative, ma costituisce, tra gli altri modelli utilizzati, uno sfondo teorico utile per verificare la coerenza concettuale delle dimensioni misurate. Il confronto con strumenti consolidati consente infatti di valutare se le dimensioni di *Slash* mostrino correlazioni attese con tratti teoricamente affini (validità convergente) e assenza di correlazioni indebite con costrutti non pertinenti (validità discriminante).

Il campione attuale della fase *beta* conta circa 200 risposte, numero ancora preliminare ma sufficiente per avviare analisi esplorative sulla struttura e sulla coerenza interna degli *item*. In questa fase, la priorità metodologica non è affermare conclusioni definitive, bensì verificare se le domande effettivamente intercettino le dimensioni teoriche che dichiarano di misurare. La sfida non è tanto costruire un questionario coerente sul piano intuitivo, quanto dimostrare empiricamente che esso discrimini in modo affidabile tra livelli differenti di competenza.

Questo processo di validazione progressiva colloca *Slash* in una posizione intermedia tra sperimentazione e consolidamento: lo strumento si può dichiarare operativo, ma il suo perfezionamento è ancora in corso. Da un punto di vista metodologico, tale trasparenza rappresenta un elemento di rigore scientifico, poiché riconosce che la misurazione delle *soft skills* richiede un lavoro continuo di raffinamento e confronto con la letteratura esistente. Tuttavia, anche un processo di validazione rigoroso non elimina la questione epistemologica più ampia: cosa significa, in termini organizzativi, misurare una competenza situata?

#### **4.2.3 Il problema della misurazione delle soft skills nel caso Slash**

Se l'architettura del questionario evidenzia attenzione metodologica nella costruzione degli *item* e nella gestione dei *bias* più evidenti, il nodo critico emerge nel momento in cui il risultato entra nei processi organizzativi concreti. La questione non riguarda tanto la coerenza interna dello strumento, quanto il modo in cui la sua sintesi viene interpretata e utilizzata nelle pratiche decisionali.

Nel caso delle *hard skills* legate all'uso dell'IA, la valutazione mantiene un ancoraggio relativamente diretto alla *performance*: familiarità operativa, capacità applicativa, comprensione funzionale degli strumenti. Pur con i limiti propri di ogni *test*, la misurazione si riferisce a dimensioni verificabili. Diverso è il caso delle *soft skills*

intercettate da *Slash*, che non corrispondono a prestazioni osservabili in senso stretto, ma a orientamenti comportamentali inferiti da scelte situazionali. Il punteggio non descrive un'azione compiuta, bensì una tendenza probabile.

La restituzione dei risultati attraverso una rappresentazione *radar*, accompagnata da *insight* interpretativi, è coerente con l'intento dichiarato di fornire uno strumento di consapevolezza. Tuttavia, la traduzione di tali orientamenti in valori differenziati espone a un primo rischio: la reificazione del profilo. Ciò che nasce come indicazione orientativa può essere percepito, soprattutto in contesti organizzativi strutturati, come descrizione stabile dell'identità professionale. La distanza tra "livello relativo attuale" e "caratteristica personale" tende a ridursi nel momento in cui il dato assume forma grafica e comparativa.

Un secondo nodo riguarda l'uso decisionale dei risultati. Sebbene *Slash* non si proponga come sostituzione del giudizio umano e mantenga formalmente la centralità dell'HR nel processo di interpretazione, la presenza di uno *score* sintetico introduce una dinamica organizzativa ricorrente: la trasformazione della complessità in criterio rapido. Il rischio non è insito nello strumento, ma nella possibilità che la sintesi venga utilizzata come filtro preliminare o come legittimazione *ex post* di decisioni già orientate. In tale passaggio, la funzione conoscitiva può scivolare verso una funzione selettiva non dichiarata.

La dimensione evolutiva dichiarata dal progetto apre inoltre un ulteriore interrogativo. La ri-somministrazione a distanza di un anno presuppone che le competenze possano mostrare una traiettoria di sviluppo misurabile. Tuttavia, nel caso delle *soft skills*, un incremento di punteggio non costituisce un indicatore univoco di crescita sostanziale. Può riflettere un effettivo cambiamento comportamentale, ma anche una maggiore familiarità con il formato dello strumento o un adattamento alle logiche percepite di valutazione. La distinzione tra trasformazione reale e adeguamento strategico rimane metodologicamente aperta.

Infine, nel corso del colloquio con la Dott.ssa Marinelli è stata discussa, in termini esplorativi, la possibilità di un'evoluzione verso un modello di *assessment* conversazionale basato su IA. Tale prospettiva, ancora in fase di riflessione e non formalmente implementata, introdurrebbe una variabile ulteriore. Se il questionario strutturato mantiene infatti una distanza procedurale tra soggetto e sistema,

un'interazione dialogica potrebbe modificare la percezione del momento valutativo, rendendolo più personale e potenzialmente più esposto alla sensazione di giudizio diretto. In tal caso, la dimensione relazionale dello strumento diventerebbe parte integrante della sua efficacia o inefficacia.

Il caso *Slash* mostra dunque come la misurazione delle *soft skills* in contesti *AI-driven* sia intrinsecamente attraversata da tensioni tra sintesi e complessità, tra orientamento e classificazione, tra supporto e standardizzazione. La questione non è eliminare tali tensioni — operazione impossibile — ma riconoscerle e governarle. È in questo spazio di gestione consapevole che si gioca la differenza tra uno strumento tecnicamente funzionante e uno strumento organizzativamente responsabile.

#### **4.2.4 Slash tra governance e co-creazione**

Le tensioni evidenziate nella sezione precedente non restano, nel caso *Slash*, sul piano teorico, ma vengono affrontate attraverso scelte organizzative e procedurali specifiche. Il progetto non si presenta come una soluzione automatica al problema della valutazione delle competenze, bensì come un dispositivo inserito in una rete di relazioni umane che ne orientano l'uso e ne delimitano l'impatto.

Un primo elemento rilevante riguarda infatti la centralità del ruolo HR. L'*assessment* non è concepito come strumento autonomo a disposizione del singolo dipendente, ma come processo mediato: sono le risorse umane a gestire l'invio, la lettura e l'eventuale restituzione dei risultati. Questa configurazione introduce una forma di supervisione organizzativa che riduce il rischio di automatizzazione incontrollata del giudizio. Il punteggio non circola in modo neutro e impersonale, ma è inserito in un contesto interpretativo.

In secondo luogo, la fase *beta* del progetto prevede una raccolta sistematica di *feedback* e la possibilità di consulenza personalizzata. Tale scelta indica che lo strumento non viene trattato come prodotto chiuso, ma come oggetto in evoluzione, aperto a revisioni. La co-creazione non riguarda soltanto la fase iniziale di progettazione, ma continua nell'interazione con le aziende utilizzatrici. In questo senso, la misurazione non è concepita come atto definitivo, ma come parte di un processo dialogico.

Un ulteriore aspetto riguarda la costruzione degli *insight* interpretativi che accompagnano i punteggi. La presenza di spiegazioni testuali accanto alla rappresentazione grafica attenua il rischio di lettura puramente numerica e incoraggia una riflessione qualitativa sul significato del risultato. Il dato viene così collocato in una cornice discorsiva, che ne limita l'uso meccanico.

Queste scelte possono essere lette in parallelo con il principio di supervisione umana discusso nei paragrafi precedenti. In *Slash*, la tecnologia non sostituisce l'interpretazione, ma produce un supporto strutturato che rimane subordinato alla responsabilità decisionale umana. L'*accountability* non è delegata all'algoritmo, ma resta collocata all'interno dell'organizzazione.

In questa prospettiva, *Slash* non si configura come *tech solution* autosufficiente, bensì come tentativo di costruire un modello ibrido di *assessment*. L'efficacia dello strumento non dipende esclusivamente dalla sua architettura tecnica, ma dalla qualità delle pratiche che ne accompagnano l'uso: supervisione, contestualizzazione, dialogo, revisione nel tempo. È in questa interazione tra dispositivo tecnologico e *governance* organizzativa che si gioca la sostenibilità del modello.

In questo senso, il caso *Slash* offre una possibile declinazione concreta di ciò che può significare "il futuro delle IA nelle organizzazioni". Non un futuro inteso come progressiva automatizzazione del giudizio, ma come costruzione di modelli ibridi in cui tecnologia e responsabilità organizzativa si co-determinano. La questione non è se l'IA possa misurare, ma a quali condizioni tale misurazione venga inserita in processi governati, interpretati e continuamente rinegoziati. Il futuro, in questa prospettiva, non coincide con l'autonomia della macchina, ma con la capacità delle organizzazioni di strutturare dispositivi tecnologici che restino iscritti in pratiche di costruzione di senso.

## CONCLUSIONI

Il percorso analitico sviluppato in questa tesi ha preso le mosse da una domanda di fondo: in quale modo un progetto come *Lead-Alliance*, nato dalla convergenza volontaria di professionisti con *background* eterogenei, riesce a costruire senso attorno all'integrazione dell'intelligenza artificiale nelle pratiche di *leadership*? La risposta, come si è cercato di argomentare nei quattro capitoli che compongono il presente elaborato, non è riducibile a una formula tecnica né a un modello organizzativo standardizzabile. Essa risiede piuttosto nel carattere processuale, negoziale e comunicativo di quella costruzione, oltre che nelle tensioni, a volte irrisolte, che essa porta con sé.

Il Capitolo 1 ha offerto le coordinate teoriche necessarie per leggere questa dinamica. La *leadership*, nel quadro del *sensemaking* weickiano, non è un attributo individuale ma un processo distribuito di attribuzione di significato all'azione collettiva. In contesti attraversati dall'intelligenza artificiale, tale processo si complica ulteriormente: l'IA non si limita a ottimizzare flussi decisionali, ma interviene sui modi in cui gli attori organizzativi percepiscono, interpretano e comunicano la realtà. La nozione di co-intelligenza, sviluppata a partire dalla letteratura su *human-in-the-loop* e collaborazione uomo-macchina, ha permesso di superare una visione duale e oppositiva tra intelligenza umana e artificiale, mostrando come il valore dell'IA emerga dall'interazione con pratiche comunicative condivise, contesti relazionali e capacità critiche situate. La co-intelligenza non è dunque una proprietà intrinseca della tecnologia, ma una qualità emergente che si costruisce nel tempo.

Il Capitolo 2 ha applicato queste categorie alla seconda fase del progetto *Lead-Alliance*, avviata formalmente nel settembre 2025. L'analisi ha messo in evidenza come il passaggio da *hub* esplorativo a "base operativa" strutturata abbia comportato non soltanto una riorganizzazione funzionale — con la definizione di un senato ristretto e di *team* tematici — ma anche una ridefinizione identitaria del progetto stesso. L'emergere di una *governance* interna più formalizzata ha introdotto nuovi equilibri tra autonomia e coordinamento, tra visione condivisa e specializzazione operativa. In questa fase, la costruzione dell'identità di progetto non ha seguito un percorso lineare: è stata il risultato di negoziazioni continue, di aggiustamenti progressivi e di una tensione

mai del tutto risolta tra la dimensione laboratoriale originaria e le esigenze di stabilità e riconoscibilità istituzionale.

