

LUISS



Dipartimento di Economia

Cattedra di Progettazione Organizzativa

L'impatto dell'Intelligenza Artificiale su coordinamento, competenze, divisione del lavoro e strutture organizzative nelle aziende knowledge-based: un'analisi teorico-empirica dei meccanismi organizzativi e dei riflessi sul benessere lavorativo.

Daniele Mascia

RELATORE

Sara Lombardi

CORRELATORE

Elena Neri

781001

CANDIDATO

Anno Accademico 2024/2025

Indice della Tesi

INTRODUZIONE	4
CAPITOLO PRIMO: L'intelligenza artificiale nella progettazione organizzativa	6
1.1 Evoluzione dell'Intelligenza Artificiale nelle organizzazioni	6
1.1.1 Automazione vs. Aumentazione	7
1.1.2 Impatto sui modelli organizzativi	10
1.2 I meccanismi di coordinamento all'interno dell'organizzazione e l'AI	14
1.2.1 Teoria di Mintzberg (1979) e nuovi modelli di coordinamento	16
1.2.2 L'IA come nuovo meccanismo di coordinamento	20
1.3 L'IA e la trasformazione dell' <i>expertise</i>	23
1.3.1 Ridefinizione del concetto di <i>expertise</i>	24
1.3.2 Relazione tra AI e professionisti <i>knowledge-based</i>	26
CAPITOLO SECONDO: Il nuovo modello teorico	30
2.1 IA e ridefinizione del coordinamento	30
2.1.1 Confronto tra meccanismi tradizionali e <i>AI-driven</i>	31
2.2 IA e divisione del lavoro	34
2.2.1 Divisione algoritmica del lavoro e meta-competenze emergenti	35
2.2.2 <i>Human-in-the-loop</i> e nuove forme di <i>expertise</i>	40
2.3 IA e struttura organizzativa	43
2.3.1 Modelli organizzativi abilitati dall'IA	44
2.3.2 Equilibrio tra controllo algoritmico e autonomia umana	48

CAPITOLO TERZO: Analisi sperimentale	51
3.1 Obiettivi della ricerca empirica	51
3.2 Metodologia della ricerca	53
3.2.1 Il settore	53
3.2.2 Disegno della ricerca e campione	53
3.2.3 Il questionario	57
3.3 Risultati dell'analisi empirica	59
3.3.1 Analisi dei dati	59
3.3.2 Analisi descrittiva e correlazione	61
3.3.3 Regressione multipla con la frequenza di utilizzo dell'AI come variabile dipendente	64
3.3.4 Regressione multipla con <i>Job Satisfaction</i> come variabile dipendente	69
3.3.5 Analisi di mediazione	73
3.3.6 Analisi di moderazione	75
3.4 Limiti della ricerca quantitativa	79
Discussione dei risultati finali della ricerca	80
Conclusioni	86

Introduzione

Il presente studio verte su come la progettazione organizzativa viene influenzata e trasformata dall'ingresso di strumenti di intelligenza generativa. L'obiettivo principale è rispondere alla seguente domanda di ricerca: “*Come l'Intelligenza Artificiale impatta sui meccanismi di coordinamento, competenze richieste, divisione del lavoro e strutture organizzative?*”

Questo studio esplora l'Intelligenza Artificiale non solo come strumento di automazione, ma come un nuovo meccanismo di coordinamento nelle organizzazioni, attraverso la ridefinizione di processi e ruoli interni, supportando il paradigma della complementarità uomo- macchina. La ricerca analizza come l'IA influisca sulla gerarchia, sulle competenze e sulle dinamiche di controllo, generando nuove forme di collaborazione tra esseri umani e algoritmi. Infatti, tale cambiamento ridefinisce *l'expertise* umana, creando ruoli ibridi e nuovi modelli di *leadership* e controllo. Successivamente, lo studio esplora il modo in cui tali cambiamenti organizzativi influenzano gli atteggiamenti e il benessere lavorativo dei dipendenti.

La tesi adotta un approccio prettamente quantitativo, supportato anche da un'analisi qualitativa dei dati. È stato somministrato un questionario sottoposto a 105 dipendenti tra cui *manager* e funzionari di aziende *knowledge-based* del settore della consulenza, i quali adottano l'AI come strumento di lavoro quotidiano. È stato preso in considerazione questo settore perché ritenuto strategico in ambito di trasformazione digitale, essendo già plasmato dall'IA generativa.

I dati raccolti sono stati analizzati mediante analisi descrittive, analisi di correlazionali e modelli di regressione multipla per comprendere le relazioni più significative tra le variabili. Inoltre, sono state effettuate analisi di mediazione per spiegare tramite quali variabili si manifesta la soddisfazione lavorativa e l'impegno verso l'organizzazione. Infine, tramite le analisi di mediazione moderata sono state testate le relazioni significative per verificare se queste cambiassero in base a variabili sociodemografiche come il genere, la generazione di appartenenza dei rispondenti o la posizione lavorativa ricoperta da essi. Le risposte aperte fornite dai partecipanti offrono un quadro esaustivo riguardo le opportunità che l'intelligenza artificiale offre nel mondo lavorativo, ma anche le perplessità dovute al rischio di riduzione dell'autonomia individuale e alla possibilità di *bias* algoritmici.

Tale elaborato evidenzia come l'impatto dell'intelligenza artificiale sul benessere lavorativo dipenda dalle modalità di implementazione. La nuova tecnologia può potenziare coordinamento ed l'efficienza, ma i suoi benefici emergono nel momento in cui avviene un ripensamento organizzativo che mantenga centrale il valore umano e la sua autonomia. Dal punto di vista

pratico- manageriale, la ricerca fornisce suggerimenti per una corretta implementazione dell'IA in azienda. È opportuno accompagnare tale trasformazione tecnologica con interventi di *change management* che definiscano chiaramente nuovi ruoli e processi, formino il personale alle competenze digitali emergenti. Politiche di questo tipo, possono prevenire effetti indesiderati come la sensazione di una sorveglianza opprimente da parte dell'algoritmo o resistenze al cambiamento. È importante favorire un'integrazione tecnologica sostenibile, in cui l'intelligenza artificiale sia accettata come leva di miglioramento organizzativo, senza erodere la motivazione e la fiducia dei dipendenti.

In conclusione, il contributo di questa ricerca è duplice. Da un lato, offre una base teorica per comprendere l'Intelligenza Artificiale come leva di trasformazione organizzativa; dall'altro, fornisce linee guida pratiche per aziende che intendono implementare l'IA in contesti complessi, bilanciando efficienza e valore delle competenze umane.

CAPITOLO PRIMO: L'Intelligenza Artificiale nella Progettazione Organizzativa

1.1 Evoluzione dell'Intelligenza Artificiale nelle organizzazioni

Oggi ci troviamo di fronte ad un punto di svolta epocale, caratterizzato da un progresso senza precedenti nella tecnologia digitale poiché l'Intelligenza Artificiale risulta essere una delle innovazioni più rivoluzionarie nella storia della tecnologia applicata alle organizzazioni aziendali. Questa fase è stata definita la "seconda era delle macchine"¹ in quanto esse assumono progressivamente ruoli legati al lavoro cognitivo, da sempre considerati prerogativa umana, senza limitarsi più a eseguire compiti fisici e ripetitivi. Tale implementazione nelle organizzazioni determina trasformazioni profonde poiché la tecnologia diventa un elemento abilitante che influisce in modo sostanziale sulla progettazione organizzativa, ovvero sulla definizione delle strutture formali, dei processi operativi e dei ruoli funzionali. Nel corso degli ultimi decenni, le imprese hanno adottato approcci progressivamente più articolati per incorporare le tecnologie intelligenti. Da un iniziale orientamento all'automazione di compiti circoscritti, si sono evolute verso modelli cognitivi avanzati mirati all'amplificazione delle capacità decisionali e operative degli individui. Sono stati raggiunti importanti traguardi in questo ambito come i rapidi progressi nell'apprendimento automatico, grazie alla capacità di elaborazione di grandi quantità di dati, e significativi traguardi nella comprensione del linguaggio naturale e nella robotica avanzata. Nonostante ciò, l'IA presenta ancora notevoli limiti come svolgere compiti che richiedono creatività, intuizione e competenze socio-emotive. Questo scenario ha determinato una profonda trasformazione nella relazione tra uomo e macchina poiché l'interazione tra questi due mondi si configura come una cooperazione sinergica. Il paradigma emergente propone un modello di complementarità in cui le macchine agiscono come strumenti di potenziamento delle capacità umane piuttosto che fondarsi su una logica dicotomica basata sulla sostituzione. In questa prospettiva i due poli agiscono su attività diverse. L'IA può migliorare la qualità delle decisioni aziendali, supportare i *manager* nell'analisi dei dati e ottimizzare i processi organizzativi, mentre l'intervento umano continua ad essere imprescindibile per fornire interpretazioni contestuali, intuizioni strategiche e capacità critiche per guidare le strategie aziendali. Pertanto, l'adozione dell'Intelligenza Artificiale nelle organizzazioni non deve essere interpretata come una minaccia all'occupazione, bensì come un'opportunità per aumentare il potenziale umano e ridefinire il modo in cui il lavoro

¹ Brynjolfsson, E., & McAfee, A. (2014). *The Second Machine Age: Work, Progress, and Prosperity in a Time of Brilliant Technologies*. W. W. Norton & Company.

manageriale viene concepito e svolto. Le imprese stanno progressivamente adottando soluzioni intelligenti per automatizzare le attività di routine che richiederebbero tempo e sforzo da parte dell'essere umano e tale approccio consente di ridurre le attività svolte dalle risorse umane con basso valore aggiunto e al contempo permette di reimpiegare tali risorse in attività più complesse, che richiedono pensiero critico e creatività. In tal modo, l'Intelligenza Artificiale contribuisce al miglioramento produttivo dell'azienda, poiché porta con sé numerosi benefici in termini di qualità, produttività, costi e tempistiche. Infatti, tale implementazione consente una riduzione significativa degli errori umani in quanto, a differenza delle persone, l'IA non è incline a distrazioni o affaticamento nei compiti ripetitivi e questo si traduce in un miglioramento della *customer experience* garantendo maggiore affidabilità e precisione. Inoltre, migliora i tempi di gestione grazie alla sua operatività continua 24/7, minimizzando le inefficienze e accelerando i processi, contestualmente si ottiene il superamento delle risorse limitate e il rallentamento del flusso di lavoro generato da componenti esterne. Dal punto di vista economico, tale implementazione, si concretizza in una riduzione dei costi operativi e un ritorno sull'investimento in tempi brevi, il tutto permette di realizzare un modello facile e veloce grazie al riutilizzo di elementi, *dataset* e infrastrutture esistenti che rende lo sviluppo più efficiente e scalabile.

1.1.1 Automazione vs. Aumentazione

L'Intelligenza Artificiale nelle organizzazioni può essere applicata secondo due principali approcci: automazione e aumentazione. Queste due strategie, sebbene distinte, non sono necessariamente alternative, ma rappresentano modalità complementari di integrazione dell'IA nei processi aziendali. I sistemi intelligenti all'interno dell'azienda si sono evoluti attraverso un progresso graduale, iniziato da impieghi focalizzati sull'automazione di attività ripetitive e standardizzate fino ad arrivare alle attuali applicazioni orientate al potenziamento delle capacità decisionali e creative del personale. In un primo momento, l'introduzione di strumenti digitali e algoritmi nelle imprese rispondeva principalmente all'esigenza di incrementare l'efficienza attraverso la sostituzione del lavoro umano in compiti specifici e formalizzabili come nel caso di operazioni di calcolo, gestione massiva di dati o funzioni di controllo automatizzato della qualità. L'automazione implica che le macchine prendano in carico un compito che di solito viene svolto dall'uomo, riducendo o eliminando del tutto il suo intervento, grazie a sistemi che seguono regole e processi standardizzati, replicando attività ripetitive con elevata precisione. Tali sistemi, quindi, eseguono compiti predefiniti in modo coerente, migliorando così

l'efficienza operativa e minimizzando il rischio di errore.² Tuttavia, un orientamento eccessivo verso l'automazione può portare a una progressiva riduzione della forza lavoro in determinati settori, rendendo necessaria una riorganizzazione dei ruoli e delle responsabilità aziendali. Un esempio significativo è rappresentato dall'industria manifatturiera, dove l'automazione ha storicamente sostituito mansioni manuali con macchinari programmati per operare in autonomia, riducendo la necessità di supervisione diretta. Infatti, tali soluzioni, spesso costruite su basi algoritmiche deterministiche, si inseriscono all'interno di un contesto essenzialmente razionale del lavoro, simile a quello osservato nei processi di automazione industriale del secolo scorso. Con l'avvento di tecnologie più avanzate, come l'apprendimento automatico e le reti neurali artificiali, si rileva una trasformazione concettuale in cui l'obiettivo non è più la mera sostituzione dell'uomo, ma piuttosto l'integrazione tra le due realtà in un'ottica di potenziamento delle *performance* individuali e collettive. In questo paradigma, l'IA funge da strumento di aumentazione, supportando le attività decisionali, analizzando grandi volumi di dati e fornendo *insight* rilevanti, in modo che l'uomo possa utilizzarli per prendere decisioni più informate.³ Questo approccio richiede un'interazione costante tra le due realtà e presuppone investimenti nella formazione del personale per sviluppare competenze complementari all'IA, promuovendo una cultura organizzativa incentrata sulla collaborazione uomo-macchina.⁴ L'aumentazione, per esempio, viene applicata in ambiti come la gestione aziendale attraverso l'uso di algoritmi predittivi oppure tramite piattaforme di analisi finanziaria per migliorare il processo decisionale, mantenendo comunque da parte del *manager* il controllo strategico finale. Scegliere tra automazione e aumentazione non è un processo binario, bensì una decisione strategica dipendente dal contesto organizzativo e dagli obiettivi aziendali. L'efficacia di questa decisione è subordinata alla capacità dell'azienda di integrare in modo coerente tali tecnologie con la propria missione e i propri valori, garantendo un equilibrio tra efficienza operativa e valorizzazione del capitale umano. L'automazione migliora le prestazioni e riduce i costi, mentre l'aumentazione garantisce maggiore flessibilità, adattabilità e capacità di rispondere a situazioni complesse. Questa dicotomia, sebbene intuitiva, non è netta poiché le due applicazioni si influenzano reciprocamente nel tempo e nello spazio, generando una tensione

² Davenport, T. H., & Kirby, J. (2016). *Only Humans Need Apply: Winners and Losers in the Age of Smart Machines*. Harper Business.

³ Brynjolfsson, E., & McAfee, A. (2014). *The Second Machine Age: Work, Progress, and Prosperity in a Time of Brilliant Technologies*. W. W. Norton & Company.

⁴ Wilson, H. J., & Daugherty, P. R. (2018). *Human + Machine: Reimagining Work in the Age of AI*. Harvard Business Review Press.

che può essere meglio compresa attraverso la teoria del paradosso.⁵ Quest'ultima afferma che due elementi apparentemente contraddittori coesistono e si influenzano reciprocamente nel tempo e nello spazio, anziché escludersi a vicenda. In questo contesto, automazione e aumentazione non sono semplicemente due alternative tra cui scegliere, ma componenti di una relazione dinamica e interdipendente. Su scala temporale, l'aumentazione può evolversi verso l'automazione poiché un processo inizialmente supportato dall'interazione tra uomo e tecnologia può, nel tempo, diventare completamente automatizzato, grazie alla crescente capacità dei sistemi di apprendere, adattarsi e prendere decisioni in autonomia. Nel settore della gestione delle risorse umane, i sistemi di selezione del personale basati su IA erano originariamente utilizzati per fornire raccomandazioni ai *recruiter* in base all'analisi dei dati storici, ma con l'integrazione delle tecniche di *machine learning* e l'affinarsi della precisione degli algoritmi, queste applicazioni si sono evolute fino ad automatizzare l'intero processo di *screening* e valutazione, riducendo progressivamente il bisogno dell'intervento umano.

Allo stesso modo, su scala spaziale, un'innovazione inizialmente confinata a un singolo processo può diffondersi ad altri ambiti aziendali, generando effetti a catena che trasformano l'intera struttura organizzativa. Tale cambiamento si può osservare nell'ambito della *customer service*, nel quale l'IA è stata impiegata inizialmente per gestire richieste di base ed è stata progressivamente estesa alla gestione delle vendite, alla personalizzazione dell'esperienza utente fino ad arrivare alla formulazione di strategie di marketing *data-driven*. L'enfasi eccessiva su uno dei due poli può infatti creare effetti indesiderati: un'automazione spinta all'estremo può portare alla perdita di competenze umane e alla rigidità dei processi decisionali, mentre un *focus* esclusivo sull'aumentazione può limitare il potenziale di scalabilità e ridurre i vantaggi competitivi delle organizzazioni. Solo integrando queste due realtà, l'Intelligenza Artificiale può donare un valore aggiunto ad ogni realtà imprenditoriale.

L'applicazione della teoria del paradosso nel contesto della gestione dell'Intelligenza Artificiale aiuta a comprendere come queste due forze possano operare simultaneamente, creando cicli di rinforzo che influenzano il funzionamento aziendale. Le organizzazioni che privilegiano l'automazione spesso sperimentano una riduzione del coinvolgimento umano nei processi decisionali e di conseguenza portano ad una maggiore dipendenza dagli algoritmi e, potenzialmente, a una minore capacità di adattamento a situazioni impreviste. Allo stesso modo, le aziende che puntano solo sull'aumentazione possono incorrere in una complessità gestionale crescente poiché il mantenimento dell'interazione uomo-macchina richiede risorse, formazione

⁵ Schad, J., Lewis, M. W., Raisch, S., & Smith, W. K. (2016). *Paradox Research in Management Science: Looking Back to Move Forward*. *Academy of Management Annals*, 10(1), 5–64.

continua e strategie di coordinamento più sofisticate. Approcci rigidi o unilaterali possono portare a circoli viziosi, mentre una gestione consapevole di questa dualità può favorire dinamiche virtuose di innovazione e adattamento. La teoria del paradosso offre quindi una chiave interpretativa preziosa per navigare le complessità della trasformazione digitale, consentendo alle aziende di trovare un equilibrio tra efficienza tecnologica e valore delle competenze umane.

1.1.2 Impatto sui modelli organizzativi

L'evoluzione dell'Intelligenza Artificiale da strumento di mera automazione a leva per l'aumentazione delle capacità decisionali e operative comporta profonde trasformazioni nei modelli organizzativi, ovvero nelle modalità con cui le imprese strutturano le proprie attività, distribuiscono le responsabilità e coordinano le risorse. Uno dei temi centrali emersi nel dibattito accademico riguarda il potenziale impatto dell'IA sulla configurazione gerarchica delle organizzazioni, in particolare se essa conduce a strutture più piatte o favorisce nuove forme di accentramento.

Una prima prospettiva sostiene che l'adozione delle nuove tecnologie può incentivare la decentralizzazione decisionale, promuovendo strutture organizzative più orizzontali. Il motivo sottostante a favore di tale tesi è che la disponibilità di strumenti intelligenti, in grado di elaborare dati in tempo reale e fornire supporto analitico, consente ai lavoratori di prendere decisioni operative senza ricorrere costantemente ai livelli gerarchici superiori. Questo tipo di potenziamento tecnologico può rafforzare l'autonomia dei dipendenti, soprattutto in contesti dove la tempestività decisionale è cruciale. Un esempio emblematico si riscontra nei *team* tecnici o di manutenzione, dove le tecnologie possono identificare anomalie e suggerire interventi in modo proattivo, permettendo agli operatori di agire velocemente e autonomamente, senza dover attendere istruzioni dal *management*. Generalizzando questa dinamica, si configura una struttura meno piramidale, con una linea di comando più corta e una potenziale riduzione dei livelli intermedi di controllo.

Un'evidenza pratica di questo fenomeno si riscontra nei modelli organizzativi adottati da alcune aziende digitali e piattaforme, come *Uber* o *Airbnb*, dove il coordinamento di milioni di interazioni avviene grazie a infrastrutture algoritmiche, con un minimo intervento umano diretto. In tali casi, le tecnologie svolgono funzioni tipiche della supervisione umana come l'allocazione dei compiti, il monitoraggio delle *performance* e la somministrazione di *feedback*, permettendo a pochi dirigenti di controllare una vasta rete di collaboratori. Questo approccio aumenta l'ampiezza del controllo manageriale e può rendere obsoleti alcuni livelli della

gerarchia tradizionale. Al contempo, una visione alternativa suggerisce che l'IA possa invece promuovere nuove forme di centralizzazione, in particolare intorno ai detentori di competenze tecniche e ai proprietari dei dati. L'implementazione efficace di sistemi di IA richiede infatti infrastrutture digitali complesse, investimenti consistenti e risorse umane altamente qualificate, come *data scientist* e ingegneri del *machine learning*. Di conseguenza, molte organizzazioni scelgono di creare centri di competenza o unità specialistiche, accentrando il controllo strategico e tecnologico. Inoltre, poiché gli algoritmi operano sulla base di parametri predefiniti, spesso stabiliti dal *top management*, l'IA può diventare uno strumento per rafforzare l'allineamento tra esecuzione operativa e obiettivi aziendali definiti dall'alto.

Numerose analisi in ambito organizzativo hanno rilevato che la concentrazione delle competenze digitali in specifiche unità aziendali può generare *silos* tecnici, nei quali gli esperti sviluppano soluzioni tecnologiche in assenza di un dialogo efficace con le aree di *business*, con il rischio di produrre applicazioni tecnologiche non allineate alle reali esigenze operative.⁶

Tale scenario può favorire una fase iniziale di accentramento delle decisioni in ambito tecnologico, che necessita successivamente di essere riequilibrata attraverso strategie mirate alla diffusione delle competenze e all'accesso condiviso dei dati, come programmi formativi trasversali o strumenti di analisi *self-service*, al fine di promuovere una maggiore integrazione tra funzione tecnica e ambiti gestionali.

Oltre al tema della distribuzione del potere decisionale, l'impatto dell'IA si riflette anche sulla ridefinizione dei ruoli e delle competenze all'interno delle organizzazioni.

Il successo di un'organizzazione dipende dalla capacità di adattarsi e investire nelle competenze digitali, infatti, attività di *upskilling* e il *reskilling* del personale sono cruciali per colmare i divari tecnologici. Per cui l'introduzione di processi aziendali aumentati dalla tecnologia richiede una revisione profonda dell'organigramma: da un lato, emergono nuove figure professionali legate alla gestione e all'interpretazione dei sistemi intelligenti; dall'altro, si rende necessario un aggiornamento continuo delle competenze esistenti per garantire una collaborazione efficace tra esseri umani e macchine.⁷ Tra i profili professionali emergenti si possono individuare diverse figure chiave, tra queste, vi sono i professionisti incaricati della progettazione e manutenzione degli algoritmi, come i *data scientist* e i *data engineer*; i *trainer* dell'IA, responsabili della preparazione dei *set* di dati per l'addestramento dei modelli e del

⁶ Puranam, P., Alexy, O., & Reitzig, M. (2019). *What's "new" about new forms of organizing?* Academy of Management Review, 49(2), 1–16

⁷ Wilson, H. J., & Daugherty, P. R. (2018). *Human + Machine: Reimagining Work in the Age of AI*. Harvard Business Review Press.

loro continuo perfezionamento, svolgendo attività come la personalizzazione delle risposte di un *chatbot*; gli *explainer* ovvero mediatori algoritmici, figure che traducono gli algoritmi dei modelli in informazioni fruibili per chi assume decisioni; infine, gli esperti che supervisionano l'etica dell'Intelligenza Artificiale con il compito di garantire la conformità ai principi di equità, trasparenza e tutela dei dati. Tali professioni non si inseriscono agevolmente negli assetti organizzativi tradizionali, perciò, richiedono una revisione dei ruoli e dei percorsi di carriera. In questo contesto, quindi, le aziende sono chiamate a ristrutturare il proprio modello organizzativo, riducendo il numero di mansioni classiche e riorganizzando il personale in grandi aree di competenza trasversali, favorendo così una maggiore adattabilità interna.

Un ulteriore effetto significativo dell'introduzione dell'IA riguarda la trasformazione dei processi decisionali e dei meccanismi di coordinamento, tema che sarà oggetto di ulteriore approfondimento nella sezione successiva. Nei contesti in cui le tecnologie intelligenti sono integrate nei flussi operativi, le decisioni tendono a fondarsi sempre più su modelli analitici e algoritmici, i quali conducono a decisioni *data-driven*. Sebbene questo possa accrescere la coerenza e la razionalità delle scelte, riducendo l'influenza di pregiudizi soggettivi, si presentano anche nuove criticità, come la ridotta trasparenza dei criteri decisionali e una certa rigidità operativa. In alcuni modelli organizzativi, le decisioni quotidiane vengono assunte direttamente da sistemi automatizzati – ad esempio nella gestione predittiva degli *stock* – mutando il ruolo dei *manager*, che da esecutori diventano architetti dei processi decisionali, definendo i parametri e le logiche con cui gli algoritmi operano. Questa trasformazione richiede nuove competenze gestionali, tra cui la capacità di comprendere il funzionamento dei modelli predittivi e intervenire sui loro criteri di funzionamento. Per affrontare tali sfide, molte imprese stanno istituendo *task force*, composti da rappresentanti delle funzioni tecnologiche, manageriali, legali e di *compliance* etica e tali comitati interdisciplinari rappresentano un ulteriore elemento strutturale che contribuisce alla riconfigurazione degli assetti organizzativi. È fondamentale evidenziare che l'effetto dell'Intelligenza Artificiale sulle strutture organizzative non segue un percorso univoco né predeterminato poiché le trasformazioni indotte da tali tecnologie variano significativamente in funzione del settore di appartenenza, delle scelte strategiche adottate dalle imprese e del livello di maturità digitale raggiunto. In tal senso, Bailey e Barley, nel loro articolo "*Beyond design and use: How scholars should study intelligent technologies*" pubblicato nel 2020 sulla rivista *Information and Organization*, invitano ad adottare una prospettiva critica nei confronti delle narrazioni deterministiche che presentano l'IA come una forza inevitabile capace di generare trasformazioni organizzative senza precedenti, prefigurando scenari estremi, quali l'eliminazione massiva del lavoro umano

o l'emergere di imprese completamente autonome. Gli autori osservano come tali visioni tendano a riaffiorare ciclicamente in corrispondenza di importanti discontinuità tecnologiche; tuttavia, l'evidenza empirica mostra che i cambiamenti sono spesso più gradualisti e influenzati da dinamiche sociotecniche complesse, in cui fattori organizzativi, culturali e istituzionali mediano l'effettiva adozione e l'impatto delle innovazioni tecnologiche.

In numerosi contesti organizzativi, l'adozione dell'Intelligenza Artificiale tende a rafforzare traiettorie evolutive già in corso, infatti, agisce da acceleratore di modelli organizzativi agili e flessibili che hanno un orientamento verso processi orizzontali e lavorano in piccoli gruppi di lavoro multidisciplinari. Ciò accade perché quest'ultimi riescono a testare rapidamente applicazioni di IA su segmenti circoscritti dell'attività aziendale, e qualora i risultati si rivelino efficaci, l'organizzazione può scalare tali soluzioni a livello sistemico. Questo approccio, sperimentale e iterativo, risulta coerente con configurazioni meno gerarchiche e orientate alla progettualità, un assetto già adottato da molte imprese, ma che diventa ancora più strategico in un contesto segnato dalla rapida evoluzione delle tecnologie intelligenti.

In sintesi, si può affermare che l'Intelligenza Artificiale rappresenta un catalizzatore nella trasformazione dei modelli organizzativi, senza tuttavia configurarsi come un fattore univoco o rigidamente prescrittivo. Al contrario, l'IA stimola una revisione critica dei presupposti organizzativi tradizionali, inducendo le imprese a introdurre forme di maggiore adattabilità. Questa esigenza di flessibilità si manifesta su più piani: nella formazione continua dei lavoratori- attraverso percorsi di aggiornamento e sviluppo di nuove competenze digitali-, nella definizione di ruoli professionali ibridi o inediti, nella gestione dei processi decisionali- sempre più orientati ai dati e, in alcuni casi, automatizzati - e nella stessa struttura organizzativa, che deve essere in grado di accogliere algoritmi non solo come strumenti, ma come veri e propri attori organizzativi, capaci di interagire con il lavoro umano.

Le aziende che riusciranno a ristrutturare il proprio assetto organizzativo per integrare in modo armonico le potenzialità dell'IA, saranno in grado di accelerare i tempi di risposta al cambiamento senza sacrificare il valore aggiunto delle capacità umane, ponendosi in una posizione privilegiata per acquisire vantaggi competitivi sostenibili.

1.2 I meccanismi di coordinamento all'interno dell'organizzazione e l'AI

L'organizzazione aziendale viene definita da Treccani come il “processo di predisposizione di risorse (umane, fisiche, informative e altre ancora) in una conformazione strutturata, al fine di portare avanti piani e realizzare gli obiettivi dell'impresa” – e continua – “l'organizzazione aziendale è il processo di selezione e strutturazione dei mezzi con i quali tali obiettivi vengono realizzati”.⁸

Possiamo, quindi, definire l'organizzazione come un sistema complesso, costituito da persone e da risorse materiali e immateriali, intenzionalmente connesse e coordinate per perseguire obiettivi condivisi. Essa presenta confini relativamente definiti e opera in modo continuativo, interagendo con l'ambiente esterno in un processo di adattamento e trasformazione costante. Le organizzazioni sono entità sociali guidate da obiettivi, progettate come sistemi di attività strutturate e coordinate in modo deliberato. Il loro funzionamento implica la gestione efficiente delle risorse per la produzione di beni e servizi, la promozione dell'innovazione anche attraverso l'uso di tecnologie produttive *computer-based*. Inoltre, esse non si limitano a rispondere ai cambiamenti dell'ambiente, ma contribuiscono attivamente a plasmarlo, influenzando dinamiche economiche e sociali. La capacità di generare valore per gli azionisti, clienti, dipendenti e la società nel suo complesso, dipende dall'efficacia con cui un'organizzazione struttura i propri processi e adotta modelli innovativi. Nel contesto attuale, le organizzazioni affrontano sfide complesse e l'analisi di quest'ultima è quindi fondamentale per comprendere come essa possa evolvere attraverso nuove metodologie gestionali e tecnologiche. Per un'azienda, la specializzazione rappresenta un elemento essenziale per il raggiungimento di elevati livelli di efficienza e competitività, ma per essere efficace richiede un adeguato coordinamento. In qualsiasi attività organizzativa, due aspetti fondamentali ne determinano il funzionamento: la suddivisione del lavoro in compiti specifici e l'integrazione coordinata di tali attività, al fine di garantire coerenza operativa e il conseguimento degli obiettivi prefissati.

Il processo organizzativo si articola in fasi ben definite che consentono di strutturare le attività e le risorse in modo efficiente. La prima fase, come già detto, riguarda la definizione degli obiettivi strategici, ovvero la direzione verso cui l'organizzazione intende muoversi e le strategie da adottare per il loro raggiungimento. A questa segue la fase di divisione del lavoro, che porta alla specializzazione verticale e orizzontale. La prima si riferisce alla stratificazione

⁸ <https://www.treccani.it/enciclopedia/organizzazione-aziendale/>

dei ruoli all'interno di un'organizzazione e implica la suddivisione delle responsabilità in livelli gerarchici per cui le decisioni vengono prese ai livelli superiori della gerarchia e trasmesse verso il basso determinando un flusso verticale di comunicazione e controllo. Ciò consente di assegnare compiti specifici a persone con competenze specialistiche e di garantire una chiara distribuzione dell'autorità, sebbene possa rallentare i processi decisionali e ridurre l'autonomia operativa dei livelli inferiori. La specializzazione orizzontale, invece, riguarda la suddivisione del lavoro in compiti distinti, assegnati a individui o unità con competenze specifiche, senza necessariamente coinvolgere una gerarchia verticale. Questo tipo di specializzazione consente di aumentare l'efficienza e la produttività, ma, se non adeguatamente coordinata, può generare problemi di frammentazione e mancanza di integrazione tra le diverse funzioni aziendali.

Di pari passo con la crescita dell'organizzazione, aumenta la necessità di una suddivisione sistemica del lavoro, rendendo indispensabile l'adozione di strumenti di coordinamento volti a garantire il corretto funzionamento dell'organizzazione e la sinergia tra le attività dei singoli individui. Distinguiamo due tipologie di coordinamento: verticale e orizzontale.

Il coordinamento verticale si fonda sulla gerarchia e sull'attribuzione di autorità decisionale, consentendo una rapida esecuzione delle attività, sebbene possa limitare il coinvolgimento e la creatività dei lavoratori. In situazioni che richiedono tempestività e reattività, il meccanismo gerarchico rappresenta una soluzione efficace. Tuttavia, quando è necessario favorire l'innovazione e la collaborazione, il coordinamento laterale assume un ruolo cruciale.

Il coordinamento laterale si realizza attraverso diversi strumenti, tra cui il confronto diretto tra le parti, l'uso di sistemi informativi per la condivisione delle informazioni e la standardizzazione di procedure e processi. Ulteriori strumenti di coordinamento laterale includono i collegamenti formali, come il lavoro in *team*, e i collegamenti informali, quali le *task force* temporanee e i ruoli integratori, figure professionali incaricate di facilitare la comunicazione tra diverse unità organizzative. Un ulteriore collegamento informale riguarda lo sviluppo di comunità di pratica, ovvero gruppi di individui che collaborano spontaneamente per risolvere problemi organizzativi, che l'azienda non ha ancora identificato, basandosi su una cultura di fiducia reciproca e apprendimento condiviso. Questi meccanismi risultano particolarmente efficaci per affrontare le sfide contemporanee del *management*, coniugando efficienza operativa e adattabilità a un ambiente in continua evoluzione.

Attualmente i meccanismi di coordinamento tradizionali all'interno delle organizzazioni stanno subendo dei profondi cambiamenti a causa dell'avvento dell'Intelligenza Artificiale, la quale sta ridefinendo il modo in cui le attività vengono gestite e integrate nei diversi contesti aziendali. L'avvento delle nuove tecnologie ha introdotto un paradigma alternativo, in cui il

coordinamento non si limita più a processi gestiti direttamente dagli individui, ma è progressivamente affidato ad algoritmi capaci di analizzare grandi quantità di dati, prendere decisioni e ottimizzare l'allocazione delle risorse in tempo reale. Le organizzazioni *knowledge-based*, caratterizzate da una crescente complessità operativa e dalla necessità di coordinare *team* distribuiti geograficamente, stanno adottando sempre più frequentemente strumenti di IA per supportare le attività di gestione e supervisione.⁹ Questa trasformazione impatta direttamente sui meccanismi di coordinamento tradizionali, ridefinendo il ruolo del *management* e la distribuzione delle responsabilità decisionali, favorendo l'emergere di strutture organizzative più flessibili e decentralizzate.

1.2.1 Teoria di Mintzberg (1979) e nuovi modelli di coordinamento

La Teoria di Mintzberg rappresenta uno dei concetti fondamentali della progettazione organizzativa, sviluppato da Henry Mintzberg nel suo libro *The Structuring of Organizations* (1979). L'autore identifica cinque meccanismi fondamentali di coordinamento che spiegano come le organizzazioni gestiscono le attività e le interazioni tra i membri:

1. **Mutuo adattamento**, in cui il coordinamento avviene attraverso la comunicazione informale tra gli individui. Le persone collaborano a stretto contatto, si scambiano conoscenze e negoziano in tempo reale sulle azioni da intraprendere. Tale meccanismo è tipico dei contesti organizzativi di piccole dimensioni come le *startup*, *team* di ricerca interdisciplinari e organizzazioni innovative con elevata autonomia. In questi contesti la flessibilità comunicativa diventa essenziale per affrontare l'incertezza e favorire la collaborazione.
2. **Supervisione diretta**, in cui il coordinamento è ottenuto attraverso il *manager*, il quale è responsabile della supervisione e del controllo delle attività dei subordinati. Egli decide i ruoli, controlla i progressi e interviene in caso di inadempienze. Tale approccio è comune nelle strutture gerarchiche tradizionali e nelle organizzazioni con operazioni standardizzate. L'imprenditore coordina personalmente i collaboratori per cui si ha un controllo centralizzato, tuttavia, l'efficacia di questo modello diminuisce all'aumentare del numero di subordinati o della specializzazione delle attività, rendendolo meno sostenibile in contesti complessi.

⁹ Van den Broek, J., Sergeeva, A., & Huysman, M. (2021). *When the machine meets the expert: An ethnography of developing AI for hiring*. *Organization Science*, 32(6), 1415–1437 pg.