Il Capitolo 3 ha approfondito le scelte comunicative attraverso cui *Lead-Alliance* ha costruito la propria presenza pubblica. L'analisi della comunicazione interna ha mostrato come la fiducia e il coordinamento si sedimentino in pratiche condivise, spesso informali, che richiedono costante manutenzione. Sul fronte esterno, la definizione di una *brand identity* coerente — con un proprio *tone of voice* (t.o.v.), un posizionamento riconoscibile e una *content strategy* articolata su più formati — ha rappresentato uno sforzo significativo di capitalizzazione simbolica. Il sito *web*, il *podcast*, la *newsletter* e il libro, che si stanno costruendo, configurano un ecosistema editoriale che non si limita a diffondere contenuti, ma costruisce una cornice narrativa attraverso cui il progetto legittima la propria esistenza e recluta pubblici. Allo stesso tempo, l'analisi ha evidenziato un rischio latente: la proliferazione di *output* e iniziative può generare dispersione, rendendo difficile mantenere la coerenza identitaria in assenza di una formalizzazione strutturale più solida. L'ipotesi di integrazione con *OpenHS* emerge in questo contesto come possibile risposta organizzativa a tale fragilità, senza tuttavia esaurire le questioni aperte.

Il Capitolo 4 ha spostato l'analisi su un piano più sistematico, affrontando la questione dell'efficacia dell'IA nei contesti organizzativi con un'attenzione critica ai suoi limiti strutturali: *bias* algoritmici, rischi di *overfitting*, derive nel tempo, asimmetrie tra *hard skills* e *soft skills* nella misurazione computazionale. Il *case study* relativo al progetto *Slash* — una piattaforma di *soft-skill assessment* e formazione — ha offerto un esempio concreto di come sia possibile costruire modelli ibridi in cui la tecnologia non sostituisce il giudizio umano, ma ne costituisce un supporto strutturato, inserito in processi di supervisione, contestualizzazione e revisione continua. In *Slash*, l'*accountability* non è delegata all'algoritmo ma resta ancorata alle pratiche organizzative: sono le risorse umane a mediare l'uso dello strumento, e persino la fase *beta* è concepita come un dispositivo aperto al dialogo e alla co-creazione. Questo modello offre una possibile declinazione concreta di ciò che può significare integrare l'IA nelle organizzazioni in modo responsabile: non come progressiva automatizzazione del giudizio, ma come costruzione di architetture ibride in cui tecnologia e responsabilità si co-determinano.

Rispetto all'impianto metodologico adottato, vale la pena soffermarsi su una riflessione conclusiva. La ricerca si è avvalsa di un approccio partecipante che ha consentito un accesso diretto ai processi decisionali, alle tensioni interne e alle negoziazioni di senso che hanno accompagnato la seconda fase del progetto. Questo posizionamento ha offerto indubbi vantaggi in termini di profondità analitica e densità empirica. Ha comportato, tuttavia, un rischio di parzialità che non può essere del tutto neutralizzato, nemmeno attraverso il rigore delle procedure di analisi. La prospettiva di chi ha contribuito attivamente alla costruzione dell'oggetto studiato resta inevitabilmente situata. Per questa ragione, i risultati presentati in questo elaborato vanno letti non come descrizione esaustiva di *Lead-Alliance*, ma come un'interpretazione contestuale e riflessiva, consapevole dei propri limiti.

In conclusione, ciò che il progetto *Lead-Alliance* — con le sue potenzialità e le sue fragilità — sembra suggerire è che l'integrazione dell'intelligenza artificiale nelle organizzazioni non è anzitutto una questione tecnica, ma una questione di senso. Richiede la capacità di costruire linguaggi condivisi, di negoziare valori e priorità, di sostenere pratiche comunicative che tengano insieme la complessità della condizione umana con le opportunità offerte dall'automazione. In questo quadro, la *leadership* non si esaurisce nella *governance* di processi e risultati: è chiamata a farsi custode del significato, in un contesto in cui il confine tra decisione umana e suggerimento algoritmico diventa sempre più sottile e, per questo, sempre più degno di attenzione critica.

Il contributo specifico di questo lavoro risiede nell'aver osservato l'integrazione tra intelligenza artificiale e *leadership* non dal punto di vista dell'adozione tecnologica, ma dal punto di vista comunicativo e organizzativo. L'attenzione alle pratiche di costruzione del senso, ai dispositivi narrativi e alle forme di *governance* emergenti consente di spostare il *focus* dalla *performance* dell'algoritmo alla qualità delle interazioni che ne accompagnano l'uso. In tal modo, la tesi propone una chiave interpretativa che può risultare utile non solo per comprendere il caso analizzato, ma per leggere criticamente processi analoghi in altri contesti organizzativi.

Le traiettorie aperte da questo lavoro sono molteplici. Sul piano empirico, sarebbe di grande interesse seguire l'evoluzione di *Lead-Alliance* nelle fasi successive. Sul piano teorico, la nozione di co-intelligenza merita approfondimenti ulteriori, tanto rispetto alle

sue condizioni di possibilità quanto alle sue implicazioni per i processi formativi e per il disegno organizzativo. La domanda che rimane aperta — e che questa tesi ha inteso contribuire ad articolare, senza pretesa di rispondervi definitivamente — è se e in quale misura le organizzazioni contemporanee abbiano sviluppato le competenze culturali e comunicative necessarie a governare, e non soltanto subire, la trasformazione che l'intelligenza artificiale porta con sé.

## BIBLIOGRAFIA

Aguinis, H., Beltran, J. R., & Cope, A. (2024). How to use Generative AI as a human resource management assistant. *Organizational Dynamics*, 53(1), Article 101029.

<https://doi.org/10.1016/j.orgdyn.2024.101029>

Aaker, D. A. (1996). *Building Strong Brands*. Free Press.

[https://irp.cdn-website.com/e38aeb7a/files/uploaded/%5BM%5DDavid\\_A.\\_Aaker\\_Building\\_Strong\\_Brands.pdf](https://irp.cdn-website.com/e38aeb7a/files/uploaded/%5BM%5DDavid_A._Aaker_Building_Strong_Brands.pdf)

Albert, S., & Whetten, D. A. (1985). Organizational identity. *Research in Organizational Behavior*, 7, 263–295. <https://psycnet.apa.org/record/1986-02640-001>

altermAInd. (2026). *AM Slash: Brochure generale* [Brochure aziendale]. Allegato A.

Ancona, D., & Caldwell, D. (1992). Demography and design: Predictors of new product team performance. *Organization Science*, 3(3), 321–341. [https://www.researchgate.net/publication/5175846\\_Demography\\_And\\_Design\\_Predictors\\_Of\\_New\\_Product\\_Team\\_Performance](https://www.researchgate.net/publication/5175846_Demography_And_Design_Predictors_Of_New_Product_Team_Performance)

Angelelli, A., Bellucci, S., Mancini, D., Pacelli, M., Panai, E., Primanni, M., Rolando, S., Sodano, G., & Alessandrini, A. (2021). *To be digital*. Moondo. [https://cudriec.com/download/libri/To\\_Be\\_Digital.pdf](https://cudriec.com/download/libri/To_Be_Digital.pdf)

Argyris, C., & Schön, D. A. (1978). *Organizational learning: A theory of action perspective*. Addison-Wesley. <https://archive.org/details/organizationalle00chri>

Bala, S. (2024). Impact of technology on multidisciplinary collaboration. *Jagannath University Journal of Research and Review*, 1(1), 36-41. [https://www.jagannathuniversityncr.ac.in/assets/docs/journal/Suman\\_Bala\\_Impact\\_of\\_Technology\\_on\\_Multidisciplinary\\_Collaboration.pdf](https://www.jagannathuniversityncr.ac.in/assets/docs/journal/Suman_Bala_Impact_of_Technology_on_Multidisciplinary_Collaboration.pdf)

Barocas, S., & Selbst, A. D. (2016). Big Data's disparate impact. *California Law Review*, 104(6), 671–732. <https://www.cs.yale.edu/homes/jf/BarocasSelbst.pdf>

Barrick, M. R., & Mount, M. K. (1991). The Big Five personality dimensions and job performance: A meta-analysis. *Personnel Psychology*, 44(1), 1–26. <https://doi.org/10.1111/j.1744-6570.1991.tb00688.x>

Benanti, P. (2018). *Homo Faber: The techno-human condition*. EDB – Edizioni Dehoniane Bologna.

- Benanti, P. (2020, June 21). *Homo sapiens e macchina sapiens*. UCSI.  
<https://www.ucsì.it/paolo-benanti-homo-sapiens-e-macchina-sapiens/>
- Benanti, P. (2022). *Human in the loop. Decisioni umane e intelligenza artificiale*. Milano: Mondadori.
- Benanti, P. (2017). *Il lavoro nell'epoca della Macchina sapiens*. Oikonomia.  
<https://www.oikonomia.it/images/pdf/2017/ottobre/03-Benanti.pdf>
- Benanti, Paolo. (2022). *La condizione tecno-umana. Domande di senso nell'era della tecnologia. Nuova edizione*. EDB Bologna.
- Benanti, Paolo. (2018). *Le macchine sapienti : intelligenze artificiali e decisioni umane*. Bologna: Marietti 1820.
- Benanti, Paolo. (2025). *L'uomo è un algoritmo? Il senso dell'umano e l'intelligenza artificiale*. Castelveccchi.
- Bourdieu, P. (1991). *Language and symbolic power*. Polity Press.  
[https://www.miguelangelmartinez.net/IMG/pdf/1991\\_bourdieu\\_language\\_ch1.pdf](https://www.miguelangelmartinez.net/IMG/pdf/1991_bourdieu_language_ch1.pdf)
- Boyatzis, R. E. (1982). *The competent manager: A model for effective performance*. John Wiley & Sons.  
[https://www.researchgate.net/publication/247813294\\_The\\_Competent\\_Manager\\_A\\_Model\\_For\\_Effective\\_Performance](https://www.researchgate.net/publication/247813294_The_Competent_Manager_A_Model_For_Effective_Performance)
- Bruno, M., Cerase, A., Binotto, M., Dominici, P., Maggi, S. P., Olivieri, C., & Picone, S. (2008). *Oltre la discriminazione*. Edigraf.  
[https://www.lavoro.gov.it/sites/default/files/archivio-doc-pregressi/AreaSociale\\_Immigrazione/Focus\\_N3\\_Allegato7\\_Oltre\\_la\\_discriminazione.pdf](https://www.lavoro.gov.it/sites/default/files/archivio-doc-pregressi/AreaSociale_Immigrazione/Focus_N3_Allegato7_Oltre_la_discriminazione.pdf)
- Bryman, A., Collinson, D., Grint, K., Jackson, B., & Uhl-Bien, M. (a cura di). (2011). *The SAGE Handbook of Leadership*. SAGE Publications Ltd.
- Burns, T., & Stalker, G. M. (1961). *The management of innovation*. Tavistock.  
<https://turkusowesniadania.pl/wp-content/uploads/2018/11/The-Management-of-Innovation.pdf>
- Cantoni Mamiani, V. (2021). *Leadership di cura: Dal controllo alle relazioni*. Milano: Vita e Pensiero.
- Caragnano, R., & Garofalo, D. (2025). I nuovi paradigmi del lavoro tra digitalizzazione, intelligenza artificiale e metaverso. *ADAPT Labour Studies e-Book Series*, n. 108. ADAPT University Press.

[https://www.adaptuniversitypress.it/wp-content/uploads/2025/10/vol\\_108\\_2025\\_Caragnano.pdf](https://www.adaptuniversitypress.it/wp-content/uploads/2025/10/vol_108_2025_Caragnano.pdf)

Christiano, P. F., Leike, J., Brown, T. B., Martic, M., Legg, S., & Amodei, D. (2017). *Deep reinforcement learning from human preferences*. arXiv.

<https://doi.org/10.48550/arXiv.1706.03741>

Cnaan, R. A., & Goldberg-Glen, R. S. (1991). Measuring motivation to volunteer in human services. *Journal of Applied Behavioral Science*, 27(3), 269–284.

<https://psycnet.apa.org/doi/10.1177/0021886391273003>

Costa, P. T., & McCrae, R. R. (1992). *Revised NEO Personality Inventory (NEO PI-R) and NEO Five-Factor Inventory (NEO-FFI) professional manual*. Psychological Assessment Resources. <https://psycnet.apa.org/doi/10.4135/9781849200479.n9>

Covey, S. R. (2023). *Fiducia e ispirazione. Come i veri leader liberano la grandezza negli altri*. Milano: Franco Angeli.

Covey, S. R. (2021). *Le 7 regole per avere successo: The 7 habits of highly effective people* (4<sup>a</sup> ed.). Milano: Franco Angeli.

Crowne, D. P., & Marlowe, D. (1960). A new scale of social desirability independent of psychopathology. *Journal of Consulting Psychology*, 24(4), 349–354.

<https://doi.org/10.1037/h0047358>

Cruciani, F. (2023). Intelligenza artificiale e intelligenza umana nell'interazione tra regole e creatività. *Filosofi(e)Semiotiche*, 10(1), 15–22.

<https://www.ilsileno.it/filosofiesemiotiche/wp-content/uploads/2023/08/2-Cruciani.pdf>

Daugherty, P. R., & Wilson, H. J. (2018). *Human + machine: Reimagining work in the age of AI*. Harvard Business Review Press.

Deci, E. L., & Ryan, R. M. (2000). The "what" and "why" of goal pursuits: Human needs and the self-determination of behavior. *Psychological Inquiry*, 11(4), 227–268.

[https://doi.org/10.1207/S15327965PLI1104\\_01](https://doi.org/10.1207/S15327965PLI1104_01)

Di Dio Roccazzella, M., & Pagano, F. (2023). *Intelligenza artificiale: Arte e scienza nel business*. Il Sole 24 Ore.

Dobelli, R. (2014). *The Art of Thinking Clearly*. HarperCollins.

<https://xqdoc.imedao.com/166eb7278f35556e3fe9dc3ef.pdf>

Drucker, P. F. (2008). *Management*. Harper Business.

- Drucker, P. F. (1963). *Managing for Business Effectiveness*. Harvard Business Review. <https://psycnet.apa.org/record/1965-08862-001>
- Drucker, P. F. (1966). *The Effective Executive: The Definitive Guide to Getting the Right Things Done*. Harper Business.  
<https://dtleadership.my/wp-content/uploads/2019/05/Drucker-2006-The-Effective-Executive-The-Definitive-Guide-to-Getting-the-Right-Things-Done.pdf>
- Drucker, P. F. (2011). *The practice of management*. Routledge.  
[https://api.pageplace.de/preview/DT0400.9781136356223\\_A23840485/preview-9781136356223\\_A23840485.pdf](https://api.pageplace.de/preview/DT0400.9781136356223_A23840485/preview-9781136356223_A23840485.pdf)
- Edmondson, A. C. (2012). Teamwork on the fly. *Harvard Business Review*, 90(4), 72–80. <https://hbr.org/2012/04/teamwork-on-the-fly-2>
- Fairhurst, G. T. (2007). *Discursive Leadership: In Conversation with Leadership Psychology*. Sage. <https://psycnet.apa.org/record/2007-00682-000>
- Fairhurst, G. T., & Connaughton, S. L. (2014). Leadership: A communicative perspective. *Leadership*, 10(1), 7-35.  
[https://www.researchgate.net/publication/275476854\\_Leadership\\_A\\_communicative\\_perspective](https://www.researchgate.net/publication/275476854_Leadership_A_communicative_perspective)
- Farayola, O. A., & Olorunfemi, O. L. (2024). Ethical decision-making in IT governance: A review of models and frameworks. *International Journal of Science and Research Archive*, 11(2), 130–138. <https://doi.org/10.30574/ijrsra.2024.11.2.0373>
- Floridi, L. (2021). *Ethics of Artificial Intelligence*. Oxford University Press.
- Floridi, L. (2016). On human dignity as a foundation for the right to privacy. *Philosophy & Technology*, 29(4), 307-312.  
[https://www.researchgate.net/publication/301640092\\_On\\_Human\\_Dignity\\_as\\_a\\_Foundation\\_for\\_the\\_Right\\_to\\_Privacy](https://www.researchgate.net/publication/301640092_On_Human_Dignity_as_a_Foundation_for_the_Right_to_Privacy)
- Floridi, L. (2013). *The Ethics of Information*. Oxford University Press.  
<https://doi.org/10.1093/acprof:oso/9780199641321.001.0001>
- Floridi, L. (2014). *The Fourth Revolution: How the Infosphere is Reshaping Human Reality*. Oxford University Press.
- Floridi, L. (2019). *The Logic of Information*. Oxford University Press.  
<https://doi.org/10.1093/oso/9780198833635.001.0001>

Floridi, L. (2019). Translating principles into practices of digital ethics: Five risks of being unethical. *Philosophy & Technology*, 32(2), 185–193.

[https://www.researchgate.net/publication/333332997\\_Translating\\_Principles\\_into\\_Practices\\_of\\_Digital\\_Ethics\\_Five\\_Risks\\_of\\_Being\\_Unethical](https://www.researchgate.net/publication/333332997_Translating_Principles_into_Practices_of_Digital_Ethics_Five_Risks_of_Being_Unethical)

Fombrun, C. J. (1996). *Reputation: Realizing value from the corporate image*. Harvard Business School Press. <https://archive.org/details/reputationrealiz0000fomb>

Gama, J., Žliobaitė, I., Bifet, A., Pechenizkiy, M., & Bouchachia, A. (2014). A survey on concept drift adaptation. *ACM Computing Surveys*, 46(4), Article 44. <https://doi.org/10.1145/2523813>

Gartner. (2025, November 10). *Gartner survey finds artificial intelligence will touch all information technology work by 2030* [Press release].

<https://www.gartner.com/en/newsroom/press-releases/2025-11-10-gartner-survey-finds-artificial-intelligence-will-touch-all-information-technology-work-by-2030>

Gartner. (2024, May 7). *Gartner survey finds generative AI is now the most frequently deployed AI solution in organizations* [Press release].

<https://www.gartner.com/en/newsroom/press-releases/2024-05-07-gartner-survey-finds-generative-ai-is-now-the-most-frequently-deployed-ai-solution-in-organizations>

Giamminola, G. (2024). *Il manager potenziato: Come l'intelligenza artificiale rende il management più efficace, creativo e strategico*. StreetLib.

Goffman, E. (2003). *La vita quotidiana come rappresentazione* (M. Ciacci, Trad.). Il Mulino. (Opera originale pubblicata nel 1959)

Goldberg, L. R. (2006). Doing it all Bass-Ackwards: The development of hierarchical factor structures from the top down. *Journal of Research in Personality*, 40, Issue 4, August 2006, 347-358, <https://doi.org/10.1016/j.jrp.2006.01.001>

Goleman, D. (2011). *Intelligenza emotiva: Che cos'è e perché può renderci felici* (4<sup>a</sup> ed.). Milano: Rizzoli.

Goleman, D. (2000). *Lavorare con intelligenza emotiva: Come inventare un nuovo rapporto con il lavoro*. Milano: Rizzoli.

Goodfellow, I., Bengio, Y., & Courville, A. (2016). *Deep learning* (Chapter 5: Machine learning basics). MIT Press. [https://users.wpi.edu/~kmus/ECE579M\\_files/ReadingMaterials/Read/MachineLearningBasics\\_Goodfellow.pdf](https://users.wpi.edu/~kmus/ECE579M_files/ReadingMaterials/Read/MachineLearningBasics_Goodfellow.pdf)

- Goodfellow, I. J., Mirza, M., Xiao, D., Courville, A., & Bengio, Y. (2013). *An empirical investigation of catastrophic forgetting in gradient-based neural networks*. arXiv. <https://doi.org/10.48550/arXiv.1312.6211>
- Greiner, L. E. (1972). Evolution and revolution as organizations grow. *Harvard Business Review*, 50(4), 37–46.  
<https://hbr.org/1998/05/evolution-and-revolution-as-organizations-grow>
- Grivet, L., Bruni, L., Perrone, G., Rinaldi, R., Frau, S., Mecozzi, R., Pucacco, F., De Vita, F., Franzese, M. F., Oliverio, F., Incani, N., Bakmaz, A., & Maggiano, F. (2025). *Allenare competenze e tecnologie*. EY Advisory S.p.A.  
<https://www.ey.com/content/dam/ey-unified-site/ey-com/it-it/services/consulting/documents/ey-coverciano-2025-paper.pdf>
- Gronn, P. (2002). Distributed leadership as a unit of analysis. *The Leadership Quarterly*, 13(4), 423–451.  
[https://www.researchgate.net/publication/222371529\\_Distributed\\_Leadership\\_as\\_a\\_Unit\\_of\\_Analysis](https://www.researchgate.net/publication/222371529_Distributed_Leadership_as_a_Unit_of_Analysis)
- Hastie, T., Tibshirani, R., & Friedman, J. (2009). *The elements of statistical learning: Data mining, inference, and prediction (2nd ed.)*. Springer.  
<http://www.stat.ucla.edu/~ywu/research/documents/BOOKS/ElementsLearningII.pdf>
- Hatch, M. J., & Schultz, M. (2008). *Taking Brand Initiative: How Companies Can Align Strategy, Culture, and Identity through Corporate Branding*. Jossey-Bass.  
<https://download.e-bookshelf.de/download/0000/5708/34/L-G-0000570834-0002382495.pdf>
- Heifetz, R. A. (1994). *Leadership without easy answers*. Harvard University Press.  
[https://static1.squarespace.com/static/53761774e4b0e8c9b2acd2a4/t/66a2b58f4003c9086cc4117d/1721939345918/Heifetz\\_1994\\_Leadership+without+Easy+Answers\\_Full+Text.pdf](https://static1.squarespace.com/static/53761774e4b0e8c9b2acd2a4/t/66a2b58f4003c9086cc4117d/1721939345918/Heifetz_1994_Leadership+without+Easy+Answers_Full+Text.pdf)
- Iacono, G. (2025). *Come cambia l'e-leadership con l'Intelligenza Artificiale*. Milano: FrancoAngeli.
- Ismail, S., van Geest, Y., & Malone, M. (2015). *Exponential Organizations: Il futuro del business mondiale*. Venezia: Marsilio Editori.
- Janis, I. L. (1982). *Groupthink: Psychological studies of policy decisions and fiascoes (2nd ed.)*. Houghton Mifflin. <https://archive.org/details/groupthinkpsycho00jani>

Joshi, S. (2025). Leadership in the age of AI: Review of quantitative models and visualization for managerial decision-making. *World Journal of Advanced Research and Reviews*, 26(1), 2773–2791. <https://doi.org/10.30574/wjarr.2025.26.1.1415>

Kahneman, D. (2011). *Thinking, fast and slow*. Farrar, Straus and Giroux.  
<https://psycnet.apa.org/record/2011-26535-000>

Kapferer, J.-N. (2012). *The New Strategic Brand Management* (5th ed.). Kogan Page.  
[https://www.academia.edu/12869300/The\\_New\\_Strategic\\_Brand\\_Management\\_Jean\\_N\\_oel\\_Kapferer\\_PDF](https://www.academia.edu/12869300/The_New_Strategic_Brand_Management_Jean_N_oel_Kapferer_PDF)

Laloux, F. (2014). *Reinventing organizations: A Guide to Creating Organizations Inspired by the Next Stage in Human Consciousness*. Nelson Parker.  
[https://www.researchgate.net/publication/341376251\\_Book\\_Review\\_Reinventing\\_Organizations\\_A\\_Guide\\_to\\_Creating\\_Organizations\\_Inspired\\_by\\_the\\_Next\\_Stage\\_in\\_Human\\_Consciousness\\_by\\_Frederic\\_Laloux](https://www.researchgate.net/publication/341376251_Book_Review_Reinventing_Organizations_A_Guide_to_Creating_Organizations_Inspired_by_the_Next_Stage_in_Human_Consciousness_by_Frederic_Laloux)

Latour, B. (2005). *Reassembling the social: An introduction to actor-network-theory*. Oxford University Press. <https://doi.org/10.1093/oso/9780199256044.001.0001>

Lead-Alliance. (s.d.). *Presentazione della seconda fase di Lead-Alliance* [Documento interno di progetto]. Appendice A.