3. **Standardizzazione dei processi di lavoro**, in cui le attività sono coordinate a priori attraverso procedure, regole e metodologie definite e per tale motivo, non servono comunicazioni costanti perché ogni individuo conosce il proprio compito. Questo meccanismo è tipico delle cosiddette “burocrazie meccaniche”, ovvero contesti con elevata ripetitività dove per garantire *output* consistenti si standardizzano i metodi di lavoro.
4. **Standardizzazione delle competenze**, il cui coordinamento si basa sulla condivisione di conoscenze e competenze acquisite attraverso percorsi formativi comuni. Gli individui, avendo ricevuto la stessa formazione e interiorizzato standard professionali simili, operano in modo coerente anche in assenza di istruzioni dirette. È tipico delle burocrazie professionali, come ospedali, università o studi legali, dove medici, docenti o consulenti agiscono con ampia autonomia, ma secondo criteri condivisi. Il principale vantaggio risiede nell’equilibrio tra libertà operativa e coerenza organizzativa. La principale sfida è garantire l’aggiornamento continuo delle competenze e la convergenza interpretativa degli standard nel tempo.
5. **Standardizzazione degli *output***, il cui coordinamento è basato sul raggiungimento di obiettivi specifici, definiti in anticipo, lasciando autonomia operativa ai lavoratori. Vengono definiti gli obiettivi e i risultati attesi e tutti perseguono tali *target* quantitativi o qualitativi, affinché il lavoro sia indirizzato verso tali risultati. È tipico nelle organizzazioni divisionali, come le aziende con sistemi di *performance management*, nelle quali i manager godono di autonomia, ma sono valutati su parametri standardizzati. Tale approccio favorisce la flessibilità nei metodi e incoraggia soluzioni creative a livello locale; tuttavia, richiede sistemi efficaci di valutazione dei risultati. Un limite potenziale è rappresentato dalle singole unità, le quali possono privilegiare il raggiungimento dei propri obiettivi a discapito della cooperazione inter-funzionale.

I cinque meccanismi di coordinamento proposti da Mintzberg offrono una base teorica solida per comprendere come le organizzazioni tradizionalmente orchestravano il lavoro. Ogni configurazione organizzativa privilegia un meccanismo specifico, ma nella pratica più meccanismi coesistono. In tal caso si rende inevitabile il ricorso a forme di adattamento reciproco e leadership informale al fine di gestire eccezioni o per stimolare l’innovazione.¹⁰

¹⁰ Mintzberg, H. (1979). *The Structuring of Organizations: A Synthesis of the Research*. Englewood Cliffs, NJ: Prentice-Hall.

L'intelligenza Artificiale, all'interno di questo contesto, può definirsi come strumento che rafforza e modifica i meccanismi attuali, oppure come nuovo meccanismo di coordinamento esistente. In questo paragrafo analizzeremo l'AI sotto la prima prospettiva.

Rispetto al mutuo adattamento, si stanno diffondendo nuove forme di comunicazione e collaborazione anche a distanza, grazie a piattaforme digitali come *Microsoft Teams* o ambienti arricchiti da assistenti virtuali e *chatbot*. Tali strumenti permettono a gruppi anche geograficamente distanti di comunicare in tempo reale per condividere conoscenze e coordinarsi sulle attività. Inoltre, tali supporti digitali permettono anche di individuare esperti con l'ausilio di sistemi di *knowledge management*, indirizzando automaticamente una domanda all'esperto incaricato in azienda.¹¹

In merito alla supervisione diretta, emergono forme di "supervisione algoritmica" in cui il controllo non è esercitato da un *manager*, ma da sistemi intelligenti in grado di monitorare e guidare le attività in tempo reale, segnalando deviazioni o impartendo istruzioni. Ad esempio, nei *call center*, l'IA può suggerire risposte in tempo reale all'operatore durante le conversazioni o segnalare anomalie al responsabile; nei contesti manifatturieri, la *computer vision* permette di automatizzare il monitoraggio, in quanto supervisiona la linea produttiva, individuando difetti o rallentamenti da comunicare ai tecnici.¹² Tali compiti erano tradizionalmente svolti da figure manageriali, ma attualmente con l'introduzione di tali tecnologie la capacità di monitoraggio è aumentata sia a livello temporale – grazie ad un controllo continuo e ininterrotto – sia a livello quantitativo, permettendo l'analisi simultanea di grandi volumi di dati e l'individuazione tempestiva di criticità.

Per quanto riguarda la standardizzazione dei processi, essa evolve in senso dinamico poiché l'IA consente l'adozione di procedure adattive che si aggiornano in funzione del contesto, scegliendo la miglior procedura in base alle circostanze. Un sistema intelligente può determinare quale sequenza operativa adottare in base ai dati in tempo reale, garantendo coerenza ma anche flessibilità. Ne sono esempio i sistemi di smistamento logistico in cui gli algoritmi decidono in tempo reale quale pacco assegnare a ciascun nastro trasportatore in base al volume, alle priorità di consegna, ecc. Quindi, il personale e le macchine operano secondo

¹¹ Faraj, S., Jarvenpaa, S. L., & Majchrzak, A. (2011). *Knowledge collaboration in online communities*. *Organization Science*, 22(5), 1224–1239g.

¹² Kellogg, K. C., Valentine, M. A., & Christin, A. (2020). *Algorithms at work: The new contested terrain of control*. *Academy of Management Annals*, 14(1), 366–410g.

linee guida *standard* dettate dall'algoritmo in tempo reale, mantenendo il coordinamento senza la necessità che tutti conoscano e applichino una procedura invariabile.¹³

Relativamente alla standardizzazione delle competenze, l'IA contribuisce sia come supporto alla formazione - tramite piattaforme *adaptive learning* per assicurare *skills standard* a tutti i dipendenti - sia come strumento che fornisce linee guida operative uniformi, riducendo la mancanza di competenze tra i professionisti. Molte organizzazioni stanno definendo *standard* di alfabetizzazione digitale trasversale, includendo competenze di base delle nuove tecnologie nei percorsi formativi interni, coordinando il lavoro attraverso uno strumento comune di conoscenza.

Infine, la standardizzazione degli *output* beneficia dell'IA attraverso strumenti di monitoraggio continuo e *data analytics* che permettono ai *manager* di accedere in tempo reale agli indicatori di *performance* e disaggregare obiettivi generali in *target* specifici per singoli *team* o individui. L'IA, quindi, funge da tutor organizzativo, suggerendo aree di miglioramento e aumentando la trasparenza sulle performance. Infatti, un'azienda che vuole raggiungere un determinato livello di soddisfazione per il cliente potrebbe implementare nuovi sistemi di intelligenza artificiale per analizzare le interazioni di servizio e fornire aree di miglioramento su cui lavorare per raggiungere l'obiettivo globale.¹⁴

Oltre al rafforzamento dei meccanismi classici, oggi si assiste all'emergere di nuove forme organizzative rese possibili dal contesto digitale, le quali sfuggono ai cinque meccanismi canonici individuati da Mintzberg. Come già citato in precedenza, un esempio emblematico è il coordinamento algoritmico tipico delle piattaforme digitali come *Uber* e *Deliveroo*, dove la gestione delle attività – dall'assegnazione dei compiti al monitoraggio – è affidato a sistemi algoritmici che svolgono di fatto il ruolo di *manager* collettivi per una rete di lavoratori autonomi.¹⁵ Il confine organizzativo di tali modelli è sfumato poiché il rapporto tra lavoratori e piattaforma è spesso privo di vincoli gerarchici tradizionali e il coordinamento si basa su incentivi di mercato come *rating* e tariffe e vincoli tecnologici, dando origine a nuove forme di controllo e organizzazione del lavoro. In parallelo, si stanno affermando modelli organizzativi fondati su reti collaborative digitali, come i progetti *open source* o le comunità professionali online. In questi ambienti, la cooperazione tra migliaia di individui è facilitata da piattaforme

¹³ Brynjolfsson, E., & McAfee, A. (2014). *The Second Machine Age: Work, Progress, and Prosperity in a Time of Brilliant Technologies*. New York: W. W. Norton & Company.

¹⁴ Davenport, T. H., & Ronanki, R. (2018). *Artificial intelligence for the real world*. Harvard Business Review, 96(1), 108–116.

¹⁵ Lee, M. K., Kusbit, D., Metsky, E., & Dabbish, L. (2015). *Working with Machines: The Impact of Algorithmic and Data-Driven Management on Human Workers*. Proceedings of the 33rd Annual ACM Conference on Human Factors in Computing Systems (CHI '15), 1603–1612.

digitali che abilitano un coordinamento leggero ma efficace e l'IA permette di filtrare contributi, suggerire compiti o segnalando errori, contribuendo così all'emergere di una forma di intelligenza collettiva. Sebbene queste strutture ricordino il mutuo adattamento, la loro estensione e flessibilità operativa sono rese possibili solo grazie alla dimensione tecnologica sottostante.

In conclusione, la teoria di Mintzberg conserva una forte capacità esplicativa anche nell'era dell'IA, ma i contesti digitali suggeriscono la necessità di integrare il modello con nuove categorie.

1.2.2 L'IA come nuovo meccanismo di coordinamento

Vista l'affermazione dell'Intelligenza Artificiale come un nuovo strumento di coordinamento organizzativo, le aziende devono affrontare questa trasformazione con un approccio strategico attraverso il bilanciamento tra l'efficienza operativa e la necessità di mantenere un coinvolgimento attivo delle risorse umane. Alcuni studi affermano che l'Intelligenza Artificiale stia emergendo come sesto meccanismo di coordinamento, ridefinendo le dinamiche organizzative poiché non si limita a supportare le decisioni umane, ma può operare in modo autonomo. Ciò avviene, per esempio, in ambito di allocazione delle risorse, di gestione delle attività, di monitoraggio delle performance e anche nel processo decisionale, riducendo così la necessità di un intervento manuale e garantendo una maggiore efficienza operativa.¹⁶ In tale modo, l'IA consente di superare i limiti dei modelli di coordinamento tradizionali. In particolare, le organizzazioni *knowledge-based* stanno sperimentando un passaggio da sistemi di coordinamento centralizzati e gerarchici a modelli più flessibili e decentralizzati, in cui l'IA agisce come un facilitatore dell'integrazione tra team, processi e dati.¹⁷

L'impatto dell'IA sul coordinamento si manifesta in diversi ambiti:

- **Allocazione delle risorse:** algoritmi di *machine learning* possono ottimizzare la distribuzione delle competenze e conoscenze all'interno dell'organizzazione. Ciò comporta un'efficace allocazione delle risorse in modo da stabilire i compiti con maggiore pertinenza, riducendo inefficienze e colli di bottiglia nei flussi di lavoro.

¹⁶ Raisch, S., & Krakowski, S. (2021). *Artificial intelligence and management: The automation–augmentation paradox in decision making*. Academy of Management Review, 46(1), 192–210.

¹⁷ Newell, S. (2015). *Managing knowledge and managing knowledge work: What we know and what the future holds*. Journal of Information Technology, 30(1), 1–17.

- **Monitoraggio e supervisione:** sistemi basati su IA possono analizzare grandi quantità di dati operativi e fornire suggerimenti al fine di migliorare le performance aziendali, permettendo di tracciare il flusso delle informazioni e riducendo la necessità di un controllo diretto da parte dei manager. L'autore pone attenzione in merito alla trasformazione di tale monitoraggio ad un mero esercizio di controllo, ricordando che il valore del lavoro dipende dalla spontaneità delle relazioni umane, dall'esperienza lavorativa e dalle competenze del *team* di lavoro.
- **Ottimizzazione della comunicazione interna:** *chatbot* e assistenti virtuali facilitano la condivisione della conoscenza e la circolazione delle informazioni tra dipendenti, supportando la collaborazione tra *team* e riducendo la dipendenza da interazioni gerarchiche. Tuttavia, Newell sottolinea l'importanza di interazioni sociali tra il gruppo di lavoro per condividere le conoscenze specifiche.

L'IA come meccanismo di coordinamento offre diversi vantaggi potenziali, primo fra tutti è la scalabilità poiché può coordinare attività gestendo migliaia di transazioni al secondo e milioni di utenti, impensabile per una struttura gerarchica tradizionale. Il secondo vantaggio è la coerenza dovuta all'utilizzo di regole che vengono applicate in modo uniforme senza le variazioni dovute a discrezionalità o errori umani momentanei; ciò può tradursi da un alto in un'equità e dall'altro in rigidità. Terzo punto è la reattività dovuta alla capacità di cogliere cambiamenti ambientali o operativi in tempo reale, riorganizzando il flusso di lavoro. Questo rende l'organizzazione più agile nell'affrontare i cambiamenti imprevedibili. Quarto vantaggio è la capacità analitica avanzata, superiore a quella umana, in grado di processare e coordinare numerosi dati e variabili come le performance storiche, le preferenze individuali, le previsioni di domanda; ciò implica decisioni ottimali in termini di massimizzazione della produttività, riducendo i tempi.

Questi cambiamenti sollevano nuove sfide per la progettazione organizzativa legate a questioni di trasparenza e fiducia. Quando le decisioni su turni, assegnazioni o valutazioni vengono prese da un algoritmo, i lavoratori possono percepire queste scelte meno eque e affidabili. Inoltre, se le motivazioni non sono chiare, nella forza lavoro insorgono sentimenti di alienazione o incomprensione. Alcuni studi hanno dimostrato che un monitoraggio algoritmico troppo stringente può far sentire i dipendenti sorvegliati in ogni istante, provocando effetti negativi sul morale e sul benessere degli individui provocando stress e inadeguatezza. Come riportato nell'articolo "*Algorithmic Management: The Role of AI in Managing Workforces*" di

Mohammad Hossein Jarrahi, Mareike Möhlmann, Min Kyung Lee pubblicato nel 2023, alcuni magazzinieri di Amazon hanno riportato elevati livelli di pressione dovuti a un ritmo elevato di lavoro con conseguenti problemi fisici e psicologici. In tali casi, il coordinamento algoritmico può degenerare in una forma di “*taylorismo* digitale estremo” dove l’ottimizzazione dell’efficienza avviene a discapito della sostenibilità del lavoro e della dignità personale.

Un’ulteriore criticità è rappresentata dai *bias* e dalle distorsioni che possono emergere nei sistemi di IA utilizzati per coordinare il lavoro. Potrebbe accadere, infatti, che il sistema penalizzi alcune categorie di lavoratori nell’assegnazione dei turni, o che possa attribuire valutazioni di *performance* distorte in base a variabili irrilevanti perché l’algoritmo è stato addestrato su dati parziali o riflette pregiudizi impliciti. È difficile individuare e correggere queste ingiustizie a causa dell’opacità algoritmica che, senza un attento controllo umano e criteri rigorosi di progettazione etica, può aggravare le disuguaglianze interne e replicare o amplificare i *bias* preesistenti. Alla luce di questi aspetti, diversi studiosi suggeriscono di evitare un approccio puramente sostitutivo in cui l’algoritmo surclassa completamente la figura manageriale; propongono, invece, un modello di divisione simbiotica del lavoro tra *manager* umani e *manager* algoritmici, in cui ciascuno contribuisce secondo le proprie specifiche capacità. I sistemi algoritmici sono particolarmente efficaci nel gestire spazi decisionali ristretti e ben delimitati, ovvero compiti definiti con criteri oggettivi di ottimizzazione, mentre gli esseri umani eccellono in spazi decisionali ampi e poco strutturati che richiedono intuito, visione d’insieme e gestione di variabili qualitative. Pertanto, un approccio efficace vede l’algoritmo assumere ruoli di micro-coordinamento come la raccolta dati, il monitoraggio di specifici indicatori; mentre i *manager* umani, si concentrano su attività di macro-coordinamento come definire obiettivi, gestire le eccezioni, motivare il personale, risolvere conflitti e situazioni nuove. Tale modello organizzativo pone le radici sulla flessibilità supervisionata che consente di tener conto dell’aspetto etico dell’Intelligenza Artificiale. Infatti, un algoritmo al fine di massimizzare la produttività potrebbe spingere i lavoratori in uno stato di esaurimento, tuttavia un *manager* consapevole del benessere del *team* è proiettato verso un’ottica di lungo periodo. Si stanno diffondendo politiche correttive per garantire maggiore equità e per affrontare i rischi connessi all’adozione algoritmica. Ad esempio, la rotazione sistematica dei compiti più gravosi o dei turni meno desiderati evita che l’algoritmo riproduca dinamiche discriminatorie. In parallelo, normative come il Regolamento Generale sulla Protezione dei Dati (GDPR) spingono le aziende verso una maggiore rendicontabilità delle decisioni algoritmiche e una supervisione attiva del loro impatto. Tali evoluzioni mostrano chiaramente che l’efficacia dell’Intelligenza Artificiale dipende dalla sua integrazione in un sistema sociotecnico costituito da regole, ruoli

e cultura organizzativa e, solo all'interno di tale cornice, quest'ultima può davvero contribuire a migliorare il funzionamento delle organizzazioni. Quindi, ad oggi, la vera sfida consiste nel progettare modelli organizzativi in cui il coordinamento algoritmico coniuga l'efficienza operativa con equità, trasparenza e attenzione al benessere delle persone. In quest'ottica, l'Intelligenza Artificiale non si pone come un'alternativa ai meccanismi tradizionali, ma li ridefinisce e li rafforza, aprendo la strada a forme organizzative più flessibili.

1.3 L'IA e la trasformazione dell'*expertise*

Attualmente stiamo assistendo alla cosiddetta “crisi dell'*expertise*”, titolo del libro di Gil Eyal, il quale esplora il crescente scetticismo e sfiducia nelle figure e nelle istituzioni, tradizionalmente associate all'autorevolezza e alla competenza nella società contemporanea. Tale crisi non si esaurisce in una semplice delegittimazione dell'autorità da parte degli esperti, ma riflette un cambiamento strutturale nel modo in cui la conoscenza viene creata, distribuita e accettata. Per affrontare questa crisi è opportuno ridefinire il ruolo degli esperti e ricostruire il rapporto di fiducia tra conoscenza tecnica e sfera pubblica.

I cambiamenti sociali, tecnologici e politici hanno progressivamente eroso la fiducia nelle istituzioni esperte, dalle università agli enti scientifici fino ai governi. Tale declino di fiducia è stato ampliato dall'odierna facilità con cui si accede alle informazioni, spesso non filtrate o non verificate, attraverso i media digitali. La conoscenza esperta viene percepita sempre di più come contestabile, con posizioni opposte, amplificate dal dibattito pubblico e dai *social media*. *Internet* e l'AI hanno, senza dubbio, democratizzato l'accesso alla conoscenza, ma al tempo stesso hanno aumentato la disinformazione dovuta alla difficoltà di distinguere le informazioni derivanti da esperti qualificati e voci non autorevoli. La crisi non è una semplice perdita di fiducia, ma un conflitto sistemico tra autorità tradizionale degli esperti e la crescente consapevolezza pubblica di interpretazioni scientifiche e tecnologiche.