Lee, J. D., & See, K. A. (2004). Trust in automation: Designing for appropriate reliance. *Human Factors*, 46(1), 50-80. [https://doi.org/10.1518/hfes.46.1.50\\_30392](https://doi.org/10.1518/hfes.46.1.50_30392)

Liu, L. (2025). The relationship between emotional trust and cultural identity in AI-created content: A case study of Chinese Australians. *Journalism and Mass Communication*, 15(5), 304–312.  
<https://www.davidpublisher.com/Public/uploads/Contribute/6927c224bb902.pdf>

Lievens, F., & Sackett, P. R. (2012). The validity of interpersonal skills assessment via situational judgment tests for predicting academic success and job performance. *Journal of Applied Psychology*, 97(2), 460–468. <https://doi.org/10.1037/a0025761>

Lu, J., Liu, A., Dong, F., Gu, F., Gama, J., & Zhang, G. (2019). Learning under concept drift: A review. *IEEE Transactions on Knowledge and Data Engineering*, 31(12), 2346–2363. <https://doi.org/10.1109/TKDE.2018.2876857>

McCloskey, M., & Cohen, N. J. (1989). Catastrophic interference in connectionist networks: The sequential learning problem. In G. H. Bower (Ed.), *Psychology of*

*learning and motivation* (Vol. 24, pp. 109–165). Academic Press.

[https://doi.org/10.1016/S0079-7421\(08\)60536-8](https://doi.org/10.1016/S0079-7421(08)60536-8)

McCrae, R. R., & Costa, P. T. (2008). The five-factor theory of personality. In O. P. John, R. W. Robins, & L. A. Pervin (Eds.), *Handbook of personality: Theory and research* (3rd ed., pp. 159–181). Guilford Press.

<https://psycnet.apa.org/record/2008-11667-005>

McDaniel, M. A., Morgeson, F. P., Finnegan, E. B., Campbell, J. P., & Braverman, E. P. (2001). Predicting job performance using situational judgment tests: A clarification of the literature. *Journal of Applied Psychology*, 86(4), 730–740.

<https://doi.org/10.1037/0021-9010.86.4.730>

Mehrabi, N., Morstatter, F., Saxena, N., Lerman, K., & Galstyan, A. (2021). A survey on bias and fairness in machine learning. *ACM Computing Surveys*, 54(6), Article 115. <https://doi.org/10.1145/3457607>

Mintzberg, H. (2009). *Managing*. San Francisco, CA: Berrett-Koehler.

Mintzberg, H. (1983). *Structure in fives: Designing effective organizations*. Englewood Cliffs, NJ: Prentice-Hall. <https://doi.org/10.2307/2393181>

Mintzberg, H. (1994). *The rise and fall of strategic planning*. Free Press. <https://theism.org/documents/Mintzberg%20%281994%29%20Fall%20and%20Rise%20of%20Strategic%20Planning.pdf>

Mintzberg, H. (1979). *The structuring of organizations*. Prentice-Hall. [https://doi.org/10.1007/978-1-349-20317-8\\_23](https://doi.org/10.1007/978-1-349-20317-8_23)

Molaschi, V. (2022). *Algoritmi e discriminazione*. Politecnico di Torino. [https://iris.polito.it/retrieve/8a3cbe23-cc03-41ed-82f9-5fe60ea31af6/Molaschi\\_algorithmi\\_discriminazione\\_2022.pdf](https://iris.polito.it/retrieve/8a3cbe23-cc03-41ed-82f9-5fe60ea31af6/Molaschi_algorithmi_discriminazione_2022.pdf)

Mollick, E. (2025). *L'intelligenza condivisa. Vivere e lavorare insieme all'AI*. Milano: Egea.

Monreale, A. (2020). Rischi etico-legali dell'Intelligenza Artificiale. *DPCE online*, (3), 3391–3398. <https://www.dpceonline.it/index.php/dpceonline/article/download/1083/1039>

Morgan, G. (1998). *Images of Organization*. Thousand Oaks, CA: Sage Publications. <https://parsmodir.com/wp-content/uploads/2019/02/morgan.pdf>

Morini, M. (2024). *NIO: Nuova Intelligenza Organizzativa - Le nuove sfide della leadership nei tempi dell'Intelligenza artificiale*. Bologna: Edizioni Pendragon.

Nonaka, I., & Takeuchi, H. (1995). *The knowledge-creating company: How Japanese companies create the dynamics of innovation*. Oxford University Press.

[https://www.researchgate.net/publication/384316693\\_The\\_knowledge-creating\\_company\\_How\\_Japanese\\_companies\\_create\\_the\\_dynamics\\_of\\_innovation\\_by\\_Nonaka\\_Ikujiro\\_Takeuchi\\_Hirotaka\\_New\\_York\\_Oxford\\_University\\_Press\\_1995\\_284\\_pp\\_1939\\_Hardcover\\_740\\_paperback\\_IS](https://www.researchgate.net/publication/384316693_The_knowledge-creating_company_How_Japanese_companies_create_the_dynamics_of_innovation_by_Nonaka_Ikujiro_Takeuchi_Hirotaka_New_York_Oxford_University_Press_1995_284_pp_1939_Hardcover_740_paperback_IS)

Northouse, P. G. (2019). *Leadership: Theory and practice* (8<sup>a</sup> ed.). Thousand Oaks, CA: Sage Publications.

[https://www.academia.edu/30456341/Book\\_Leadership\\_Theory\\_and\\_Practice](https://www.academia.edu/30456341/Book_Leadership_Theory_and_Practice)

Ofem, O. (2024). *Ethical integration of AI into organizational behavior: Introducing the AI-IOB model*. Research Square. <https://doi.org/10.21203/rs.3.rs-5272515/v1>

Orlikowski, W. J. (2007). Sociomaterial practices: Exploring technology at work. *Organization Studies*, 28(9), 1435–1448.

[https://www.researchgate.net/publication/247734657\\_Sociomaterial\\_Practices\\_Exploring\\_Technology\\_at\\_Work](https://www.researchgate.net/publication/247734657_Sociomaterial_Practices_Exploring_Technology_at_Work)

Osgood, C. E., Suci, G. J., & Tannenbaum, P. H. (1957). *The measurement of meaning*. University of Illinois Press.

<https://gwern.net/doc/psychology/1957-osgood-themeasurementofmeaning.pdf>

Ouyang, L., Wu, J., Jiang, X., Almeida, D., Wainwright, C., Mishkin, P., Zhang, C., Agarwal, S., Slama, K., Ray, A., et al. (2022). *Training language models to follow instructions with human feedback*. arXiv. <https://doi.org/10.48550/arXiv.2203.02155>

Palo Alto Networks. (n.d.). *What problems does black box AI cause?* Cyberpedia. <https://www.paloaltonetworks.com/cyberpedia/black-box-ai>

Project Management Institute. (2021). *A guide to the project management body of knowledge (PMBOK® Guide)* (7th ed.). Newtown Square, PA: PMI. <https://tegnum.edu.pe/wp-content/uploads/2023/09/Project-Management-Institute-A-Guide-to-the-Project-Management-Body-of-Knowledge-PMBOK-R-Guide-PMBOK%C2%AE%E8%8F-Guide-Project-Management-Institute-2021.pdf>

Pulizzi, J. (2013). *Epic content marketing: How to Tell a Different Story, Break through the Clutter, and Win More Customers by Marketing Less*. McGraw-Hill.

<https://contentburger.co/wp-content/uploads/2019/02/Content-Burger-Pulizzi-Joe-Epic-content-marketing.pdf>

Qwaider, S. R., Abu-Saqer, M. M., Albatish, I., Alsaqqa, A. H., Abunasser, B. S., & Abu-Naser, S. S. (2024). Harnessing artificial intelligence for effective leadership: Opportunities and challenges. *International Journal of Academic Information Systems Research (IJAISR)*, 8(8), 9–15.

<http://ijeais.org/wp-content/uploads/2024/8/IJAISR240802.pdf>

Rammstedt, B., & Beierlein, C. (2014). Big Five Inventory-SOEP (BFI-S): Zusammenstellung sozialwissenschaftlicher Items und Skalen (ZIS). *Zusammenstellung sozialwissenschaftlicher Items und Skalen*. <https://doi.org/10.6102/zis54>

Rammstedt, B., & John, O. P. (2007). Measuring personality in one minute or less: A 10-item short version of the Big Five Inventory in English and German. *Journal of Research in Personality*, 41(1), 203–220. <http://dx.doi.org/10.1016/j.jrp.2006.02.001>

Ries, A., & Trout, J. (1981). *Positioning: the battle for your mind*. N. Y. : Warner Books, a Warner Communications Company.

[https://archive.org/details/positioningbattl0000ries\\_t0v3](https://archive.org/details/positioningbattl0000ries_t0v3)

Russell, S. (2022). Human compatible: Artificial intelligence and the problem of control. *Perspectives on Digital Humanism*, 19-24.

[https://www.researchgate.net/publication/356505374\\_Artificial\\_Intelligence\\_and\\_the\\_Problem\\_of\\_Control](https://www.researchgate.net/publication/356505374_Artificial_Intelligence_and_the_Problem_of_Control)

Shneiderman, B. (2022). *Human-Centered AI*. Oxford University Press.

<https://doi.org/10.1093/oso/9780192845290.001.0001>

Siau, K., & Wang, W. (2018). Building trust in artificial intelligence, machine learning, and robotics. *Cutter Business Technology Journal*, 31(2), 47-53.

[https://www.researchgate.net/publication/324006061\\_Building\\_Trust\\_in\\_Artificial\\_Intelligence\\_Machine\\_Learning\\_and\\_Robotics](https://www.researchgate.net/publication/324006061_Building_Trust_in_Artificial_Intelligence_Machine_Learning_and_Robotics)

Skarzyńska, E. (2025). Leadership in the age of AI in the context of effectively connecting technology to teams. *Scientific Papers of Silesian University of Technology*, 229, 475–490. <https://doi.org/10.29119/1641-3466.2025.229.27>

Spencer, L. M., & Spencer, S. M. (1993). *Competence at work: Models for superior performance*. Wiley.

Stephens, M., Zalabak, M., Havens, J. C., Messler, H., Allgood, M., Berne, R. W., Bovell, S., Bridgesmith, L., Calvo, R. A., Duin, A. H., Felice, J., Harrod, J., Kurihara, K., Pedersen, I., Pena-Fernandez, A., Polgar, D. R., Prestes, E., & Whitaker, J. (2020). *The benefits of a multidisciplinary lens for artificial intelligence systems ethics*. IEEE Standards Association.

<https://standards.ieee.org/wp-content/uploads/2023/07/ead-benefits-multidisciplinary-lens.pdf>

Stinchcombe, A.L. (1965) Social Structure and Organizations. In: March, J.P., Ed., *Handbook of Organizations*, Rand McNally, Chicago, 142-193.  
[https://www.researchgate.net/publication/367369019\\_Social\\_Structure\\_and\\_Organizations](https://www.researchgate.net/publication/367369019_Social_Structure_and_Organizations)

Stryker, C. (n.d.). *What is human-in-the-loop?* IBM Think.  
<https://www.ibm.com/think/topics/human-in-the-loop>

Thaler, R. H., & Sunstein, C. R. (2008). Nudge: Improving decisions about health, wealth, and happiness. *The Social Science Journal*, 45(4), 700–701.  
[https://www.researchgate.net/publication/257178709\\_Nudge\\_Improving\\_Decisions\\_About\\_Health\\_Wealth\\_and\\_Happiness\\_RH\\_Thaler\\_CR\\_Sunstein\\_Yale\\_University\\_Press\\_New\\_Haven\\_2008\\_293\\_pp](https://www.researchgate.net/publication/257178709_Nudge_Improving_Decisions_About_Health_Wealth_and_Happiness_RH_Thaler_CR_Sunstein_Yale_University_Press_New_Haven_2008_293_pp)

Tolbert, P. S., & Zucker, L. G. (1996). The institutionalization of institutional theory. In S. R. Clegg, C. Hardy, & W. R. Nord (Eds.), *Handbook of organization studies* (pp. 175–190). Sage. <https://files01.core.ac.uk/download/pdf/5132446.pdf>

Torre, F., et al. (2025). *AI leadership for corporate boards: Leading responsible AI for value creation*. Springer.  
[https://www.researchgate.net/publication/396897368\\_AI\\_Leadership\\_for\\_Corporate\\_Boards\\_Leading\\_Responsible\\_AI\\_for\\_Value\\_Creation](https://www.researchgate.net/publication/396897368_AI_Leadership_for_Corporate_Boards_Leading_Responsible_AI_for_Value_Creation)

Trist, E. L., & Bamforth, K. W. (1951). Some social and psychological consequences of the longwall method of coal-getting. *Human Relations*, 4(1), 3–38.  
<https://doi.org/10.1177/001872675100400101>

Turkle, S. (2011). *Alone Together: Why We Expect More from Technology and Less from Each Other*. New York: Basic Books.  
[https://www.researchgate.net/publication/50382537\\_Alone\\_Together\\_Why\\_We\\_Expect\\_More\\_from\\_Technology\\_and\\_Less\\_From\\_Each\\_Other](https://www.researchgate.net/publication/50382537_Alone_Together_Why_We_Expect_More_from_Technology_and_Less_From_Each_Other)

- Turkle, S. (2017). *Reclaiming Conversation: The Power of Talk in a Digital Age*. Penguin Press.  
[https://www.researchgate.net/publication/350521529\\_Reclaiming\\_Conversation\\_The\\_Power\\_of\\_Talk\\_in\\_a\\_Digital\\_Age](https://www.researchgate.net/publication/350521529_Reclaiming_Conversation_The_Power_of_Talk_in_a_Digital_Age)
- Umoren, O., Didi, P. U., Balogun, O., Abass, O. S., & Akinrinoye, O. V. (2022). Strategic digital storytelling techniques for building authentic brand narratives and driving cross-generational consumer trust online. *Gyanshauryam: International Scientific Refereed Research Journal*, 5(3), 238-262.  
<https://gisrrj.com/paper/GISRRJ225325.pdf>
- Unione Europea. (2024). *Regolamento (UE) 2024/1689 del Parlamento europeo e del Consiglio, del 13 giugno 2024, che stabilisce norme armonizzate sull'intelligenza artificiale (AI Act)*. Gazzetta Ufficiale dell'Unione Europea, L, 2024/1689.  
<https://eur-lex.europa.eu/eli/reg/2024/1689/oj>
- Veratti, E., & Barin, A. (2025, January 15). *Ethics in AI. Un manifesto per l'intelligenza artificiale etica*. Bain & Company.  
<https://www.bain.com/contentassets/6d1d47eb578546bb939833f984a36a6b/ethics-in-ai.pdf>
- Watzlawick, P., Beavin, J. H., & Jackson, D. D. (1967). *Pragmatics of human communication: A study of interactional patterns, pathologies, and paradoxes*. W. W. Norton & Company.  
<https://comm163v.wordpress.com/wp-content/uploads/2013/01/bavelas.pdf>
- Weber, M. (1978). *Economy and society: An outline of interpretive sociology* (G. Roth & C. Wittich, Eds.). University of California Press. (Original work published 1922).  
[https://www.researchgate.net/publication/228007544\\_Economy\\_and\\_Society\\_An\\_Outline\\_of\\_Interpretive\\_Sociology\\_2\\_Vols](https://www.researchgate.net/publication/228007544_Economy_and_Society_An_Outline_of_Interpretive_Sociology_2_Vols)
- Weick, K. E. (1995). *Sensemaking in organizations*. Sage Publications.
- Weick, K. E., Sutcliffe, K. M., & Obstfeld, D. (2005). Organizing and the process of sensemaking. *Organization Science*, 16(4), 409–421.  
<https://doi.org/10.1287/orsc.1050.0133>
- Weizenbaum, J. (1976). *Computer power and human reason: From judgment to calculation*. W. H. Freeman.

[https://www.researchgate.net/publication/286058724\\_Computer\\_Power\\_and\\_Human\\_Reason\\_From\\_Judgment\\_to\\_Calculation](https://www.researchgate.net/publication/286058724_Computer_Power_and_Human_Reason_From_Judgment_to_Calculation)

Weizenbaum, J. (1966). ELIZA—A computer program for the study of natural language communication between man and machine. *Communications of the ACM*, 9(1), 36–45. <https://doi.org/10.1145/365153.365168>

Wenger, E. (1998). *Communities of practice: Learning, meaning, and identity*. Cambridge University Press.  
[https://www.researchgate.net/publication/225256730\\_Wenger\\_E\\_1998\\_Communities\\_of\\_practice\\_Learning\\_meaning\\_and\\_identity](https://www.researchgate.net/publication/225256730_Wenger_E_1998_Communities_of_practice_Learning_meaning_and_identity)

Werner, J., World Economic Forum, OECD, Microsoft Research, Stanford HAI, Richards, D., Microsoft & LinkedIn, McKinsey & Company, Deloitte, UNESCO, & NIST. (2024). *AI literacy for everyone: Building trust, resilience, and competitive advantage*. [https://ldi.njit.edu/sites/ldi/files/AI%20Literacy%20for%20Everyone\\_0.pdf](https://ldi.njit.edu/sites/ldi/files/AI%20Literacy%20for%20Everyone_0.pdf)

Zarifis, A. (2025). *Leadership with AI and trust: Adapting popular leadership styles for AI*. De Gruyter. <https://doi.org/10.1515/9783111630137>

## SITOGRAFIA

Altermaind. (n.d.). *Home page*. <https://www.altermaind.com/> (Ultima Consultazione: 19 febbraio 2026)

Altermaind. (n.d.). *Slash*. <https://www.altermaind.com/slash> (Ultima Consultazione: 19 febbraio 2026)

Intelligenza Artificiale Spiegata Semplice. (n.d.). *Entra nella Community sull'Intelligenza Artificiale più grande d'Italia*. <https://www.iaspiegatasemplice.it/> (Ultima Consultazione: 19 febbraio 2026)

JECoMM – Junior Enterprise dell'Università degli Studi di Milano. (n.d.). *JECoMM*. <https://www.jecomm.it/> (Ultima Consultazione: 19 febbraio 2026)

Lead-Alliance. (n.d.). *Lead-Alliance*. <https://www.leadalliance-ai.com/> (Ultima Consultazione: 13 marzo 2026)

Lead-Alliance. (25 Febbraio 2026). *Questionario Lead-Alliance: Misurazione del Sentiment*. <https://forms.gle/WBFtiJnHdgAkvcWS7> (Ultima Consultazione: 27 febbraio 2026)

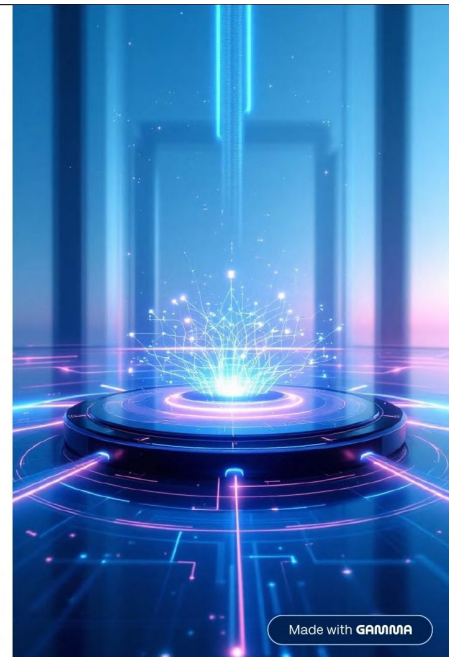
LinkedIn. (n.d.). *Lead-Alliance*. <https://www.linkedin.com/company/lead-alliance-ai/about/?viewAsMember=true> (Ultima Consultazione: 13 marzo 2026)

OpenHS. (n.d.). *OpenHS*. <https://www.openhs.it/> (Ultima Consultazione: 27 febbraio 2026)

## APPENDICE A: Presentazione della seconda fase di Lead-Alliance

### Lead-Alliance: i prossimi passi

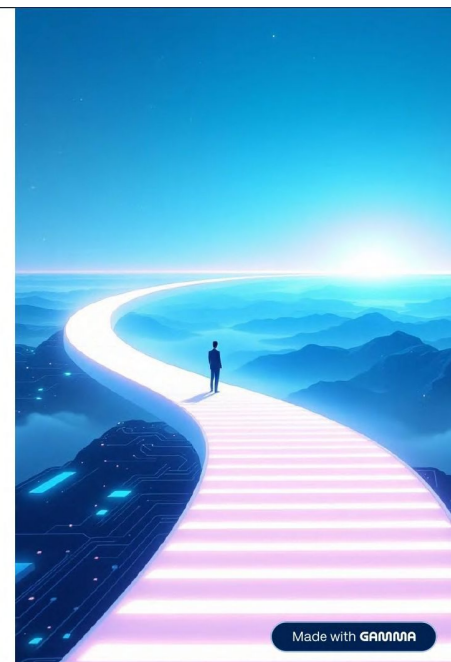
Costruiamo insieme un hub innovativo per trasformare il futuro della leadership attraverso l'intelligenza artificiale.



### La Nostra Unicità

#### Cosa ci distingue nel panorama AI-Leadership?

La combinazione di **competenza tecnica**, **esperienza manageriale** e **visione strategica** dei membri Lead Alliance ci posiziona come il ponte ideale tra innovazione tecnologica e leadership pratica.



## Due Livelli di Crescita Strategica

### Primo Livello

#### Fattibilità Immediata

Contenuti digitali, eventi online e pubblicazioni per costruire autorevolezza nel settore AI-Leadership.