In questo scenario, l'adozione pervasiva dell'Intelligenza Artificiale nelle organizzazioni non solo modifica processi e assetti strutturali, ma incide anche sulla natura stessa del sapere esperto. Evolve la definizione di competenza specialistica e si riconfigura la relazione tra i sistemi intelligenti e i professionisti basati sulla conoscenza – medici, avvocati, ingegneri, analisti, ricercatori – la cui attività si fonda su esperienza, giudizio e capacità interpretativa. In questa fase transitoria, vengono messi in discussione i confini tradizionali tra conoscenza umana e capacità computazionale, generando nuove forme di complementarità, ma anche tensioni rispetto al ruolo dell'esperto. La sfida, quindi, non è solo tecnica, ma culturale e organizzativa:

ridefinire l'*expertise* in un contesto in cui l'autorevolezza non può più essere data per scontata, ma va costruita in dialogo con sistemi intelligenti e con una società più critica e informata.

1.3.1 Ridefinizione del concetto di *expertise*

Il concetto di *expertise* è stato tradizionalmente associato al possesso individuale di un sapere approfondito e di competenze sviluppate attraverso percorsi formativi e pratica sul campo. Secondo tale approccio, l'esperto viene descritto come colui che dispone di un patrimonio cognitivo e di capacità pratiche, incluse intuizioni ed esperienza, che lo rendono in grado di affrontare con efficacia problemi complessi in uno specifico ambito.¹⁸

L'avvento di sistemi di intelligenza artificiale ha posto nuove domande sul significato stesso di competenza, poiché tali sistemi sono in grado di eseguire compiti specialistici con elevati livelli di accuratezza. Se un algoritmo di *deep learning* riesce a diagnosticare malattie cutanee con prestazioni simili a quelle di un dermatologo esperto, o se un sistema è in grado di analizzare contratti, progettare codici o tradurre testi complessi, è lecito chiedersi se tali macchine possano essere considerate esperte nei rispettivi domini. Questo tema pone in rilievo una questione centrale, ovvero ci si chiede se le competenze si possano trasferire a un'entità artificiale, qualora essa disponga di dati e capacità computazionali sufficienti.

Le narrative più diffuse in ambito manageriale e tecnologico tendono a rispondere positivamente, immaginando l'IA come un'estensione automatizzata delle capacità umane in grado di operare in domini sempre più complessi.¹⁹ Tali visioni si fondano su una concezione dell'*expertise* come sapere codificabile e replicabile: un costrutto mentale individuale che può essere emulato se la macchina dispone di informazioni adeguate e di una potenza di elaborazione sufficiente. Tuttavia, studi recenti mettono in discussione questo paradigma, proponendo una lettura dell'*expertise* non come entità statica, ma come processo relazionale. Secondo Pakarinen e Huising (2023), la competenza non risiede esclusivamente nell'individuo, ma si costituisce attraverso una rete di relazioni che coinvolgono altri professionisti, utenti, strumenti e istituzioni. In questa prospettiva, un esperto è tale non soltanto per il suo sapere tecnico, ma perché opera in un contesto sociale e organizzativo che ne riconosce e legittima il ruolo poiché interagisce con colleghi e destinatari del suo lavoro, segue protocolli condivisi e partecipa a una comunità epistemica che ne certifica la competenza.

¹⁸ Freidson, E. (1970). *Professional dominance: The social structure of medical care*. New York: Atherton Press.

¹⁹ Pakarinen, J., & Huising, R. (2023). *Relational expertise and the challenge of AI integration in professional work*.

Dunque, l'introduzione di sistemi di IA in contesti professionali non equivale semplicemente ad aggiungere una nuova fonte di calcolo, bensì comporta una trasformazione nelle dinamiche attraverso cui il sapere viene esercitato, distribuito e riconosciuto poiché modifica le relazioni tra attori, criteri di valutazione e responsabilità decisionali.

Pakarinen e Huising (2023) individuano tre dimensioni critiche che emergono nel passaggio da una visione sostanzialista a una relazionale dell'*expertise* in presenza dell'IA:

1. **Opacità:** l'impossibilità, da parte degli esperti umani, di comprendere a fondo i meccanismi decisionali degli algoritmi può compromettere la fiducia e la legittimazione reciproca. Un medico potrebbe esitare ad affidarsi ad una raccomandazione dell'IA se non è in grado di spiegarla a sé stesso o al paziente, violando così uno dei fondamenti della responsabilità professionale.
2. **Traduzione:** l'integrazione efficace dell'IA richiede che i problemi professionali vengano tradotti in un linguaggio interpretabile dalla macchina, e che gli output algoritmici vengano reinterpretati nel contesto pratico. Questo lavoro di traduzione è esso stesso una forma di *expertise*, che valorizza nuove figure come i *data translator* o i *prompt engineer*. L'*expertise*, quindi, non scompare, ma evolve verso forme meta-cognitive e intermediali.
3. **Responsabilità:** nei contesti professionali tradizionali, l'esperto è responsabile delle proprie decisioni, ma quando entra in gioco l'IA, si pone il problema di chi risponda in caso di errore. Il professionista umano rimane formalmente responsabile, ma il ruolo dell'algoritmo può complicare il processo di decisione, generando nuove tensioni etiche e operative.

Tale cambio di paradigma porta a rivedere l'idea stessa di professione in quanto si passa da una concezione centrata sull'individuo e il suo bagaglio di conoscenze ad una visione focalizzata sulle interazioni, sulle pratiche condivise e sul contesto in cui esse si sviluppano. In questo senso, le professioni non vengono sostituite, ma si adattano progressivamente per incorporare l'IA all'interno delle loro routine operative, ridefinendo ciò che conta come sapere esperto. Esempi concreti confermano questa trasformazione: nell'ingegneria del *software*, alcuni strumenti spostano l'*expertise* dalla scrittura manuale del codice alla supervisione architeturale, al testing e all'ottimizzazione; in ambito medico, la competenza si orienta verso l'integrazione tra dati generati da dispositivi intelligenti e il giudizio clinico complessivo; nel settore legale, l'avvocato combina l'*output* dell'IA con la capacità retorica e strategica di

persuadere una giuria. Questa evoluzione implica un cambiamento nella natura stessa delle competenze richieste; tuttavia, le qualità che rendono l'uomo insostituibile risiedono sempre più nella sfera relazionale, etica e interpretativa. Capacità come l'empatia, il pensiero critico, la creatività e la sintesi interdisciplinare diventano centrali, proprio perché difficilmente replicabili da sistemi algoritmici. Non sorprende, quindi, che si affermi con sempre maggior frequenza che le cosiddette *soft skills* siano oggi considerate competenze chiave, al pari – se non superiori – delle abilità tecniche. Il valore professionale, in quest'ottica, non dipende più solo dal sapere specialistico, ma dalla capacità di dialogare in modo critico e consapevole con le tecnologie intelligenti.

In sintesi, l'*expertise* nell'era dell'Intelligenza Artificiale va compresa come un processo ibrido e dinamico, costituito dall'intersezione tra agenti umani, sistemi intelligenti e contesti organizzativi. Sono profonde le implicazioni per la formazione, la valutazione e l'organizzazione delle professioni poiché si affermano nuove competenze e quelle tradizionali devono essere reinterpretate in chiave relazionale e adattiva. La sfida non è stabilire se l'IA potrà essere “esperta”, ma comprendere come la sua presenza ridefinisca il modo in cui si diventa, si agisce e si è riconosciuti come esperti.

1.3.2 Relazione tra AI e professionisti *knowledge-based*

Uno dei contesti più rilevanti in cui si osserva l'interazione tra l'Intelligenza Artificiale e i professionisti è quello dell'*expertise* aumentata, in cui i sistemi intelligenti affiancano il lavoro umano offrendo supporto decisionale. In queste configurazioni, la relazione è di tipo collaborativo-complementare poiché da un lato, l'algoritmo aumenta la produttività e l'accuratezza attraverso l'elaborazione di grandi volumi di dati e l'individuazione di *pattern* complessi, mentre dall'altro lato, l'individuo dedica maggior tempo ad aspetti qualitativi interpretando e adattando questi risultati al contesto operativo.²⁰

Tuttavia, questa cooperazione non è sempre fluida e scontata. Ad esempio, Van den Broek, Sergeeva e Huysman hanno documentato l'implementazione di un sistema di *machine learning* in ambito HR, per due anni in una grande organizzazione.²¹ L'obiettivo iniziale era creare un algoritmo in grado di valutare i candidati in modo oggettivo, automatizzando il processo di *screening* e riducendo, se non eliminando, il coinvolgimento dei *recruiter* umani. Il sistema di

²⁰ Daugherty, P. R., & Wilson, H. J. (2018). *Human + Machine: Reimagining work in the age of AI*. Harvard Business Press.

²¹ Van den Broek, E., Sergeeva, A., & Huysman, M. (2022). *When the machine meets the expert: A field study of AI integration in HR decision-making*.

IA elaborava punteggi e classifiche dei candidati sulla base di dati storici, ma i *recruiter*, attraverso analisi qualitative come colloqui, osservazioni e conoscenza del contesto, interpretavano e, in alcuni casi, contestavano i risultati proposti. L'algoritmo non veniva percepito né come un'entità infallibile, né come uno strumento marginale, ma si configurava piuttosto come un agente decisionale co-partecipativo, la cui influenza era bilanciata dall'esperienza dei professionisti. Gli autori descrivono questo assetto come una forma di interdipendenza dialettica tra Intelligenza Artificiale ed *expertise* di dominio, ben lontana dalla visione dicotomica "IA contro esperto" che aveva inizialmente ispirato il progetto.

Questo studio empirico mette in luce alcuni elementi centrali per comprendere l'evoluzione del rapporto tra IA e professionisti. In primo luogo, dimostra che le macchine, seppur potenti, non sono in grado di cogliere pienamente le sfumature di una decisione complessa senza l'intervento umano. Nel caso della selezione del personale, l'IA apprende dalle decisioni del passato, ma non necessariamente riflette le esigenze strategiche attuali dell'organizzazione o le competenze emergenti richieste dal mercato per cui i dati storici, seppur ricchi, possono risultare incompleti o distorti da *bias* preesistenti.

In secondo luogo, emerge con chiarezza che il coinvolgimento attivo dei professionisti nella progettazione e nell'adattamento dei sistemi IA è fondamentale. Nel caso osservato, i *recruiter* non si sono limitati a subire l'introduzione dell'algoritmo, ma hanno contribuito a plasmarlo, fornendo riscontri e partecipando alla definizione dei parametri decisionali. Questo processo ha rafforzato il senso di *ownership* tecnologica e ha ridotto la distanza tra tecnologia e pratica quotidiana. Le evidenze suggeriscono che le organizzazioni dovrebbero promuovere un approccio partecipativo alla progettazione dell'IA, per favorirne l'accettazione e l'efficacia.

Il terzo aspetto è la trasformazione dei ruoli professionali. I *recruiter* hanno acquisito conoscenze di base in *data science* per dialogare efficacemente con i progettisti, mentre questi ultimi hanno imparato a comprendere le logiche operative del contesto HR. Questa convergenza ha dato origine a nuove figure professionali ibride, capaci di muoversi tra dominio tecnico e applicativo e tale fenomeno è stato riscontrato in diversi ambiti.

Infine, l'adozione dell'IA introduce una dimensione di verifica incrociata e apprendimento continuo. Il confronto tra la valutazione algoritmica e quella umana genera uno spazio riflessivo che stimola l'approfondimento critico. Infatti, quando i due giudizi coincidono, si rafforza la fiducia nel risultato, ma quando divergono, si apre un'opportunità di riconsiderazione e questo confronto può portare sia a correggere errori umani, sia a individuare limiti o *bias* nel modello algoritmico. Nel caso studiato, i *recruiter* hanno talvolta modificato il proprio giudizio grazie al supporto dell'IA, mentre, in altri casi hanno segnalato *pattern* discriminatori appresi dal

sistema, promuovendo così un miglioramento congiunto delle competenze umane e delle performance tecniche.

Iniziative di automazione totale spesso si sono scontrate con limiti operativi e culturali, obbligando a reintrodurre l'esperienza umana nel ciclo decisionale dopo resistenze e risultati deludenti. Tali episodi hanno talvolta incrinato la fiducia dei professionisti nelle tecnologie, alimentando un atteggiamento scettico. Tuttavia, la storia dell'introduzione dell'IA è punteggiata da entusiasmi e delusioni, ma occorre trovare un equilibrio di fiducia in modo tale da non alimentare lo scetticismo che può impedire di adottare tecnologie in contesti critici e neanche supportare l'adozione acritica e ingenua delle nuove tecnologie.

Inoltre, un altro aspetto da sottolineare è l'impatto diretto sui percorsi formativi poiché automatizzando attività di *routine*, esiste il rischio che i giovani professionisti non acquisiscano le basi necessarie per sviluppare una competenza piena. È quindi essenziale bilanciare l'uso dell'IA con momenti formativi tradizionali in modo da evitare una sorta di *deskilling* o mancato *upskilling*.

Infine, l'IA contribuisce a ridefinire i confini tra professioni. Alcune attività prima riservate a esperti possono essere svolte da operatori meno qualificati grazie al supporto algoritmico, con effetti di democratizzazione delle competenze. Questa realtà, però, potrebbe far nascere tensioni tra le professioni, a causa della minaccia di perdita del proprio status professionale. Allo stesso tempo, nuove specializzazioni emergono, come il *data ethicist* e il *data analyst*, cambiando l'assetto di potere e prestigio tra figure professionali all'interno delle organizzazioni *AI-driven*. Coloro che gestiscono tali piattaforme potrebbero acquisire più influenza rispetto ai detentori dell'*expertise* tradizionale, finché questi ultimi non integrano nel proprio ruolo anche la padronanza di tali strumenti.

In conclusione, la relazione tra IA e professionisti *knowledge-based* non si configura come una semplice sostituzione, ma come un processo di interdipendenza e adattamento reciproco. Le migliori evidenze suggeriscono che quando i professionisti abbracciano l'IA come alleato e ne guidano l'implementazione con la loro esperienza, si ottiene un miglioramento delle *performance* e nascono nuove pratiche di lavoro ibride più efficaci di quelle precedenti.

Al contrario, quando l'IA viene percepita dai dipendenti come sostituto antagonista, si generano resistenze, errori e perdita di fiducia che possono vanificare i potenziali benefici.

Pertanto, si rende necessario per le organizzazioni ridefinire i ruoli e le responsabilità nei processi decisionali misti e adottare strategie di *change management* come progetti co-pilota dove l'esperto lavora a fianco con i *data scientist*.

È altresì fondamentale comunicare con trasparenza che l'obiettivo di tale trasformazione digitale non è quello di annichilire la componente umana, ma di potenziarne la portata. Come sintetizzato da Davenport e Kirby (2016)²², il valore umano rimane imprescindibile, a condizione che sia in grado di evolvere nella direzione di una collaborazione consapevole con le macchine.

I professionisti *knowledge-based*, dal canto loro, dovranno sviluppare meta-competenze orientate alla gestione dell'IA e all'interpretazione critica dei suoi *output*, aggiungendo un nuovo capitolo al proprio repertorio professionale. In questo scenario, l'Intelligenza Artificiale non si configura come agente di discontinuità, ma come catalizzatore di una trasformazione organizzativa in cui le potenzialità umane e quelle computazionali si combinano per raggiungere risultati altrimenti irraggiungibili.

²² Davenport, T. H., & Kirby, J. (2016). *Only Humans Need Apply: Winners and Losers in the Age of Smart Machines*. Harper Business.

CAPITOLO SECONDO: Il Nuovo Modello Teorico

2.1 IA e ridefinizione del coordinamento

Con lo sviluppo crescente di sistemi basati sull'apprendimento automatico, si assiste alla nascita di un nuovo paradigma alternativo in cui il coordinamento non pone più le radici esclusivamente sull'interazione e sulle decisioni umane, ma su forme di coordinamento basate sui dati e sugli algoritmi capaci di operare con una velocità, pervasività e granularità prima impensabili. Proporre un nuovo modello teorico per integrare l'Intelligenza Artificiale nella progettazione organizzativa significa riconoscere che non è solo uno strumento tecnologico, ma parte integrante del sistema di coordinamento dell'organizzazione. In tal senso, l'IA può assumere il ruolo di agente organizzativo che coadiuva o sostituisce l'uomo in alcune funzioni di coordinamento come nell'assegnazione dei compiti, nel monitoraggio dell'avanzamento del lavoro, nella trasmissione di informazioni e direttive. Questo modello teorico deve quindi ampliare i tradizionali paradigmi organizzativi includendo i meccanismi *algorithmic-driven*. Una sfida chiave è comprendere come bilanciare questi nuovi meccanismi con quelli preesistenti, mantenendo coerenza interna ed efficacia. In letteratura emergono concetti come il *coordination latency* e l'*ambient coordination*,²³ utili per interpretare tali trasformazioni. Il primo termine può essere definito come la "latenza di coordinamento", ossia il tempo che intercorre tra il momento in cui sorge la necessità di coordinare un'attività e il momento in cui tale coordinamento effettivamente avviene. Nei modelli tradizionali, questa latenza può essere elevata, si pensi ai ritardi dovuti a riunioni di allineamento, catene decisionali lunghe, comunicazioni asincrone via e-mail. L'IA promette di ridurre drasticamente questa latenza grazie alla capacità di processare grandi volumi di informazioni in tempo reale e di fornire *feedback* istantanei.²⁴ Il sistema organizzativo è reso più istantaneo, data la velocità e la reattività che superano di gran lunga le capacità umane tradizionali. Parallelamente, attraverso la tecnologia intelligente si raggiunge il cosiddetto *ambient coordination*, ossia un coordinamento pervasivo benché poco visibile, che non si attua tramite istruzioni esplicite e formali, ma attraverso un ecosistema digitale che coinvolge tutte le persone. Tale concetto si traduce, quindi, in segnali, notifiche e raccomandazioni generati da sistemi IA integrati negli strumenti di lavoro, i quali guidano quotidianamente il comportamento dei lavoratori. Questi ultimi, infatti, possono essere influenzati e guidati nelle loro azioni senza essere pienamente

²³ Kellogg, K. C., Valentine, M. A., & Christin, A. (2020). *Algorithms at work: The new contested terrain of control*. *Academy of Management Annals*, 14(1), p. 5

consapevoli di questo condizionamento. Ciò significa che parte del coordinamento avviene sullo sfondo poiché l'ambiente tecnologico in cui l'individuo opera, come sensori IoT e applicazioni intelligenti, provvede ad allineare le attività e a segnalare deviazioni, riducendo il bisogno di interventi gerarchici diretti. Tale concetto si può applicare per ottimizzare la gestione del *workflow management*: un membro del team, invece di attendere il *meeting* mattutino con il responsabile per sapere quali attività svolgere, può ricevere suggerimenti aggiornati da un sistema AI che tiene conto dello stato di avanzamento degli altri membri del *team* e delle esigenze emergenti. Questa forma di coordinamento è distribuita, in quanto non c'è un unico coordinatore umano ed è continua nel tempo.