### Secondo Livello

#### Visione a Lungo Termine

Piattaforma completa, corsi strutturati e master per diventare punto di riferimento globale.

La nostra roadmap si ispira ai modelli di successo di Singularity University e alle teorie delle *Exponential Organizations* di Ismail.

Made with GAMMA

## Primo Livello: Contenuti e Comunità

### 1 Podcast Specializzato

Serie mensile "AI Leadership Explained" per democratizzare le conoscenze avanzate con linguaggio accessibile.

### 2 Eventi Online Mensili

Webinar interattivi con leader del settore, casi studio e sessioni Q&A per la community.

### 3 White Paper Autoriale

Pubblicazione annuale che consolida le migliori pratiche e tendenze emergenti dell'AI-Leadership.

### 4 Newsletter Esclusiva

Aggiornamenti settimanali su innovazioni, strumenti pratici e opportunità per i membri Lead Alliance.



Made with GAMMA



## Secondo Livello:

### 1) La Piattaforma Digitale



#### "AI Play" Style Platform

Spazio digitale collaborativo per contenuti multimediali, ricerche e progetti condivisi tra appassionati di AI-Leadership.



#### Pillole Video Periodiche

Contenuti brevi e coinvolgenti per mantenere alta l'attenzione su trend emergenti e casi di successo.

Made with GAMMA



### 2) Formazione Strutturata



#### Corso Base AI-Leadership

Modulo e-learning SCORM con spunti pratici, casi d'uso reali e applicazioni concrete per manager moderni.



#### Master Specialistico

Percorso avanzato con faculty interna Lead Alliance, posizionato strategicamente nel mercato della formazione executive.

Made with GAMMA

### 3) Networking e Connessioni

#### Modello Intro.co

Piattaforma dedicata alle connessioni professionali tra esperti di AI e leadership, facilitando scambi di competenze e collaborazioni strategiche.

- Matching intelligente tra profili
- Condivisione expertise specializzate
- Opportunità di mentoring



Made with GAMMA

### 4) Ecosistema di Conoscenza



#### Archivio Dati e Ricerche

Database centralizzato con protocolli di ricerca condivisi per informazioni affidabili e fruibili.



#### Spazio per le Storie

Testimonianze reali di organizzazioni che applicano AI nella pratica quotidiana con successo.



#### Forum Collaborativo

Ambiente di confronto per co-creare idee e stimolare progettualità comuni tra membri.

Made with GAMMA



## 5) Community-Driven Excellence

01

### Ascolto Attivo

Valorizzazione delle proposte spontanee degli utenti per costruire un percorso realmente partecipato.

02

### Co-creazione

Sviluppo collaborativo di contenuti e iniziative basate sui bisogni reali della community.

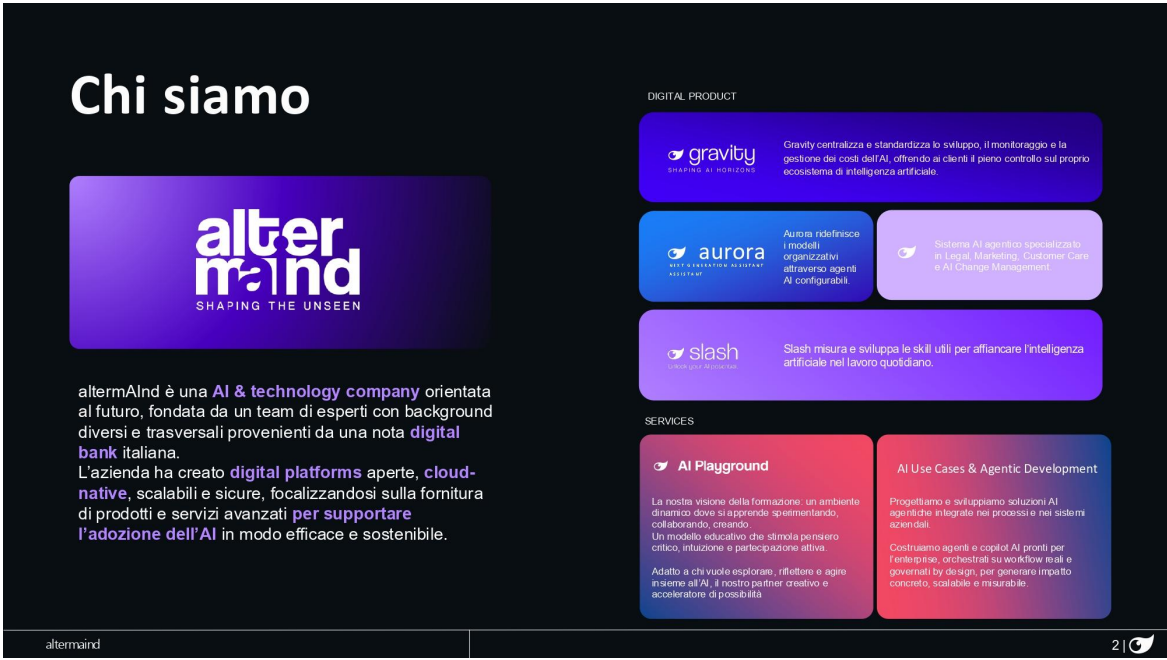
03

### Kit Freemium

Strumenti pratici per organizzazioni, bilanciando valore gratuito e servizi premium.

Made with **GAMMA**

ALLEGATO A: Brochure altermAInd Slash



## AI READINESS

# A che punto sono le aziende italiane?

altermaind



## AI readiness

L'AI readiness in Italia è ancora bassa: il sistema produttivo resta ai blocchi di partenza.

Solo il **6% delle aziende** utilizza l'intelligenza artificiale in modo trasversale e orientato al futuro. La maggior parte del tessuto imprenditoriale è ancora fermo alla fase di sperimentazione.

Fonte: AI Leadership Survey by altermaind – 231 feedback collected among SME and corporate innovators

### AI-driven innovator

L'AI è parte integrante dell'organizzazione: sistemica, distribuita, con visione futura.

6%

### Strategico

L'adozione dell'AI è guidata da una strategia chiara, con leadership e processi solidi.

25%

### Strutturato

L'AI è presente e organizzata in alcune aree, ma non ancora pienamente integrata.

22%

### In transizione

Interesse crescente, ma approccio frammentato o ancora esplorativo.

5%

### Tradizionale

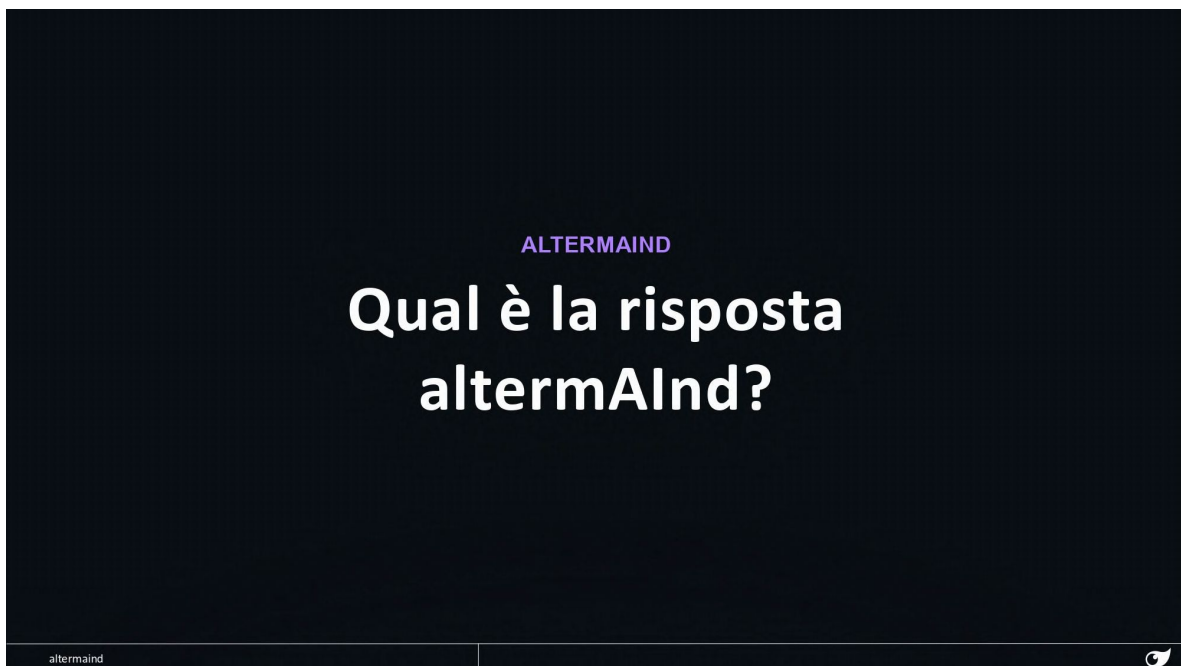
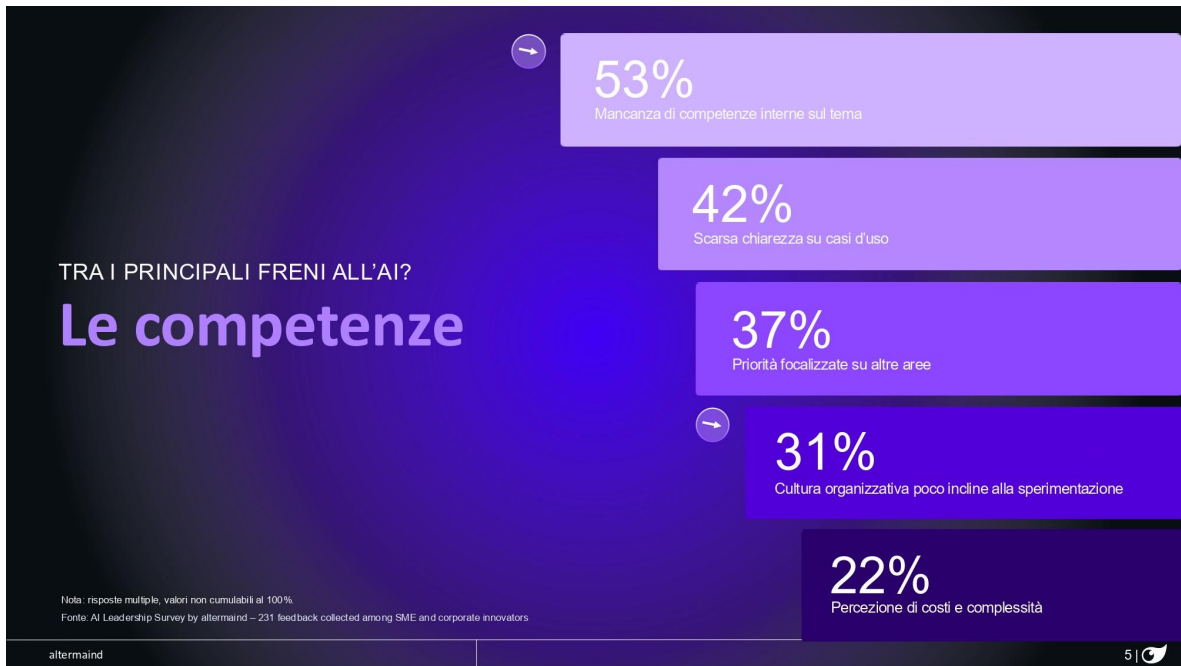
L'AI non è ancora un tema strutturato. Mancano visione, processi e governance.