In ultima analisi, è opportuno rivedere alcuni assunti classici sull'organizzazione, infatti sempre più si riduce il margine tra le fasi di pianificazione, comunicazione ed esecuzione. Il coordinamento diventa un processo continuo, nel quale pianificazione ed esecuzione si sovrappongono costantemente mediati dall'algoritmo. Infine, è necessario considerare le implicazioni etiche e manageriali che stanno emergendo: un coordinamento invisibile può essere efficiente, ma potrebbe ridurre la trasparenza percepita dai lavoratori, i quali potrebbero non comprendere le attività assegnate dall'algoritmo, creando asimmetrie informative tra chi controlla gli algoritmi e chi ne è controllato.²⁵ Nel modello teorico proposto, l'Intelligenza Artificiale è concepita come un ulteriore attore di coordinamento, caratterizzato da proprietà distintive quali l'elevata rapidità di elaborazione, la diffusione capillare all'interno dei processi organizzativi e un potenziale grado di opacità, elementi che richiedono un'integrazione attenta nei tradizionali assetti strutturali.

2.1.1 Confronto tra meccanismi tradizionali e *AI-driven*

È utile confrontare il coordinamento *AI-driven* con i meccanismi tradizionali per comprendere a fondo il suo contributo. Questo è possibile attraverso l'analisi di alcune dimensioni chiave: velocità, trasparenza e *accountability*.²⁶

²⁵ Kellogg, K. C., Valentine, M. A., & Christin, A. (2020). *Algorithms at work: The new contested terrain of control*. *Academy of Management Annals*, 14(1), p. 5

²⁶ Pakarinen, J./M., & Huising, R. (2023). *Relational expertise and the challenge of AI integration in professional work*. Working paper

1. Velocità

Nei meccanismi tradizionali, le capacità umane e le strutture formali limitano la velocità di coordinamento. Le decisioni avvengono spesso tramite scambi comunicativi che richiedono tempo, come in riunioni periodiche, per cui la risoluzione di conflitti oppure il coordinamento e la comunicazione tra diversi *team* richiede spesso l'intervento dei rispettivi superiori, con inevitabili ritardi dovuti al processo di approvazione. L'introduzione degli algoritmi ha modificato radicalmente questa dinamica poiché permettono un controllo istantaneo, fornendo *feedback* immediati. Tali sistemi intelligenti sono in grado di raccogliere enormi moli di dati sullo stato delle operazioni, elaborandoli continuamente, consentendo un coordinamento complessivo e continuo rispetto a quello umano. Tale evoluzione riduce drasticamente la *coordination latency*, infatti le azioni correttive che in un sistema umano richiederebbero ore o giorni, possono essere implementate in pochi secondi o minuti. Un esempio concreto di questo fenomeno si individua nell'uso di *software* di *project management* basati su intelligenza artificiale, i quali sono capaci di riassegnare priorità o ridistribuire risorse in risposta a imprevisti, senza dover attendere l'esito di riunioni formali, monitorando l'andamento dei progetti in tempo reale. Così facendo, la latenza tra problema e risposta si riduce notevolmente, limitando l'impatto negativo.

2. Trasparenza

La questione della trasparenza rappresenta un'ulteriore area di differenziazione tra i modelli di coordinamento tradizionali e quelli algoritmici. Nei contesti umani, la trasparenza è spesso garantita poiché i decisori sono noti e possono spiegare le motivazioni delle loro scelte, adattando il messaggio ai diversi interlocutori mentre l'Intelligenza Artificiale, in particolare nelle sue declinazioni basate su *machine learning*, introduce un livello di opacità in quanto gli algoritmi risultano complessi e poco interpretabili da parte degli utenti. I sistemi basati sull'apprendimento automatico, aumentano la disponibilità dei dati per *audit* successivi, grazie alle registrazioni e al tracciamento di tutte le operazioni, ma ciò non garantisce una percezione di maggiore trasparenza tra i lavoratori, soprattutto se le informazioni raccolte non sono rese accessibili o comprensibili. Per cui, sono necessari interventi mirati come spiegazioni semplificate dei processi decisionali e canali di interazione uomo-macchina per ridurre tale opacità in quanto la trasparenza algoritmica si configura non solo come una questione tecnica, ma anche organizzativa e percettiva. Nel confronto tra coordinamento umano e *AI-driven* emerge dunque un *trade-off*: da un lato, l'Intelligenza Artificiale abilita livelli inediti di raccolta

e analisi dei dati, dall'altro introduce nuove forme di opacità e asimmetrie informative che possono erodere la fiducia e la comprensione delle decisioni operative.

3. *Accountability*

Anche l'attribuzione della responsabilità decisionale subisce una profonda trasformazione. Nei modelli convenzionali, l'*accountability* è facilmente identificabile in quanto si conoscono le persone a capo dell'azione eseguita, per cui eventuali fallimenti progettuali o inefficienze vengono imputati in modo diretto ai responsabili gerarchici, come *manager* o capi progetto. Con l'avvento dei sistemi IA, la catena di attribuzione della responsabilità diviene meno intuitiva poiché a fronte di un errore di coordinamento generato da un algoritmo, risulta complesso stabilire la responsabilità: ricade sullo sviluppatore che ha programmato il sistema, sul dirigente che decide di implementarlo o sull'utente che ha eseguito passivamente l'azione? Tale dilemma trova una risposta tramite il concetto di *moral crumple zone*²⁷, secondo cui l'operatore umano viene posto in una posizione di assorbitore delle responsabilità, assumendosi la colpa per decisioni su cui aveva in realtà un controllo limitato. In considerazione di tali criticità, diviene indispensabile adottare nuove strategie di *accountability*, tra cui *audit* algoritmici periodici, ruoli specifici di supervisione dei sistemi intelligenti e procedure che consentano agli utenti di interrogare e contestare o revisionare le decisioni automatizzate.

Kellogg, Valentine, & Christin (2020)²⁸ individuano quattro proprietà fondamentali dell'*algorithmic control* che ne descrivono l'impatto sulla natura del coordinamento: comprensività, istantaneità, interattività e opacità. Le prime due, la capacità di raccogliere dati su vasta scala e di reagire istantaneamente agli *input*, potenziano la reattività dell'organizzazione; l'opacità, al contrario, solleva interrogativi sulla comprensibilità e sulla tracciabilità delle decisioni automatizzate, generando potenziali vulnerabilità sul fronte della trasparenza e della responsabilità.

Un ulteriore elemento distintivo introdotto dai sistemi *AI-driven* riguarda la modalità di distribuzione del coordinamento. A differenza dei modelli tradizionali, in cui il coordinamento è centralizzato e gestito da figure specifiche, l'Intelligenza Artificiale consente una gestione

²⁷ Elish, M. C. (2019). *Moral crumple zones: Cautionary tales in human-robot interaction*. Engaging Science, Technology, and Society, 5, 40–60.

²⁸ Valentine, M. A., Kellogg, K. C., & Christin, A. (2020). *Algorithms at work: The new contested terrain of control*. Academy of Management Annals, 14(1).

diffusa, attraverso una rete di agenti autonomi che prendono decisioni a livello locale. Tale configurazione richiama i principi dell'auto-organizzazione e del coordinamento olistico. Piuttosto che affidarsi esclusivamente all'interazione sociale, come avviene nei modelli di *Holacracy*, il coordinamento è supportato e regolato da un'infrastruttura tecnologica che guida e assiste l'attività dei diversi nodi organizzativi. Studi come quelli di Bernstein et al. (2016)²⁹ suggeriscono modelli ibridi in cui elementi di autonomia si combinano con supporti strutturati, piuttosto che modelli di auto-organizzazione radicale, i quali possono presentare criticità operative. In questo quadro, il concetto di *ambient coordination* assume rilievo poiché il coordinamento non si realizza attraverso ordini espliciti, ma è incorporato nell'ambiente digitale del lavoro. Tutto ciò conduce ad una organizzazione sempre più agile e meno dipendente dall'intervento diretto dei *manager*, i quali mantenendo comunque un elevato livello di orchestrazione delle attività. In conclusione, un modello teorico efficace dovrà bilanciare l'efficienza tecnologica con la salvaguardia delle prerogative umane attraverso strutture di *governance* solide che includono procedure per la verifica delle decisioni algoritmiche. Tutto ciò per attenuare le preoccupazioni legate alla trasparenza dell'algoritmo, all'attribuzione della responsabilità e alla percezione del controllo da parte dei lavoratori.

2.2 IA e divisione del lavoro

La divisione del lavoro sta subendo delle profonde trasformazioni, in particolare nel modo in cui le attività vengono scomposte, distribuite tra unità operative o individui e successivamente integrate per il raggiungimento degli obiettivi comuni. Nella progettazione organizzativa tradizionale, la divisione del lavoro si fonda su principi di specializzazione. I compiti omogenei vengono concentrati all'interno di ruoli e funzioni e i diversi livelli di responsabilità vengono assegnati in base ad una gerarchia delle competenze, in cui l'esperienza e la conoscenza rappresentano i criteri decisionali. L'adozione dell'automazione intelligente introduce la cosiddetta divisione algoritmica del lavoro, caratterizzata da un'elevata atomizzazione dei compiti e da una riallocazione di essi dinamica, in tempo reale. In tale configurazione, l'IA assume il ruolo di coordinatore operativo, suddividendo processi complessi in micro-attività e assegnandole continuamente all'agente, umano o artificiale, che risulta più adatto a eseguirle. Tale trasformazione comporta l'emergere di fenomeni inediti e tra essi si annovera il concetto di *shadow work* algoritmico, ovvero quell'insieme di attività non formalmente riconosciute nei

²⁹ Bernstein, E., Bunch, J., Canner, N., & Lee, M. (2016). *Beyond the Holacracy Hype: The Overwrought Claims—and Actual Promise—of the Next Generation of Self-Managed Teams*. Harvard Business Review, 94(7/8)

modelli organizzativi classici, ma richieste o generate dall'interazione con i sistemi intelligenti. Parallelamente, si evidenzia una crescente necessità di acquisire meta-competenze, intese come capacità di interpretare, interrogare e integrare gli output dell'IA nei propri processi decisionali, competenze sempre più cruciali per operare efficacemente in ambienti digitalizzati. Infine, vengono sradicate le tradizionali gerarchie di competenza, poiché alcune forme di sapere tecnico o esperienziale perdono centralità, mentre emergono nuove configurazioni di *expertise*, ridefinendo i criteri di autorità basata sulla conoscenza all'interno delle organizzazioni.

2.2.1 Divisione algoritmica del lavoro e meta-competenze emergenti

È importante analizzare, in vista dell'utilizzo delle nuove tecnologie, come la suddivisione e l'assegnazione del lavoro si evolve e si modifica in forme sempre più granulari e flessibili. Un algoritmo può analizzare un processo complesso, scomponendolo in singole azioni elementari, le cosiddette micro-attività, che possono poi essere distribuite tra diversi soggetti. Tale atomizzazione dei compiti consente di estendere la rete operativa al di fuori dei confini organizzativi, ad esempio attraverso pratiche di *crowdsourcing*. Un esempio emblematico si riscontra nello sviluppo *software*, in cui il progetto si può suddividere in centinaia di *micro-task* - la scrittura di funzioni specifiche, il testing unitario, la correzione di *bug* isolati - che vengono assegnati a programmatori selezionati dinamicamente sulla base di competenze, disponibilità e carico di lavoro, anche in ambito esterno. L'allocazione dinamica dei compiti, resa possibile dall'IA, comporta che gli incarichi non siano più legati stabilmente a una posizione o funzione organizzativa, ma possono essere interscambiabili tra individui o unità in base a criteri variabili. Questo approccio offre vantaggi evidenti in termini di efficienza poiché permette di sostituire intervalli non produttivi dei dipendenti con attività utili in altre aree, e di mobilitare competenze specialistiche in modo trasversale, senza necessità di passaggi gerarchici intermedi. Inoltre, consente una scalabilità notevole perché i sistemi intelligenti grazie alla loro visione aggiornata delle competenze e degli attori, reagiscono rapidamente a variazioni nella domanda attivando risorse *freelance* o riallocando personale interno. Tuttavia, questo modello non è privo di criticità, infatti l'estrema frammentazione del lavoro rischia di compromettere la visione d'insieme da parte dei lavoratori. Gli individui, eseguendo *micro-task* isolati, possono perdere il senso complessivo del prodotto o servizio cui contribuiscono e tale perdita di contesto può intaccare la motivazione intrinseca e ostacolare i processi di apprendimento professionale.

Kellogg, Valentine e Christin (2020)³⁰ evidenziano come la pratica della micro-suddivisione del lavoro, può ridurre le opportunità di formazione di reti sociali e di sviluppo di capitale sociale interno, elementi tradizionalmente associati alla crescita professionale e alla costruzione dell'influenza individuale nelle organizzazioni. In secondo luogo, l'assegnazione algoritmica delle attività può creare un senso di precarietà interna, infatti se gli incarichi sono decisi dinamicamente dall'algoritmo, i confini dei ruoli organizzativi tendono a dissolversi. Da un lato, ciò favorisce l'agilità organizzativa, ma dall'altro, può alimentare nei lavoratori una percezione di instabilità, data l'assenza di una chiara progressione di carriera e la competizione costante, sia interna che esterna, per accedere ai compiti più qualificanti. Alcuni individui potrebbero percepirsi come ingranaggi facilmente sostituibili in un sistema più ampio, con conseguenze negative sul clima organizzativo e sul senso di appartenenza.

Un concetto emergente, in ambito algoritmico, associato alla divisione del lavoro è quello di *shadow work*. Tradizionalmente, il termine "lavoro ombra" indicava quell'insieme di attività non riconosciute formalmente nei contratti di lavoro o nelle strutture organizzative, spesso svolte senza remunerazione aggiuntiva da clienti o dipendenti, come nel caso del *self-service* o delle attività amministrative svolte fuori dall'orario di lavoro. Con l'avvento dell'IA, il *shadow work* assume nuove connotazioni. Quest'ultimo termine può riferirsi al lavoro nascosto svolto dagli algoritmi stessi, riducendo la visibilità delle attività lavorative tradizionali, d'altro canto, l'efficienza dei sistemi intelligenti richiede spesso un supporto umano non formalizzato. I lavoratori sono chiamati a svolgere compiti supplementari che non facevano parte delle loro mansioni originarie, come il *training* degli algoritmi attraverso *feedback* e correzioni, la raccolta e preparazione di dati di alta qualità. Il rischio del *shadow work* algoritmico è duplice: da una parte, sovraccarica i lavoratori con compiti ulteriori, spesso non riconosciuti né formalmente né economicamente; dall'altra rende più complesso misurare e riconoscere il reale contributo umano nei risultati complessivi, poiché una parte significativa del lavoro diventa invisibile agli strumenti di monitoraggio tradizionali. Tale situazione comporta che il merito di un prodotto o servizio finalizzato viene attribuito prevalentemente al sistema automatizzato, senza considerare l'intervento umano indispensabile nella fase di supporto, correzione e supervisione. Affinché le organizzazioni possano gestire in modo equo ed efficace il *shadow work* algoritmico, sarà necessario sviluppare nuovi strumenti di rilevazione e riconoscimento del

³⁰ B Valentine, M. A., Kellogg, K. C., & Christin, A. (2020). *Algorithms at work: The new contested terrain of control*. *Academy of Management Annals*, 14(1).

lavoro invisibile, nonché aggiornare le descrizioni di ruolo includendo esplicitamente le attività di *training*, supervisione e monitoraggio dei sistemi intelligenti. In tal caso, sarà possibile preservare la motivazione, il riconoscimento professionale e la sostenibilità organizzativa in ambienti a forte componente automatizzata.

Inoltre, è opportuno sottolineare che l'Intelligenza Artificiale all'interno delle organizzazioni comporta un cambiamento nel profilo delle competenze richieste ai lavoratori dato dall'emergere di nuove esigenze formative. La capacità di eseguire un compito specifico diventa meno centrale rispetto alla capacità di adattarsi, apprendere rapidamente e collaborare efficacemente con sistemi automatizzati. Tali capacità, definite meta-competenze, rappresentano abilità di ordine superiore, le quali permettono agli individui di navigare con successo in ambienti lavorativi dinamici e tecnologicamente complessi. Un esempio di tali competenze riguarda l'alfabetizzazione dei dati e degli algoritmi. I lavoratori devono essere in grado di comprendere i principi di base del funzionamento degli algoritmi, interpretare correttamente gli *output* generati dai sistemi e riconoscere eventuali distorsioni o limiti insiti nei dati. Non si tratta necessariamente di acquisire competenze di programmazione avanzata, ma piuttosto di sviluppare una comprensione delle logiche sottostanti agli strumenti digitali utilizzati. Studi recenti, come quello di Kaynak (2023)³¹, mostrano come anche individui privi di formazione tecnica possano acquisire competenze operative significative attraverso percorsi di apprendimento come *bootcamp* intensivi, *peer learning* e piattaforme di risorse aperte.³² Tale evidenza sottolinea come l'auto-apprendimento rapido di concetti tecnici, rappresenti una competenza chiave per il successo professionale, resa possibile dall'utilizzo strategico delle risorse distribuite e accessibili. Un'altra meta-competenza cruciale è la flessibilità cognitiva, ovvero la capacità di adattarsi rapidamente a cambiamenti di compiti, strumenti e procedure. In un contesto in cui l'allocazione dei *task* è dinamica e guidata da algoritmi, il lavoratore deve essere pronto a modificare il proprio *set* di attività, ad apprendere nuove procedure operative e a riconfigurare il proprio contributo all'interno dei flussi di lavoro in evoluzione. Questa flessibilità si accompagna a una forma di resilienza organizzativa, intesa come capacità di affrontare l'incertezza e di percepire il cambiamento non come minaccia, bensì come opportunità di crescita. Inoltre, diventa essenziale sviluppare la capacità di collaborazione

³¹ Kaynak, E. (2023). *Leveraging learning collectives: How novice outsiders break into an occupation*. *Organization Science*, 35(3), 948–973

³² Rostain, M. & Huising, R. (2023). *Vicarious Coding: Breaching Computational Opacity in the Digital Era*. *Academy of Management Journal* (in stampa journals.aom.org pg 42).

uomo-macchina. L'efficacia operativa in ambienti tecnologici dipende dalla capacità dei lavoratori di integrare i suggerimenti e le analisi fornite dall'Intelligenza Artificiale nei propri processi decisionali, mantenendo al contempo un'autonomia di giudizio sufficiente per riconoscere quando è necessario intervenire criticamente. Collaborare con sistemi di apprendimento automatico non implica una delega passiva delle decisioni, bensì un'interazione consapevole che valorizzi i punti di forza della tecnologia senza rinunciare al discernimento umano. Un ulteriore aspetto di rilievo è rappresentato dalla capacità di comunicazione interdominio, ossia dall'abilità di fungere da ponte tra esperti tecnici e specialisti di dominio. In molte organizzazioni, si assiste a una crescente esigenza di figure che sappiano tradurre le esigenze operative in specifiche tecniche per i *team* di sviluppo IA e, inversamente, interpretare gli *output* algoritmici in chiave utile per il business. La letteratura, in particolare Levina e Vaast (2005),³³ ha definito questa capacità come *boundary spanning competence*, sottolineandone il ruolo strategico nell'integrazione dei diversi saperi e nella gestione delle interfacce tra domini disciplinari differenti. Infine, il pensiero critico ed etico verso l'IA rappresenta una meta-competenza imprescindibile. L'interazione quotidiana con sistemi intelligenti richiede la capacità di valutare criticamente gli *output* algoritmici, di riconoscere potenziali errori e di riflettere sulle implicazioni etiche delle decisioni. Un lavoratore deve essere in grado di interrogarsi sulla correttezza e sull'impatto delle raccomandazioni generate da un sistema predittivo, sollevando dubbi e proponendo alternative quando necessario. Lo sviluppo di questa sensibilità critica diventa fondamentale non solo per la tutela dell'equità e della trasparenza, ma anche per il mantenimento della responsabilità professionale in ambienti automatizzati. La costruzione di queste meta-competenze avviene sempre più attraverso modalità di apprendimento continuo, informale e distribuito. In ultima analisi, investire su di esse rappresenta una strategia fondamentale per trasformare l'introduzione dell'intelligenza artificiale da fattore di rischio a leva di *empowerment* e crescita organizzativa.