42%

Livello di AI maturity

altermaind

4 |



CHE COS'È

Slash è la prima piattaforma SaaS  
che aiuta a misurare,  
comprendere e potenziare  
le skills soft e hard per  
lavorare in sinergia con l'AI

altermaind

7 | 

## Progettata per soddisfare le esigenze delle aziende

### Sei un HR?

Monitora in modo integrato le competenze trasversali dei dipendenti attraverso una visualizzazione aggregata di trend, gap e potenziale, per orientare strategie di sviluppo e migliorare la readiness complessiva all'uso dell'AI.

### Sei un manager?

Comprendi la tua readiness verso l'AI e quella delle tue persone, individua punti di forza e aree di squilibrio nei tuoi team e favorisci un apprendimento continuo.

### Sei un dipendente?

Scopri le tue skills per approcciare e utilizzare al meglio l'AI, colma i tuoi gap con la formazione personalizzata e continua a crescere.

Slash è in Beta e disponibile gratuitamente per i primi 3 mesi. Basta solo registrarsi e si può accedere e iniziare a utilizzare (e far utilizzare) la piattaforma

altermaind

8 | 

# Benefici

Fotografia  
chiara e  
condivisa delle  
competenze

Linguaggio  
comune tra HR  
e business

Insight concreti  
per orientare  
priorità

Assessment  
come strumento  
decisionale

altermaind

9 | 

# Come funziona Slash

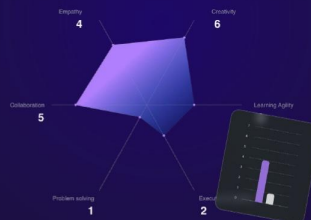
1

Skill assessment



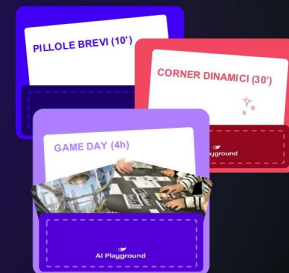
2

Profilo personale



3

Training personalizzato



altermaind

10 | 

# Come funziona Slash

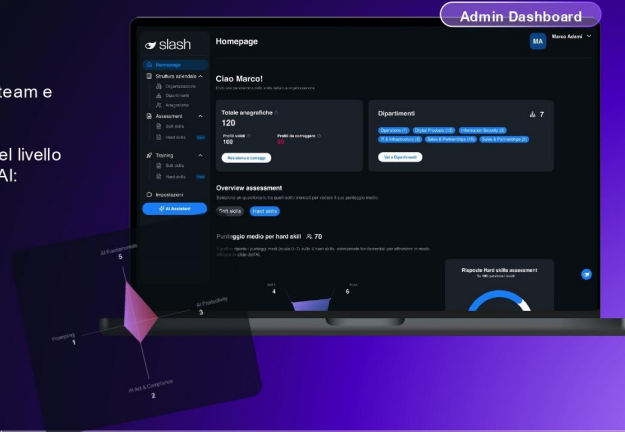
BETA

H1 2026

H2 2026

## BETA SLASH

- Slash gratuito per i primi 3 mesi
- Soft skill assessment con visibilità aziendale, team e singolo dipendente (a discrezione HR)
- **Hard skills assessment generic**: valutazione del livello di maturità nelle competenze tecniche legate all'AI:
  - AI Act & compliance
  - AI fundamentals
  - AI prompting
  - AI productivity



altermaind

11

# Come funziona Slash

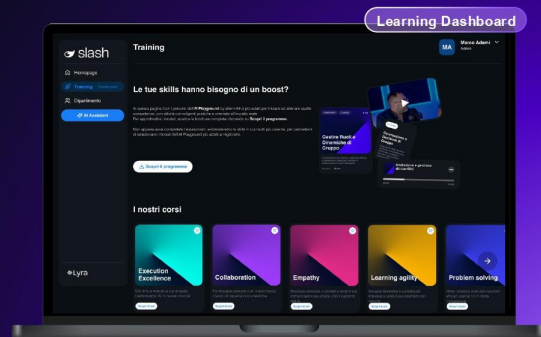
BETA

H1 2026

H2 2026

## BETA SLASH

- Accesso all'**AI Playground** in presenza per lavorare sulle soft skill da migliorare maggiormente
- **Catalogo training online (in arrivo)**: video pillole, manuali di approfondimento, quiz, esercitazioni e simulazioni interattive per trasformare soft e hard skills in performance concrete



altermaind

12

# Come funziona Slash

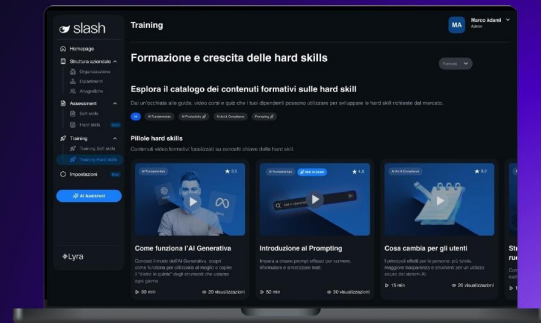
BETA

H1 2026

H2 2026

## EVOLUZIONE COMPETENZE

- **Hard skills assessment role-specific:** valutazione del livello di maturità nelle competenze tecniche legate all'AI specifiche per ruolo per potenziare rapidamente le capacità operative
- **Catalogo training online arricchito:** integrazione di contenuti formativi per skill role-specific
- **Coach AI** per arricchire l'assessment tramite simulazioni e use case connessi all'ambito professionale



altermaind

13 |

# Come funziona Slash

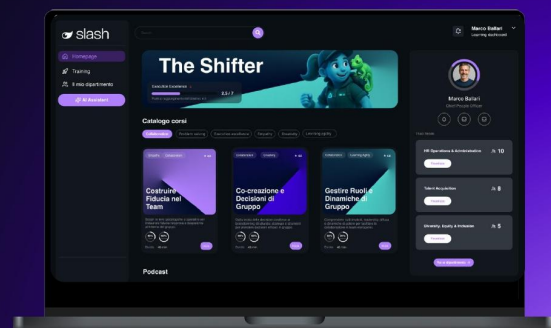
BETA

H1 2026

H2 2026

## EVOLUZIONE FORMAZIONE

- **Expected profile:** confronto tra competenze attuali e ideali per il ruolo per evidenziare gap, punti di forza e costruire percorsi di crescita sempre più personalizzati e allineati alle esigenze del singolo
- **Catalogo training potenziato:** un'ampia selezione di media realizzati da una rete di formatori



altermaind

14 |

01

# SOFT SKILLS

Fornire una diagnosi chiara delle soft skills e attivare un percorso esperienziale che, grazie all'AI, potenzi i comportamenti chiave e riduca rapidamente i gap rilevati.

altermaind



01 SOFT SKILLS ASSESSMENT

## Misura le soft skills

Il questionario (assessment) di Slash analizza sei soft skills chiave, definite dal nostro modello come le più efficaci per lavorare in modo evolutivo con l'AI: Empathy, Creativity, Collaboration, Learning Agility, Problem Solving e Execution Excellence.

L'assessment include domande situazionali che esplorano come le persone tendono ad agire in contesti diversi: ogni risposta contribuisce a costruire una fotografia del modo in cui ciascun utente interpreta e applica le proprie competenze trasversali.

3 su 24 domande

A quale tra queste parole ti senti più vicino:

Familiarità

Scoperta

Immagina di aver iniziato a imparare a suonare da autodidatta. Dopo aver fatto qualche tentativo, ti accorgi che non stai migliorando. Come reagisci?

1 su 24 domande

- ☐ Continuo ad esercitarmi guardando tutorial per capire dove sto sbagliando.
- ☐ Cerco un insegnante per farmi seguire in modo più strutturato.
- ☐ Provo con uno strumento più semplice: imparare a suonare mi interessa, ma forse la chitarra non fa per me.
- ☐ Decido di prendermi del tempo per capire se realmente mi interessa investire del tempo per imparare ad utilizzare la chitarra.

altermaind



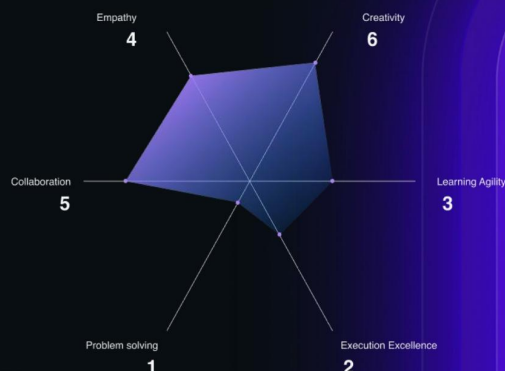
## 02 ACTUAL PROFILE

# Scopri il livello delle competenze

I risultati compongono l'Actual Profile: un insieme di dati oggettivi e insight personalizzati sulle sei soft skills.

Ogni profilo mostra il livello di sviluppo delle competenze, aiutando a individuare punti di forza e aree di crescita.

L'Actual Profile non esprime un giudizio, ma un punto di partenza per la crescita continua: una base su cui costruire percorsi di sviluppo, engagement e miglioramento misurabile nel tempo.



altermaind

## 03 TRAINING PERSONALIZZATO

# AI Playground: la nostra idea di formazione esperienziale

AI Playground è la nostra visione della formazione: **un ambiente dinamico** dove si apprende sperimentando, collaborando, creando.

Un modello educativo che **stimola pensiero critico**, intuizione e partecipazione attiva.

Adatto a chi vuole esplorare, riflettere e agire insieme all'AI, il nostro **partner creativo e acceleratore di possibilità**



altermaind

02

# HARD SKILLS

Misurare il livello di maturità nelle competenze tecniche  
legate all'AI e potenziare rapidamente le capacità  
operative

altermaind



## Hard Skills Assessment generic

**AI essentials: conoscenze base**  
(tipi di modelli, limiti, rischi, capacità)

**AI Act & compliance**  
(rischi, obblighi, classificazione sistemi AI)

**Prompting**  
(framework, contesto, cicli di raffinamento)

**AI per la produttività**  
(scrittura, analisi, riassunto, documenti)

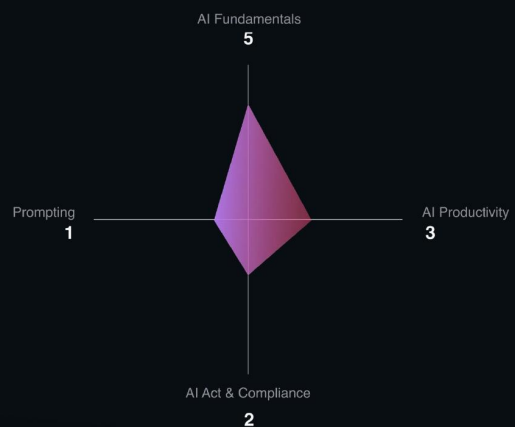
altermaind



# Aree di miglioramento

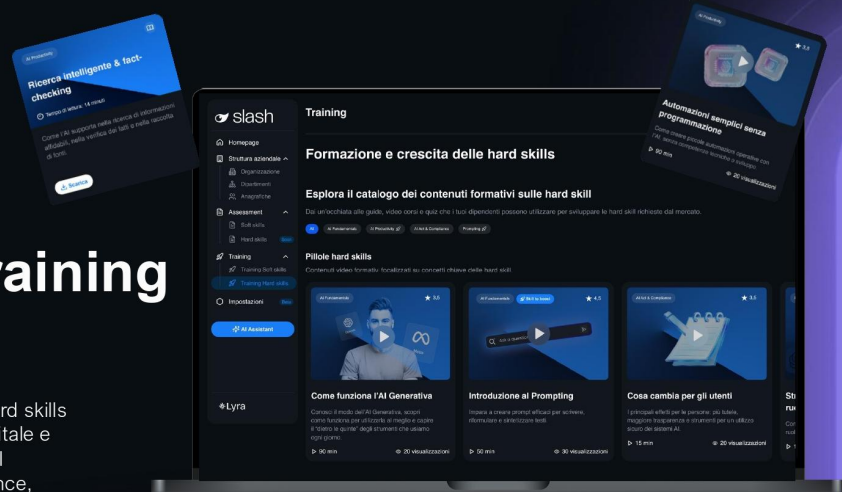
I risultati forniscono una fotografia del livello di conoscenza attuale rispetto ai 4 ambiti precedentemente indicati.