L'introduzione dell'intelligenza artificiale nei contesti organizzativi sta producendo effetti rilevanti non soltanto sulla divisione del lavoro, ma anche sulla configurazione delle gerarchie di competenza. Tradizionalmente, l'autorità professionale era strettamente correlata all'esperienza accumulata e alla padronanza di competenze tecniche o di dominio: figure *senior*, grazie al capitale di conoscenze sviluppato nel tempo, detenevano posizioni di influenza e potere decisionale. Tuttavia, l'adozione di sistemi intelligenti sta progressivamente ridefinendo

³³ Levina, N., & Vaast, E. (2005). *The emergence of boundary spanning competence in practice: Implications for implementation and use of information systems*. MIS Quarterly, 29(2), 335–363

questi equilibri. L'automazione di compiti altamente specializzati, precedentemente prerogativa degli esperti umani, riduce il valore relativo di alcune competenze tradizionali. In alcuni ambiti l'Intelligenza Artificiale è in grado di replicare, e talvolta superare, la capacità analitica dell'uomo, spostando l'accento dall'esperienza tacita alla capacità di interazione e supervisione di tali sistemi. Un ingegnere *senior*, la cui autorevolezza era basata sulla conoscenza operativa maturata in anni di servizio, può vedere il proprio ruolo trasformarsi in quello di supervisore di algoritmi, mentre nuove figure, come *data analyst* e specialisti di AI, acquisiscono crescente centralità. Parallelamente, emergono competenze ibride che combinano capacità tecniche e conoscenza dei processi aziendali. Profili capaci di interpretare i dati prodotti dagli algoritmi e di tradurli in decisioni operative acquistano un valore strategico, indipendentemente dall'anzianità anagrafica o dall'appartenenza a un dominio tradizionale. Questo fenomeno porta all'affermazione di nuove élite professionali, in cui la capacità di integrare saperi diversi risulta premiante rispetto alla specializzazione verticale classica.³⁴

Nicolini et al. (2017)³⁵ evidenziano come l'introduzione di tecnologie avanzate trasformi le dinamiche di sapere-potere nelle organizzazioni. Il sapere viene sempre più incorporato negli strumenti digitali, modificando le modalità attraverso cui l'*expertise* viene riconosciuta e valorizzata. In tale scenario, la rapidità di apprendimento, la flessibilità cognitiva e la capacità di collaborare efficacemente con i sistemi intelligenti diventano criteri fondamentali per l'avanzamento professionale, talvolta sovvertendo le gerarchie fondate sull'anzianità o sulla competenza tecnica settoriale. La creazione di forme di *expertise* duale, che uniscono competenze tecniche e conoscenze settoriali, rappresenta un fattore strategico poiché sempre più frequentemente, ruoli emergenti richiedono una combinazione di abilità che prima appartenevano a profili distinti. Un esempio emblematico di trasformazione delle gerarchie di competenza è offerto dal fenomeno del *Vicarious Coding*, analizzato da Rostain e Huising (2023).³⁶ In molte organizzazioni, operatori tradizionali acquisiscono competenze tecniche indirette osservando il funzionamento dei sistemi digitali o collaborando con sviluppatori, senza necessariamente ricevere una formazione. Tali competenze "ombra" accrescono il valore operativo degli individui e sfumano i confini tra esperto tecnico e utente operativo, generando una nuova forma di *expertise* distribuita.

³⁴ Pakarinen, M., & Huising, R. (2023). *The Relational Work of Expertise: Reconfiguring the Social Fabric of Knowledge in the Era of Artificial Intelligence*. *Organization Science*, 34(1).

³⁵ Nicolini, D., Mengis, J., & Swan, J. (2017). *Understanding the role of objects in organizing knowledge in organizations: A practice-based view*.

³⁶ Rostain, M. & Huising, R. (2023). *Vicarious Coding: Breaching Computational Opacity in the Digital Era*. *Academy of Management Journal*.

In definitiva, il nuovo assetto lavorativo impone alle organizzazioni una profonda riflessione sulla gestione e sul riconoscimento delle competenze. Da un lato, diventa imprescindibile investire nella formazione continua e nel *reskilling*, al fine di dotare l'intera forza lavoro di competenze digitali di base; dall'altro, occorre valorizzare le nuove competenze ibride e le capacità di integrazione interdisciplinare che emergono spontaneamente nei contesti *AI-rich*. Infine, per evitare che tali sistemi diventino un fattore di ampliamento delle disuguaglianze interne, i modelli organizzativi dovranno favorire lo sviluppo inclusivo delle competenze, utilizzando le stesse tecnologie intelligenti come strumenti di formazione, *mentorship* e crescita diffusa.

2.2.2 *Human-in-the-loop* e nuove forme di expertise

L'espressione *Human in the loop (HITL)* si utilizza per indicare l'intervento umano all'interno di sistemi di Intelligenza Artificiale e sottolinea come il contributo dell'uomo non venga sostituito completamente, ma permane nelle diverse fasi del ciclo decisionale e operativo, con diversi gradi di coinvolgimento. Nonostante le tecnologie siano altamente autonome, è importante individuare il ruolo attivo dell'essere umano nella supervisione, nell'addestramento, nella validazione o nella correzione dell'*output* algoritmico. I modelli HITL si articolano secondo le diverse architetture e sono generalmente riconducibili a tre paradigmi principali.³⁷

- Nei sistemi **AI simbolici**, basati su regole ed ontologie definite esplicitamente, l'uomo è coinvolto a monte come *knowledge engineers* che codifica le regole e aggiorna la base di conoscenza. La competenza umana risiede nell'esperto del dominio che esplicita le regole e nell'ingegnere che le traduce in un sistema formale mentre l'IA simbolica segue pedissequamente ciò che è stato programmato. L'uomo è all'interno di tale processo soprattutto nella fase di progettazione e manutenzione delle regole e l'autonomia dell'IA simbolica è limitata ai confini precedentemente definiti.
- Nei sistemi **connessionisti o sub-simbolici**, come il *machine learning* moderno e reti neurali profonde, il ruolo dell'essere umano assume caratteristiche diverse. In fase di *training*, gli operatori predispongono i dati, etichettano esempi - nel caso di apprendimento supervisionato- e selezionano gli algoritmi da impiegare. Successivamente, il modello elabora autonomamente strutture e *pattern* nei dati. Durante l'impiego operativo,

³⁷ Russell, S. J., & Norvig, P. (2021). *Artificial intelligence: A modern approach* (4th ed.). Pearson.

l'algoritmo è in grado di generare decisioni, ad esempio classificando immagini o formulando previsioni, spesso senza intervento umano istantaneo. In questo contesto, il concetto di *Human-in-the-loop* emerge in altre forme. Per prevenire derive indesiderate, l'uomo viene inserito come supervisore con il compito di monitorare l'*output* dei sistemi e di intervenire in caso di errore o incertezza. In altri casi, si richiede attivamente il coinvolgimento dell'uomo per etichettare nuovi dati ritenuti ambigui, mantenendo così l'essere umano nel ciclo di apprendimento. In questi contesti, l'*expertise* richiesta risulta duplice: da un lato richiede competenza sul dominio per prendere decisioni sui casi complessi, dall'altro necessita di una comprensione del funzionamento algoritmico affinché si possano interpretare le indicazioni.

- I modelli **ibridi** di IA integrano componenti simboliche e connessioniste, o più in generale cercano di fondere l'efficienza dell'automazione con l'intuito umano. In tali sistemi, spesso si disegna deliberatamente un ruolo attivo e costante per l'uomo. Un esempio è nei sistemi di *decision support*: l'IA formula una proposta di decisione, che viene successivamente rivista, approvata o modificata dall'operatore umano; tale *feedback* può inoltre essere incorporato dal modello per raffinare le prestazioni future. L'*expertise* umana richiesta consiste nell'avere una forte competenza in merito decisionale, come ad esempio, capire se un candidato possiede *soft skills* adeguate, e al tempo stesso essere in grado di interagire efficacemente con l'interfaccia algoritmica e interpretarne i dati.

Indipendentemente dal modello specifico, tre esigenze motivano il mantenimento dell'uomo nel processo decisionale: assicurare il controllo e la sicurezza, prevenire o correggere gli errori dell'IA, e integrare le competenze tacite che l'algoritmo non possiede. Questo ultimo punto è di cruciale importanza, infatti molte professioni sono fondate su conoscenze tacite e difficilmente codificabili in regole o dati. L'esclusione totale dell'essere umano dai processi decisionali comporterebbe il rischio di una progressiva perdita del patrimonio cognitivo. I modelli *Human in the loop*, al contrario, si propongono di integrare il contributo intuitivo e qualitativo dell'operatore con le capacità analitiche proprie dei sistemi algoritmici, promuovendo una forma di collaborazione sinergica tra intelligenza naturale e artificiale. Tuttavia, la presenza mal gestita dell'umano nel circuito, può anche avere effetti negativi come il *deskilling*. Tale termine si riferisce al fenomeno per cui i lavoratori perdono le abilità acquisite, o non ne sviluppano di nuove, perché il sistema automatizzato svolge gran parte del lavoro. Ad esempio, nell'industria manifatturiera automatizzata, dato il ruolo centrale delle macchine, gli operatori possono col tempo disimparare certi procedimenti manuali; oppure nel

caso di piloti d'aereo, l'uso intensivo del pilota automatico solleva timori riguardo alla perdita di abilità di volo manuale, cosicché in situazioni di emergenza essi non siano più pronti a reagire con prontezza e maestria. Elish (2019)³⁸ parlando di automazione avanzata avverte non solo del rischio di perdita di abilità, ma introduce anche il concetto di *moral crumple zone*: ossia, l'umano viene tenuto formalmente nel loop, come supervisore ultimo, ma finisce per essere degradato a un ruolo passivo, salvo poi essere considerato responsabile quando vi sono errori. Questa situazione si delinea doppiamente problematica poiché il lavoratore non mantiene aggiornate le sue abilità a causa del lavoro della macchina e in aggiunta potrebbe essere accusato di un errore che probabilmente non era davvero in grado di prevenire attivamente. Ciò può portare a una perdita di fiducia nelle tecnologie e a un senso di ingiustizia tra i lavoratori.

Il concetto di *Vicarious Coding* studiato da Rostain & Huising (2023)³⁹ è un esempio di come i lavoratori possano sviluppare competenze prima considerate fuori dalla loro portata, grazie all'osservazione e alla collaborazione con esperti e tecnologie. Il loro studio ha dimostrato che nonostante gli operatori non avessero una formazione tecnica, imparavano indirettamente a interpretare e capire il codice delle macchine, semplicemente osservando i programmatori. In questo modo, gli operatori sviluppavano la capacità di intervenire su guasti o discrepanze del programma, apportando piccole modifiche al codice. Tale risultato mostra come l'automazione non necessariamente elimina competenze, ma l'interazione costante con coloro che la gestiscono può creare nuova *expertise* diffusa. Rostain e Huising evidenziano come tale forma di competenza tecnica emergente, possa manifestarsi anche in soggetti non qualificati grazie a processi di apprendimento e all'osservazione diretta dell'attività tecnica sul campo. Tale situazione, porta con sé due implicazioni importanti: i lavoratori possono evolvere insieme alla tecnologia, acquisendo *skills* un tempo riservate a specialisti, al contempo l'organizzazione beneficia di personale più versatile e capace di individuare problemi tecnologici senza fare capo necessariamente agli esperti esterni. In una prospettiva di insieme, il *vicarious learning* può essere visto come strumento capace di colmare il *gap* tra sviluppatori di IA e utilizzatori. Alimentare un nuovo tipo di *expertise* ibrida diffusa nell'organizzazione rappresenta una strategia vincente e mitiga l'effetto *silo* in cui gli specialisti capiscono gli algoritmi. Ciò è possibile promuovendo la *job rotation* tra ruoli tecnici e operativi, o creando coppie di lavoro esperto di dominio e *data scientist*, i quali lavorano congiuntamente, trasferendosi conoscenze.

³⁸ Elish, M.C. (2019). *Moral Crumple Zones: Cautionary Tales in Human-Robot Interaction*. Engaging Science, Technology, and Society, 5 pg 40

³⁹ Rostain, M. & Huising, R. (2023). *Vicarious Coding: Breaching Computational Opacity in the Digital Era*. Academy of Management Journal

In caso di errore algoritmico, l'attribuzione della responsabilità a capo dell'uomo rischia di disincentivare l'uso degli strumenti tecnologici, per cui diventa fondamentale ridisegnare le pratiche di responsabilizzazione. Pertanto, è necessario formalizzare la condivisione della responsabilità tra l'attore umano e la macchina, evitando che le conseguenze ricadano interamente sull'operatore. È opportuno che quest'ultimo abbia il tempo e i mezzi adeguati per validare il risultato finale, al contrario la responsabilità risulterebbe fittizia.

In conclusione, il nuovo modello teorico vede l'IA come strumento di amplificazione delle capacità umane, ma con la necessità di ridefinire ruoli professionali e responsabilità in modo da non perdere né le abilità né l'equità nei carichi di responsabilità.

2.3 IA e struttura organizzativa

Dopo aver esplorato le implicazioni dell'IA sul coordinamento e sulla divisione del lavoro, è doveroso inquadrare come le innovazioni si riflettano a livello della struttura organizzativa complessiva. Quest'ultima definisce formalmente la distribuzione dei compiti, delle unità, le correlazioni tra di esse e i meccanismi di potere. Storicamente, per rispondere a esigenze di maggiore flessibilità o innovazione, l'organizzazione ha assistito alla mutazione della sua struttura da forme funzionali e divisionali fino ad arrivare a modelli a matrice, a rete e basati sulle piattaforme. L'avvento dell'Intelligenza Artificiale pone il quesito riguardo alla necessità di ripensare alle strutture tradizionali e la costituzione di modelli organizzativi nuovi o ibridi. Nella pratica emergente, si possono identificare alcune configurazioni strutturali che integrano l'IA in modalità differenti:

- ***AI-Augmented Matrix***: strutture a matrice in cui l'IA supporta il coordinamento tra le dimensioni strutturali, mitigando alcuni problemi classici della matrice come i conflitti di priorità e il sovraccarico decisionale.
- ***AI-Hierarchy***: gerarchie tradizionali nelle quali l'IA è inserita come supporto manageriale e analitico, costituendo di fatto una “gerarchia ibrida” uomo-algoritmo.
- ***AI-Enabled Networks***: organizzazioni a rete rese possibili o potenziate dall'uso di piattaforme e algoritmi che coordinano i nodi della stessa.
- ***AI Governance Teams***: l'inserimento formale di nuove unità dedicate alla gestione dell'IA e dei dati per assicurare che quest'ultima sia allineata con la strategia aziendale.

L'adozione di un modello non esclude che possano coesistere elementi di più modelli in una stessa organizzazione, utilizzando la tecnologia intelligente come strumento che ridefinisce o arricchisce i meccanismi strutturali.

2.3.1 Modelli organizzativi abilitati dall'IA

- *AI-Augmented Matrix*

Nella tradizionale struttura a matrice, i dipendenti rispondono a due linee manageriali, tipicamente una funzionale e una di progetto o prodotto. Tale struttura è nota per i suoi benefici come la flessibilità intrinseca e l'utilizzo efficiente delle competenze su diversi progetti, ma anche per le criticità come la doppia subordinazione, i conflitti di priorità, e l'elevato bisogno di coordinamento. L'integrazione dell'IA rende più agevole la gestione quotidiana di questa complessità, infatti in una AI-Augmented Matrix, l'azienda potrebbe adottare un sistema algoritmico in grado di monitorare costantemente i carichi di lavoro, l'avanzamento dei progetti e le *performance*, fornendo ai *manager* e direttamente ai dipendenti, le informazioni e raccomandazioni per snellire il processo. Questo tipo di coordinamento può ridurre l'attrito tipico tra i diversi *manager* e può fornire trasparenza in tempo reale sulle allocazioni poiché i lavoratori possono essere costantemente aggiornati. Tale implementazione riduce le incomprensioni e le sovrapposizioni, ma chiaramente richiede fiducia nel sistema e una cultura collaborativa. Infatti, i *manager* devono affidarsi ai dati e non vivere ogni riassegnazione dei compiti come una riduzione del loro campo di azione. In tal senso, la matrice aumentata risulta idonea in organizzazioni già inclini alla collaborazione inter-funzionale, fornendo lo strumento tecnico per supportarla. Un altro aspetto è la valutazione delle *performance*, tema da sempre considerato ostico in ambiente matriciale a causa dell'incertezza sull'identificazione dei responsabili riguardo la valutazione del lavoratore. Tale condizione potrebbe essere facilitata dai sistemi intelligenti, i quali possono fornire *input* oggettivi per le valutazioni. Un algoritmo potrebbe analizzare il rispetto delle scadenze, i *bug* risolti, i *feedback*, fornendo al manager funzionale e al manager di progetto evidenze su cui basare la valutazione congiunta in modo che la decisione finale resti umana al fine di considerare elementi qualitativi che l'algoritmo non coglie.

- *AI-Hierarchy*

Molte organizzazioni mantengono una struttura fondamentalmente piramidale, ma l'avvento di algoritmi avanzati ha creato la cosiddetta gerarchia "ibrida" in cui i livelli di *management* tradizionali vengono coadiuvati o parzialmente sostituiti da sistemi algoritmici in alcuni processi decisionali e di controllo. Ciò accade, ad esempio, nelle aziende di logistica o nei magazzini automatizzati in cui il caporeparto assume il ruolo di supervisore delle eccezioni e

gestore di persone in quanto le sue mansioni tradizionali, vengono oramai gestite da un *software*. Siamo di fronte ad un'organizzazione dove l'algoritmo funge da “*manager* intermedio”. Kellogg et al. (2020)⁴⁰ descrivono vari meccanismi di controllo algoritmico, note come le “6R”: *restricting, recommending, recording, rating, rewarding, replacing*. Quest'ultime mostrano come molte attività prima gestite dai *manager*, come ad esempio, dare istruzioni, monitorare, valutare, premiare o sanzionare, possano essere in parte svolte dalle macchine intelligenti. Nell'attività di *recommending*, l'algoritmo raccomanda al lavoratore come ottimizzare la propria prestazione; nel *recording e rating*, il sistema registra in dettaglio tutte le attività e produce punteggi di *performance*; mentre nell'attività di *rewarding* può assegnare ricompense immediate, in particolare *bonus*, in base a regole predefinite.

Uno dei vantaggi rivendicati di questo modello è la scalabilità e coerenza delle decisioni. Un algoritmo applica la *policy* in modo uniforme su larga scala, 24/7, senza affaticarsi o essere influenzato da *bias* momentanei, gestendo un numero di subordinati superiore rispetto al manager umano. Gasser & Almeida (2017)⁴¹ osservano che l'uso di algoritmi nel *management* conduce ad un approccio quasi “*tayloristico*” digitale, dove ogni attività dei lavoratori può essere analizzata e ottimizzata centralmente. Il *Digital Taylorism* replica i principi di Frederick Taylor, ovvero il controllo scientifico del lavoro, il frazionamento e la misurazione rigorosa della *performance* attraverso la potenza dei dati e degli algoritmi. Talvolta, l'AI-Hierarchy assume tale inclinazione caratterizzata da un alto controllo centralizzato e un'autonomia locale ridotta. Esistono però delle criticità da sottolineare, infatti un aspetto chiave è l'accettazione da parte dei lavoratori poiché essere coordinati da un algoritmo impersonale può generare alienazione, resistenze o comportamenti elusivi. Ciò implica che è opportuno far percepire l'equità del sistema.⁴² In una gerarchia ibrida ben bilanciata, l'IA assume il ruolo di *decision support* o *execution support*, ma l'umano mantiene la supervisione critica. Invece, in scenari dove l'algoritmo prende il sopravvento totale sul controllo, possono emergere tensioni e calo di *engagement* dei dipendenti. Infine, tale modello pone anche domande sulla stratificazione organizzativa poiché le organizzazioni avranno sempre meno livelli gerarchici umani, data la capacità del software di gestire un numero di lavoratori superiore senza supervisori.

⁴⁰ Kellogg, K.C., Valentine, M.A., & Christin, A. (2020). *Algorithms at work: The new contested terrain of control*. *Academy of Management Annals*, 14(1), 366-41

⁴¹ Gasser, U. & Almeida, V.A.F. (2017). *A Layered Model for AI Governance*. *IEEE Internet Computing*, 21(6), 58-62.

⁴² Barati, M., & Ansari, B. (2022). *Effects of algorithmic control on power asymmetry and inequality within organizations*. *Journal of Management Control*, 33(4), 525-544.

- *AI-Enabled Networks*

Molte organizzazioni moderne si muovono verso forme a rete, sia internamente tramite strutture a *team* flessibili o *project-based*, sia esternamente attraverso reti di *partner*, alleanze, uso di *freelancer*. L'uso dell'Intelligenza Artificiale funge da abilitatore e orchestratore di queste reti. Un AI-Enabled Network si può immaginare come un insieme di nodi- *team*, individui, aziende- collegati da piattaforme digitali intelligenti che facilitano la collaborazione, lo scambio di informazioni e il coordinamento di attività condivise. Un esempio classico sono le piattaforme digitali tipo marketplace che connettono domanda e offerta in una rete di microimprenditori. L'algoritmo regola il matching tra cliente e fornitore, pricing dinamico e reputazione. La struttura dell'organizzazione è ridotta all'osso sul piano gerarchico, ma l'infrastruttura algoritmica tiene insieme migliaia di attori indipendenti. Anche in contesti aziendali tradizionali, possiamo vedere l'emergere di tale modello. Infatti, un'azienda potrebbe avvalersi di un algoritmo di *HR matching*, identificando le persone adatte per creare un *pool* di specialisti per lo svolgimento di un progetto. Valentine et al. (2017)⁴³ hanno coniato il termine *flash teams* e *flash organizations* per indicare *team* creati da piattaforme digitali, con personale interno ed esterno all'organizzazione, coordinato da strumenti intelligenti. Newell (2015)⁴⁴ afferma come i sistemi decisionali algoritmici possano abbattere le barriere tradizionali, costituendo uno spazio unitario di decisione tra entità che precedentemente operavano a silos. È importante creare una visione integrata dove i dati circolano liberamente in modo che le organizzazioni operino meno a compartimenti stagni. Attività come il *marketing* e la produzione, tipicamente separate, potrebbero collaborare strettamente tramite un algoritmo di *predictive analytics* che in tempo reale analizza le tendenze di vendita e suggerisce aggiustamenti produttivi. Chiaramente, l'AI-Enabled Network richiede fiducia reciproca e interoperabilità. Slayton & Clark-Ginsberg (2018)⁴⁵ evidenziano come la fiducia negli algoritmi sia cruciale, specialmente quando le reti includono soggetti esterni all'organizzazione come *partner* e fornitori in modo tale che essi siano a conoscenza dell'equità di tali sistemi e non pensino che questi avvantaggino indebitamente l'uno o l'altro. Ciò rimanda al bisogno di *governance*, trasparenza e accordi contrattuali chiari attorno agli algoritmi di rete.

⁴³ Valentine, M. A., Kellogg, K. C., & Christin, A. (2020). *Algorithms at work: The new contested terrain of control*. *Academy of Management Annals*, 14*(1), 366–410.

⁴⁴ Newell, S., & Marabelli, M. (2015). Strategic opportunities (and challenges) of algorithmic decision-making: A call for action on the long-term societal effects of 'datification'. *Journal of Strategic Information Systems*, 24(1), 3–14.

⁴⁵ Slayton, R., & Clark-Ginsberg, A. (2018). Beyond regulatory capture: Coproducing expertise for critical infrastructure protection. *Regulation & Governance*, 12(1), 115–130.

- *AI Governance Teams*

Un elemento strutturale, emergente in molte organizzazioni, sono i *team* o comitati dedicati alla *governance* dell'Intelligenza Artificiale. Questi sono gruppi inter-funzionali, spesso collocati a livello centrale, i quali hanno il compito di definire linee guida, supervisionare i progetti AI, assicurare conformità etica e normativa e diffondere le *best practice* nell'uso dell'automazione intelligente. È importante considerare tali gruppi all'interno della struttura organizzativa poiché formalizzano nell'organigramma una responsabilità trasversale, prima inesistente e che attualmente risulta vitale per l'organizzazione. Bailey & Barley (2020)⁴⁶ suggeriscono che per studiare integralmente l'impatto delle tecnologie intelligenti, è opportuno analizzare come esse ridefiniscono ruoli e interazioni organizzative. L'AI Governance Team è un nuovo organo nella struttura, non legato a una funzione tradizionale ma coinvolge rappresentanti dell'IT per le attività tecniche riguardanti gli algoritmi; rappresentanti legali e di compliance per la valutazione dei rischi normativi e di privacy; rappresentanti di HR per l'impatto su lavoro e competenze; *manager* di *business* di varie unità al fine di assicurare che l'IA sia orientata verso obiettivi concreti. In alcuni casi viene delineata anche la figura del Chief AI Officer, il che sottolinea come la gestione dell'IA sia diventata un ruolo dirigenziale di primo livello in alcune organizzazioni. Il compito di tali *team* è molteplice. Dal punto di vista strutturale, essi fungono da meccanismo di collegamento tra diverse parti dell'azienda, evitando che ogni divisione sviluppi soluzioni IA isolate creando duplicazioni o incoerenze. Assicurano, inoltre, che la direzione aziendale abbia una visione unitaria sugli investimenti e sulle priorità in campo computazionale. Inoltre, questi *team* spesso curano l'aspetto formativo in modo che l'adozione dell'IA sia omogenea in tutta la struttura. Promuovo, infatti, programmi di sviluppo come *workshop* per *manager* non tecnici oppure *training* per *data scientists* sulle problematiche di *business*. Nelle grandi aziende *tech* o in settori altamente sensibili vengono creati comitati etici costituiti da consulenti esterni o esperti indipendenti che hanno il compito di segnalare eventuali problemi etici nei progetti AI, attraverso raccomandazioni non vincolanti, pubblicando rapporti sulla condotta dell'azienda. Strutturalmente, significa inserire un controllo di livello istituzionale sull'uso delle nuove tecnologie. Tale modello riconoscere che l'IA non è solo questione tecnica propria del reparto IT, ma un tema organizzativo strategico da gestire con un approccio dedicato.

⁴⁶ Bailey, D. E., & Barley, S. R. (2020). *Beyond design and use: How scholars should study intelligent technologies and work*. Information and Organization, 30(2), 100286.

In conclusione, *l'AI-Augmented Matrix* e *l'AI-Hierarchy* mostrano come soluzioni di intelligenza computazionale possono modificare le forme organizzative interne esistenti, rispettivamente matrice e gerarchia, rendendole più reattive ma anche ponendo questioni di controllo e autonomia. *L'AI-Enabled Network* illustra come l'IA consente forme più fluide e interconnesse di organizzazione, anche al di fuori dei confini tradizionali, mentre gli *AI-Governance Teams* evidenziano la necessità di aggiungere nuove componenti alla struttura per assicurare un utilizzo consapevole e coordinato dell'IA.

2.3.2 Equilibrio tra controllo algoritmico e autonomia umana

Dal principio, i sistemi intelligenti sono stati impiegati all'interno dell'organizzazione per avere maggior controllo ed efficienza sui processi attraverso la standardizzazione, dovuta all'analisi dei dati, e il monitoraggio continuo. Il vantaggio principale promesso dall'Intelligenza Artificiale è liberare i lavoratori da compiti routinari, affinché si possano dedicare ad aspetti creativi e di maggior valore aggiunto. Questo porta al paradosso dell'autonomia algoritmica: se da una parte l'algoritmo può concedere autonomia, dall'altra il fatto che un algoritmo regola strettamente molte attività può di fatto ridurre l'autonomia percepita dai lavoratori. Tale paradosso può condurre il lavoratore a seguire istruzioni rigide, dettate da una macchina, perdendo la possibilità di gestire il proprio lavoro secondo il proprio giudizio.

Kellogg et al. (2020)⁴⁷ notano come alcune forme di controllo algoritmico manipolino i comportamenti in modo sottile ma penetrante. Mediante la *gamification* e messaggi di *nudging*, i dipendenti credono di agire liberamente, ma sono continuamente spinti a comportarsi in un determinato modo ed entro binari ben definiti dall'azienda. Dunque, si parla di un'autonomia sorvegliata poiché vi è un'autonomia di facciata, ma il controllo è sottostante. Il termine *nudging* indica quelle spinte gentili che indirizzano gli individui senza eliminarne la libertà di scelta. Ziewitz (2019) parla di come gli algoritmi possano fungere da *algocratic governance*, cioè agiscono spesso con meccanismi comportamentali sottili. In alcuni casi, si può arrivare a forme più forti chiamate *shoving* o *budging*, ossia spinte più decise o modifiche dell'ambiente tali da rendere di fatto obbligate certe scelte. Il bilanciamento è critico in quanto eccessivi *nudges* possono risultare manipolativi e diminuire l'autonomia percepita, mentre una presenza scarsa non avrà efficacia affinché l'Intelligenza Artificiale guidi il miglioramento dell'azienda.

⁴⁷ Bailey, D. E., & Barley, S. R. (2020). *Beyond design and use: How scholars should study intelligent technologies*. Information and Organization, 30(2), 100286.

Un tema collegato è quello del *Digital Taylorism* contrapposto all'*Adaptive Autonomy*. Come anticipato nel precedente paragrafo, il primo concetto si riferisce ai principi di controllo scientifico del lavoro formulati da Taylor oltre un secolo fa. Ogni compito viene standardizzato e assegnato in base a criteri ottimali e i lavoratori sono valutati con metriche quantitative e spinti a conformarsi ai parametri di efficienza fissati dall'algoritmo. Questo approccio massimizza il controllo e l'efficienza immediata, ma può soffocare l'iniziativa individuale, riducendo il lavoratore ad un esecutore frammentato. In conclusione, l'autonomia dell'operatore diventa minima, mentre l'efficienza complessiva risulta ampia, ma con ben noti problemi di stress e soddisfazione.

Un approccio alternativo è invece l'*Adaptive Autonomy*, in cui l'IA viene usata per rafforzare l'autonomia dei lavoratori e la capacità adattiva dell'organizzazione, invece di controllarla rigidamente. In questo modello, i dati e le analisi fungono da supporto alle decisioni locali, lasciando margini di flessibilità affinché i lavoratori possano adattare le soluzioni tecnologiche al contesto specifico. L'idea sottostante è che le macchine e l'uomo si completino reciprocamente. Infatti, l'autonomia adattiva implica che gli strumenti intelligenti gestiscano le eccezioni insieme all'uomo e non deve essere quest'ultimo a monitorare le carenze dell'IA. In aggiunta, l'*Adaptive Autonomy* riconosce che esistono contesti in cui la creatività, l'innovazione e la conoscenza tacita umana sono fonti di valore, e un eccessivo controllo digitale li soffocherebbe. Chiaramente, la soluzione ottimale è trovare un equilibrio ibrido. In concreto, un'organizzazione può delegare ai *team* locali la libertà di organizzare le attività giornaliere, fornendo in aggiunta strumenti di *feedback* automatico tramite gli algoritmi per monitorare i risultati. I sistemi di *feedback* ibridi rappresentano un meccanismo chiave per trovare questo equilibrio. Tale sistema combina *feedback* algoritmici, rapidi, frequenti e basati su dati, con *feedback* umani, i quali sono contestuali, qualitativi e basati su esperienza. Questo duplice circuito permette di arricchire i dati fornendo una spiegazione adeguata, senza che l'IA diventi una "scatola chiusa". Inoltre, il *feedback* deve essere bidirezionale poiché anche l'uomo deve poter dare un riscontro sull'operato dell'intelligenza computazionale. Nel caso di scelte subottimali dell'algoritmo, è importante segnalare il problema, in modo che tali *input* possano essere utilizzati per migliorare il sistema, affinché si crei un ciclo di apprendimento reciproco uomo-macchina.

In conclusione, la sfida centrale della progettazione organizzativa nell'era moderna è trovare l'equilibrio tra controllo algoritmico e autonomia umana. In questo nuovo contesto, nasce la

tentazione di sfruttare al massimo il potere di monitoraggio e ottimizzazione dell'Intelligenza Artificiale, ma d'altro canto, vi è la consapevolezza che le organizzazioni creano valore grazie all'ingegno e impegno delle persone che influiscono sulle decisioni. Il paradosso dell'autonomia algoritmica si risolve, integrandola in sistemi sociotecnici ben disegnati, nei quali l'autorità è condivisa tra algoritmi e persone in modo complementare.

CAPITOLO TERZO: Analisi Sperimentale

3.1. Obiettivi della ricerca empirica

Nel presente capitolo vengono illustrati gli obiettivi e gli esiti dell'indagine empirica, la quale è volta a verificare il modello teorico delineato in precedenza. Tale elaborato analizza l'impatto dell'Intelligenza Artificiale sulla progettazione organizzativa e sui lavoratori nelle aziende *knowledge-based*. L'obiettivo primario di quest'analisi è provare empiricamente come l'adozione di strumenti intelligenti nei processi aziendali, generi cambiamenti all'interno della struttura organizzativa. Nello specifico, si vuole analizzare come tali tecnologie influenzino i meccanismi di coordinamento delle attività, la divisione del lavoro e le competenze richieste. Successivamente, la ricerca si focalizza su come tali cambiamenti condizionino gli atteggiamenti e il benessere dei dipendenti.

Sono state sviluppate alcune ipotesi di ricerca a seguito della revisione della letteratura. Infatti, si parte dall'idea che l'Intelligenza Artificiale contribuisca a strutturare in modo più chiaro le attività, attraverso l'automazione dei compiti ed il supporto decisionale. Quindi, si è ipotizzato il legame tra utilizzo da parte dei lavoratori di strumenti di intelligenti e una diversa percezione organizzativa, manifestata tramite una relazione positiva tra frequenza d'uso degli strumenti IA e chiarezza nella divisione del lavoro e nella trasformazione delle competenze professionali. Inoltre, la creazione di nuove specializzazioni rende più nitidi i confini dei ruoli e aumenta la consapevolezza del valore aggiunto che i lavoratori possono apportare in un contesto contaminato dall'IA. In secondo luogo, ci si propone di indagare se tali dimensioni organizzative spieghino il meccanismo attraverso cui la nuova tecnologia influisca sul benessere lavorativo. L'ipotesi, infatti, è che tale uso non influisca direttamente sulla soddisfazione o sull'impegno verso l'organizzazione, ma abbia un effetto indiretto attraverso i cambiamenti nella struttura e nelle competenze. Il benessere dei lavoratori non viene migliorato direttamente grazie all'IA, ma aumenta nella misura in cui i nuovi strumenti contribuiscono a organizzare meglio il lavoro, a valorizzare le competenze individuali e a far percepire un impatto positivo sui risultati organizzativi. Questo meccanismo rifletterebbe un processo di complementarità uomo-IA in cui la tecnologia potenzia le capacità umane e ne accresce il senso di efficacia, con ricadute benefiche sul coinvolgimento e la soddisfazione personale. Infine, la ricerca considera possibili differenze individuali in questi fenomeni poiché si ipotizza che l'impatto dell'IA possa variare in funzione di caratteristiche socio-demografiche o di ruolo. Per esempio, i lavoratori in posizioni manageriali o con diverse esperienze potrebbero adottare l'IA con frequenza diversa e trarne benefici differenti rispetto ai ruoli operativi; oppure, variabili

come il genere e la generazione potrebbero moderare il rapporto tra IA e atteggiamenti; infine, anche da quanto tempo l'IA è presente in azienda potrebbe creare differenze percettive all'interno dei gruppi. Tali ipotesi riflettono la volontà di comprendere se alcune categorie siano maggiormente avvantaggiate o penalizzate e quindi se l'adozione dell'IA agisca in modo uniforme su tutti i lavoratori.