Anche in questo caso, la valutazione non ha una finalità giudicante, ma vuole essere il primo touchpoint su cui avviare percorsi di sviluppo che supportino individui e team in una crescita concreta ed efficace verso l'AI.



# Catalogo training Generic

Quattro moduli per allenare le Hard skills essenziali rispetto alla cultura digitale e utilizzo AI, comuni a tutti i ruoli: AI Fundamentals, AI Act & Compliance, Prompting, Productivity



# Catalogo training Generic

Per ognuno dei 4 KPI hard generic, il catalogo training si compone di:

- ~ **40 video pillole** da 5-6min ciascuna
- ~ **50 domande** che vengono proposte randomicamente tramite quiz interattivo
- ~ **10 paper e documenti** di approfondimento

|                           | Video pillole AI Act & Compliance      |
|---------------------------|----------------------------------------|
| Contesto normativo        | Perché l'AI deve essere regolata       |
|                           | AI e diritti fondamentali              |
| AI Act – le basi          | Cos'è l'AI Act                         |
|                           | A chi si applica l'AI Act              |
|                           | Sistemi AI coperti e non coperti       |
|                           | Approccio risk-based dell'AI Act       |
| Livelli di rischio        | Sistemi AI vietati                     |
|                           | Sistemi AI ad alto rischio             |
| Obblighi e responsabilità | Provider vs deployer                   |
|                           | Documentazione obbligatoria            |
|                           | Sanzioni e conseguenze                 |
|                           | Ruolo del management                   |
| Governance AI             | Cos'è una AI Governance                |
|                           | Modelli di governance AI               |
|                           | AI Risk Management                     |
|                           | Policy e procedure AI                  |
| Trasparenza & XAI         | Cos'è l'Explainable AI                 |
|                           | Trasparenza verso utenti e stakeholder |
|                           | Supervisione umana                     |
|                           | Tracciabilità e auditabilità           |

altermaind

## Hard Skills specific assessment

Le Specific Hard Skills sono le competenze tecniche legate al ruolo e al contesto in cui si lavora. Cambiano quindi da funzione a funzione: chi si occupa di marketing e comunicazione può usare tool per copy, advertising o content automation; chi lavora su prodotto e innovazione utilizza strumenti di AI per ideare, prototipare e testare; mentre nei team di sviluppo l'AI supporta attività come il coding, le analisi predittive e l'automazione dei flussi.

### Marketing

generazione copy, analisi audience, benchmark, campagne

### HR

job description, screening CV, analisi competenze, learning path

### Finance

analisi numeriche, sintesi report, forecasting assistito

### Procurement

analisi fornitori, RFQ, negoziazione assistita, valutazione rischi

### Analisi funzionale

mappatura processi, raccolta requisiti, user story, gap analysis

### Sviluppo / Tech

generazione codice, debug, documentazione tecnica, refactoring

Aree e contenuti in corso di definizione

altermaind



03

# ASSESSMENT BY AI

Evoluzione degli strumenti di assessment per offrire coaching AI-based personalizzati

altermaind



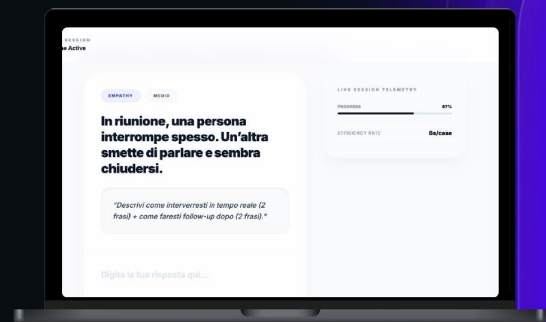
AI-BASED ASSESSMENT

## Sviluppo continuo

L'assessment di Slash si evolve con l'integrazione dell'AI, ampliando l'analisi continua, contestuale e personalizzata delle competenze soft e hard.

Gli strumenti di valutazione diventano sistemi basati su AI, capaci di fornire insight mirati, suggerimenti di sviluppo e valutazioni automatiche basate su comportamenti e risposte.

L'AI permette di usare l'assessment come strumento preliminare per orientare decisioni e ruoli e come base dinamica per aggiornare il profilo nel tempo, supportando la crescita individuale e organizzativa.



altermaind



# Vuoi saperne di più?

Contattaci per attivare la Beta o approfondire insieme la piattaforma



**Giuseppe Montella**

CHIEF DESIGN &  
MARKETING OFFICER

+39 348 2940015

giuseppe.montella@altermaind.com



**Viola Marinelli**

RESEARCH & CUSTOMER  
INSIGHTS

+39 327 5668565

viola.marinelli@altermaind.com



**Clarissa Cardullo**

ACCOUNT & PROJECT  
MANAGER

+39 327 5604264

clarissa.cardullo@altermaind.com



**Stefano Spinelli**

SALES DIRECTOR

+39 335 8293200

stefano.spinelli@altermaind.com





## **DICHIARAZIONE DI ORIGINALITÀ DEL TESTO**

### **Premessa**

Elaborato finale e tesi magistrale devono risultare dal lavoro autonomo del/la candidato/a. L'utilizzo delle fonti primarie e secondarie, in formato cartaceo e/o elettronico, deve essere chiaramente indicato secondo le modalità correnti, distinguendo le citazioni, dirette e indirette, dalle osservazioni e considerazioni del/la candidato/a. Si intendono come citazioni anche le traduzioni da testi pubblicati non in lingua italiana e vanno adeguatamente indicate anche le parafrasi. Dare per propria l'opera di altri, anche con riferimento a parte di opera che venga inserita nella propria senza indicazione della fonte, costituisce plagio.

### **Dichiarazione**

Dichiaro sotto la mia responsabilità che quanto scritto nell'elaborato finale o nella tesi magistrale risulta da elaborazione personale e che, citazioni escluse, nessuna parte è stata copiata da pubblicazioni scientifiche o divulgative in formato cartaceo, elettronico o in lavori prodotti da altri studenti, né tradotta da testi fonte in lingua straniera.

Nel caso di materiale tratto da pubblicazioni scientifiche, da Internet o da altri documenti di cui non sono l'autore, dichiaro sotto la mia responsabilità di averne espressamente e direttamente indicato la fonte alla fine della citazione o a piè di pagina.

Mi assumo questo impegno fino al termine e alla consegna finale del lavoro.

### **Consapevolezza della sanzione**

Sono consapevole che in caso di plagio sono passibile di sanzioni che possono comportare l'impossibilità di laurearmi.

Nome e cognome: Micol Maria Vittoria Bianchi

Numero di matricola: 40996A

Milano il 13/03/2026

Firma



## **DECLARATION OF ACADEMIC HONESTY**

I, the undersigned, hereby declare that I am the author of the present paper and that, except where specifically acknowledged, no parts have been copied from other authors or sources, or from papers previously submitted for assessment by myself or other students.

Any paragraph or portion of text that I have excerpted from a scientific publication, the Internet, or other sources of information has been duly placed in quotation marks and explicitly cited in a clear footnote reference.

Additionally, I declare that I have read and understood the Faculty's provisions with regard to student plagiarism, and am aware of the penalties outlined under the appropriate articles of the current (December 2007) *Student Regulations*.

Student's Name and Surname: Micol Maria Vittoria Bianchi

Student ID card no.: 40996A

Milano the 13/03/2026

Signature

A handwritten signature in black ink, reading "Micol Maria Vittoria Bianchi". The signature is written in a cursive style, with the first name "Micol" and the last name "Bianchi" being more prominent.





## **Ringraziamenti**

Desidero ringraziare con tutto il cuore ogni persona che ha contribuito, in modo diretto o indiretto, alla realizzazione di questo lavoro di ricerca e al mio percorso di studi.

Anzitutto, esprimo la mia più profonda gratitudine al mio relatore, Luigi Di Iorio. Grazie di cuore per la fiducia che continui a riporre in me, per avermi coinvolta in progetti entusiasmanti e per avermi affidato compiti di crescente responsabilità. Grazie anche al prof. Manuel Dileo che, nonostante la sfida di fare da correlatore a una tesi di ambito umanistico, si è dimostrato un valido e attento supporto.

Ringrazio con affetto tutto il team e i miei compagni di Lecco100. Non avrei mai immaginato l'impatto profondo che il master avrebbe avuto sul mio percorso accademico e professionale. Vi sono sinceramente grata. Un ringraziamento speciale va ad Alessio, per i suoi preziosi consigli e gli spunti bibliografici.

Grazie di cuore a Gaia, la mia mentore, a Martina, a Yunfei e a Montse. Siete state figure fondamentali per il mio percorso non solo per quanto riguarda la mia crescita accademica e professionale, ma soprattutto per la motivazione costante che mi avete dato. Senza di voi non sarebbe andata allo stesso modo!

Grazie anche al Centro Linguistico SLAM, a Sonja, ai miei amici cinesi del programma Marco Polo e Turandot. Al di là dell'aspetto lavorativo, vi sono soprattutto grata per ciò che abbiamo condiviso a livello umano.

Un grazie speciale va poi sempre alla mia mamma. Sarai sempre il mio più grande esempio. Grazie anche alle mie nonne, a papà e a Margherita, per il vostro sostegno costante e discreto, che non è mai mancato. Un abbraccio forte va alla mia famiglia tutta, i presenti e gli assenti, per la fiducia che avete sempre riposto in me e nel mio percorso affatto lineare.

Un grazie di cuore ai miei amici, per il vostro sostegno instancabile e la vostra pazienza decennale. Anche quando ero assente, distratta o assorbita dal lavoro o dallo studio,

sapevo di potere contare su di voi. Sarete sempre un punto di riferimento in mezzo a ogni tempesta.

Grazie a tutti i miei colleghi e compagni di corso, che con la loro presenza hanno reso il mio percorso più ricco e stimolante. Un pensiero particolare va a Giulia R., Giulia L. e Arianna: grazie per la vicinanza, le risate, i confronti e la complicità. Avete reso questi due anni meno difficili, e decisamente più belli.

Un grazie immancabile anche a tutti i professori che hanno lasciato un segno in questo mio percorso. Un pensiero particolare va a Irina Bajini, Amy Juan e Clara Bulfoni. In questi anni abbiamo condiviso molto più di un semplice rapporto accademico, e anche per questo vi terrò sempre a cuore come un esempio virtuoso di come si dovrebbe insegnare.

Infine, con tutto il cuore, ringrazio Fabio. Senza di te al mio fianco, a sostenermi ogni singolo giorno, non ce l'avrei mai fatta. La tua presenza costante e piena d'amore è stata la mia forza più grande.