In sintesi, gli obiettivi empirici di questa ricerca sono:

- misurare il grado di adozione dell'IA e valutarne le correlazioni con dimensioni organizzative chiave - il coordinamento, la divisione del lavoro, le competenze, l'impatto percepito- e con gli esiti per i lavoratori, ovvero *Commitment*, la *Work Identification* e *Job Satisfaction*;
- testare modelli esplicativi, tramite regressioni e analisi di mediazione, che verifichino l'effetto di predittori organizzativi e individuali sulla frequenza di utilizzo dell'IA e sulla soddisfazione lavorativa;
- esaminare la presenza di moderazioni, ossia se il legame tra IA e i suoi effetti varia in base a fattori come il genere, la generazione di appartenenza o la posizione ricoperta nell'azienda.

3.2 Metodologia della ricerca

3.2.1 Il settore

Il settore che è emerso come più strategico per condurre un'indagine empirica sul tema è stato il settore della consulenza e dei servizi professionali. Quest'ultimo presenta un livello di adozione di strumenti AI sempre più crescente nei processi, motivo per il quale possiede un *corpus* di pubblicazioni recenti ma non ancora saturo, per cui tale studio può apportare novità. La consulenza, inoltre, è intrinsecamente un contesto *knowledge-based* in cui coordinamento, *expertise* e divisione del lavoro sono aspetti centrali e già influenzati dai nuovi strumenti tecnologici. In una prospettiva di lungo periodo, la consulenza è destinata a essere plasmata dall'AI generativa e da soluzioni di automazione cognitiva e tale situazione garantisce attualità e longevità alla ricerca. Inoltre, è importante sottolineare la praticabilità di tale studio. Infatti, i consulenti e i professionisti sono facilmente accessibili tramite *survey* online, grazie all'ambiente giovanile e tecnologico, incline a parlare delle trasformazioni digitali in atto nel proprio lavoro. Quindi, il settore offre alte probabilità di ottenere *insight* ricchi sulle dimensioni teoriche da analizzare; un ambiente abbastanza omogeneo di *knowledge work* per generalizzare i risultati e una facilità logistica superiore rispetto a settori più chiusi. Questa decisione non esclude che altri settori siano anch'essi molto rilevanti, anzi possono essere presi in considerazioni per studi futuri o comparativi.

In conclusione, la consulenza consente di massimizzare la fattibilità empirica senza sacrificare la significatività teorica, risultando dunque la più indicata per un'indagine approfondita su come gli strumenti di intelligenza artificiale riconfigurino il coordinamento, la divisione del lavoro, le competenze e l'impatto organizzativo nelle aziende basate sulla conoscenza.

3.2.2 Disegno della ricerca e campione

Al fine di indagare gli obiettivi sopra delineati, è stato utilizzato un disegno di ricerca principalmente quantitativo insieme ad un contributo qualitativo. È stato somministrato un questionario contenente domande a risposta chiusa e una domanda aperta in cui i partecipanti potevano esprimere la propria opinione riguardo l'impatto dell'intelligenza artificiale nel proprio lavoro. Il questionario è stato somministrato a dipendenti di aziende *knowledge-based* del settore della consulenza che impiegano strumenti di IA come parte integrante del lavoro quotidiano. Al fine di catturare una visione eterogenea sull'adozione di sistemi intelligenti in azienda, il campione ha compreso personale con ruoli manageriali, ma anche funzionari e profili

operativi. I partecipanti sono stati selezionati attraverso la piattaforma *Prolific*, la quale mette in contatto ricercatori e persone disponibili a partecipare studi e sondaggi. La partecipazione è stata volontaria e anonima, garantendo la riservatezza delle risposte.

Il campione preso in esame è formato da 105 rispondenti, aventi caratteristiche disomogenee, ma rappresentative del personale impiegato in aziende *knowledge-based*. La composizione per sesso è bilanciata (55 donne, 50 uomini), mentre in ambito generazionale, i *Millennials* sono il gruppo prevalente (61.0% del campione), seguiti in misura analoga da Generazione X e Genenerazione Z e da un gruppo più ridotto di *Baby Boomers*.

Il livello di istruzione è complessivamente elevato, infatti circa il 83% dei rispondenti possiede un titolo universitario di primo o secondo livello - 47.6% *Bachelor's degree* + 35.2% *Master's degree* -, con percentuali minori ma presenti di titoli post-universitari : *Postgraduate course*, *Post graduated diploma* e *PhD*.

I ruoli aziendali coprono diversi livelli gerarchici e funzioni. Le macroaree più rappresentative sono *Technology / Digital Transformation* (26.7%), seguita dalle funzioni *HR/People & Organization* e *Strategy / Corporate Development*. Questa composizione consente di analizzare come l'adozione e l'impatto dell'IA si manifestino attraverso contesti funzionali differenti e tra ruoli con diversa responsabilità, offrendo una base solida per esplorare eterogenee dinamiche di coordinamento, divisione del lavoro e trasformazione dell'*expertise*.

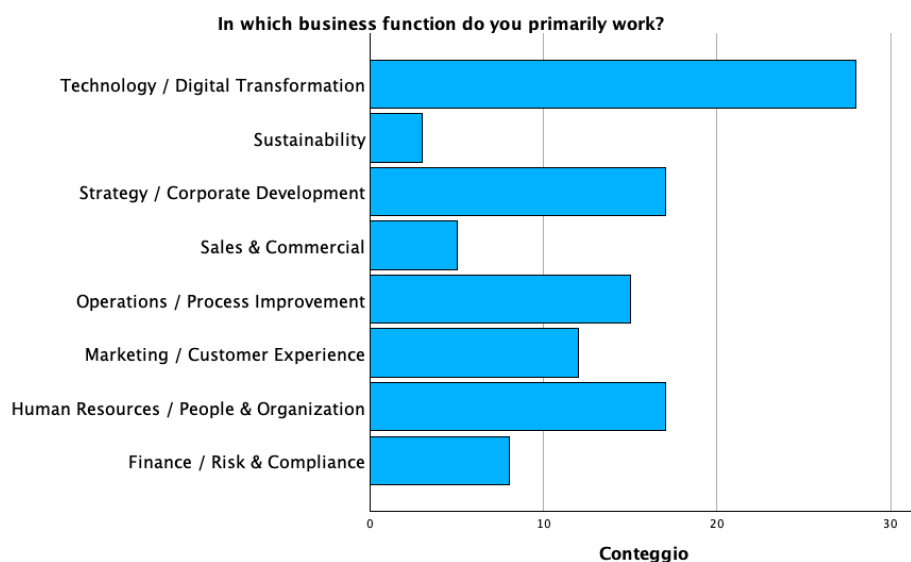


Tabella 1: Aree aziendali in cui lavorano i rispondenti, SPSS

Di seguito viene riportata una panoramica descrittiva del campione oggetto dell'indagine, con le principali caratteristiche demografiche e professionali.

Variabili	Campione (N=105)	Percentuale (%)
Genere		
Maschio	50	47,6
Femmina	55	52,4
Preferisco non dirlo	0	0
Età		
< 25 anni (Generazione Z)	14	13,3
25-44 anni (Millennials)	64	61
45-54 anni (Generazione X)	14	13,3
55- >65 anni (Baby Boobers)	13	12,4
Educazione		
Laurea triennale	50	47,6
Corso post-laurea	5	4,8
Diploma post-laurea	9	8,6
Laurea Magistrale	37	35,2
Dottorato di ricerca o equivalente	4	3,8
Ruolo		
Dirigenti	10	9,5
Managers/quadri	38	36,2
Operativi/Addetti	57	54,3
Anni in azienda		
< 1 anno	20	19
Da 1 a 5 anni	54	51,4
Da 6 a 10 anni	21	20
Da 11 a 20 anni	8	7,6
> 20 anni	2	1,9

Tabella 2: Statistiche descrittive del campione, elaborazione personale

Analizzando i grafici sotto riportati, emerge un quadro significativo sull'adozione e sull'utilizzo delle soluzioni di Intelligenza Artificiale all'interno delle organizzazioni rispondenti. È interessante sottolineare che la maggior parte dei partecipanti è familiare con tali strumenti già da diverso tempo. Infatti, il 41,9% dichiara che la propria organizzazione ha implementato l'IA da 1 a 3 anni, mentre le aziende che l'hanno introdotta più recentemente, da 6 mesi a 1 anno rappresentano il 28,6%, seguiti in misura minore da coloro che la utilizzano da diversi anni (oltre 3 anni rappresentano il 14,3%) o da meno di 6 mesi (11,4%). Solo una quota residuale costituita dal 3,8% segnala che la propria azienda non ha ancora adottato tali sistemi.

Il secondo grafico rappresenta la frequenza di interazione individuale con strumenti di IA ed emerge che la maggioranza del campione, il 41%, dichiara di utilizzarli frequentemente mentre il 37,1% ne fa un uso quotidiano. In misura minore, il 16,2% li utilizza occasionalmente e solo una piccola parte raramente o mai. Tali risultati suggeriscono che le organizzazioni hanno già intrapreso in larga parte un percorso di implementazione dell'IA e che l'uso operativo da parte dei dipendenti è ormai diffuso e sistematico, segnalando un livello di maturità tecnologica relativamente avanzato.

How long has your organization been using AI-based solutions?

105 risposte

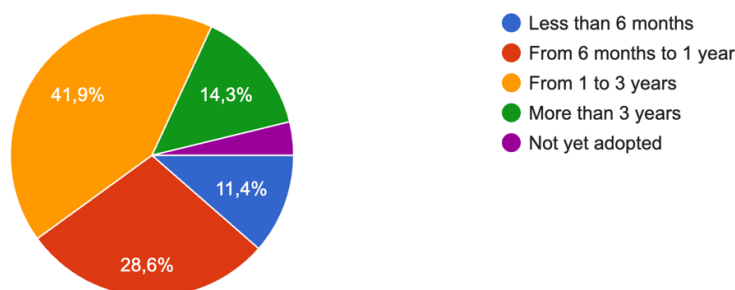


Tabella 3: Periodo di adozione dell'IA in azienda, risposte del questionario di *Google Form*

How often do you interact with AI-based tools?

105 risposte

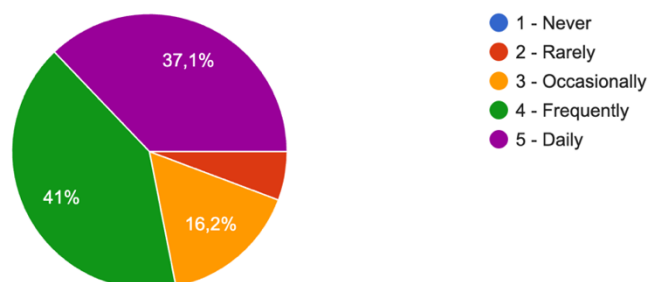


Tabella 4: Uso personale degli strumenti di IA, risposte del questionario di *Google Form*

3.3.3 Il questionario

Il questionario somministrato ai dipendenti è volto a raccogliere i dati sulle abitudini di utilizzo dell'AI e le percezioni soggettive riguardo ai meccanismi organizzativi e agli esiti attitudinali. Le domande poste all'inizio raccolgono informazioni demografiche e inerenti alla frequenza di utilizzo di strumenti di intelligenza generativa al fine di descrivere il campione. Successivamente il questionario si suddivide in diverse sezioni, ciascuna delle quali è costituita da quattro o cinque domande chiuse su scala *Likert* (1= Per niente d'accordo; 5= Completamente d'accordo). Per ciascun rispondente, è stato calcolato numericamente il valore medio della variabile di riferimento in modo tale da ottenere un indicatore sintetico della variabile latente. Infine, per arricchire l'interpretazione dei risultati quantitativi, il questionario si conclude con una domanda a risposta aperta per raccogliere percezioni più dettagliate da parte dei rispondenti. Le variabili prese in considerazione durante le analisi, sono state:

- Frequenza di utilizzo dell'IA:

questa variabile, indicata come *AI_Often*, rispondeva alla domanda: “*How often do you interact with AI-based tools in your daily work?*”. Misura la frequenza con cui il rispondente interagisce con tali strumenti e fornisce un indicatore quantitativo dell'adozione individuale dell'IA in ambito lavorativo.

- Meccanismi di coordinamento:

questa variabile misura in che modo l'IA viene percepita come meccanismo di coordinamento interno, sostituendo o integrando i tradizionali meccanismi organizzativi. Le domande indagano se l'IA automatizza la distribuzione dei compiti, invia notifiche o raccomandazioni, riduce la necessità di riunioni, funge da facilitatore invisibile delle interazioni tra colleghi e abilita forme di supervisione costante.

- Divisione del lavoro e Gerarchia:

la variabile riflette il modo in cui l'IA ridisegna la divisione del lavoro e le dinamiche gerarchiche, spostando il confine tra attività umane e automatizzate. Esplora se la nuova tecnologia modifica l'autonomia decisionale e delinea più chiaramente le mansioni, specializzandole o suddividendole. Le domande indagano sulla riorganizzazione dei flussi di lavoro e sui processi decisionali.

- Evoluzione dell'*expertise*:
misura la trasformazione della competenza professionale a seguito dell'integrazione dell'IA, tra perdita di *expertise* tradizionale e sviluppo di nuove forme di specializzazione. Tale variabile prende in considerazione l'acquisizione di nuove competenze, il ruolo dell'IA come supporto al giudizio professionale, le difficoltà di apprendimento per i nuovi assunti, la riduzione della necessità di esperti tradizionali e l'emergere di nuove figure professionali come il *data translator* e l'*AI explainer*.
- Impatto organizzativo:
rileva la percezione dei lavoratori rispetto all'impatto dell'IA sull'organizzazione, bilanciando trasparenza, fiducia e senso di controllo. Le domande esplorano la comprensione dei criteri decisionali dell'IA, il senso di monitoraggio costante, la fiducia nell'imparzialità delle decisioni e la capacità dell'azienda di gestire efficacemente l'integrazione uomo-macchina.
- Soddisfazione and percezioni dopo l'adozione dell'AI:
questa sezione valuta gli esiti soggettivi dell'adozione dell'IA, misurando il *Commitment*, la *Work Identification* e la *Job Satisfaction*. La variabile cattura la risposta soggettiva dei lavoratori, esplorando se l'introduzione dell'IA ha aumentato l'impegno verso l'organizzazione, il senso di identificazione con il lavoro e la soddisfazione generale.

3.3 Risultati dell'analisi empirica

3.3.1 Analisi dei dati

Dopo aver raccolto un numero significativo di risposte al questionario, i dati sono stati esportati ed elaborati con il software IBM SPSS. Il piano di analisi statistica ha seguito diverse fasi coerenti con gli obiettivi di ricerca:

- **Analisi descrittiva e correlazionale:**

in primo luogo, sono state calcolate le statistiche descrittive come la media, la deviazione standard, il minimo e il massimo per tutte le variabili rilevate al fine di fotografare il profilo medio del campione. Successivamente è stata condotta un'analisi delle correlazioni di *Pearson* tra tutte le variabili principali - *IA_Often*, *AI_Coordination*, *Division_of_Labor*, *Expertise*, *Organizational_Impact*, *Commitment*, *Work_Identification*, *Job_Satisfaction* - per individuare le associazioni bivariate significative. La matrice di correlazione ha permesso di verificare le ipotesi iniziali in merito alle relazioni tra adozione dell'IA, cambiamenti organizzativi e *outcome* per i lavoratori. Ciò costituisce una base per selezionare i predittori nelle analisi successive.

- **Analisi di regressione multipla:**

per testare l'effetto simultaneo di più variabili indipendenti su una dipendente, sono stati stimati modelli di regressione multipla. In particolare, si sono impostati due insiemi di modelli. Il primo esplora come variabile dipendente la frequenza di utilizzo dell'IA (*AI_Often*), per identificare quali fattori organizzativi e personali ne influenzino significativamente l'adozione. Il secondo prende in considerazione la soddisfazione lavorativa (*Job_Satisfaction*) come variabile dipendente, per determinare il peso relativo delle diverse dimensioni e delle caratteristiche individuali nello spiegare il benessere lavorativo. Per ciascuna regressione, è stata adottata una strategia gerarchica a blocchi: nel Modello 1 sono state inserite solo le variabili organizzative e percettive di interesse; nel Modello 2 si sono aggiunte, come controlli e potenziali predittori aggiuntivi, anche le variabili sociodemografiche, quali genere, istruzione, posizione, area di *business*. Tale procedura consente di verificare se l'introduzione di fattori individuali produce un miglioramento significativo nel modello predittivo. I coefficienti di regressione standardizzati e i relativi *p-value* sono stati esaminati per identificare quali predittori avessero effetti significativi sulle variabili dipendenti, ponendo

attenzione anche ad eventuali segni negativi inaspettati. Prima di interpretare i coefficienti, è stata verificata l'assenza di multicollinearità tra i predittori, tramite l'indice VIF e la tolleranza, così da assicurare la stabilità delle stime.

- **Analisi di mediazione:**

dato l'interesse per gli effetti indiretti dell'adozione dell'IA sugli *outcome*, quali il *Commitment* e la *Job Satisfaction*, si è proceduto con un'analisi di mediazione multipla usando l'approccio di Preacher e Hayes attraverso l'estensione di SPSS "Process". Si sono condotte due analisi distinte, modificando la variabile dipendente Y e mantenendo la stessa variabile X e gli stessi mediatori seriali. In particolare, si è ipotizzato un percorso mediato in cui la variabile X= *AI_Often* influisce sulla variabile Y di *outcome* attraverso una catena di mediatori rappresentata dalle dimensioni organizzative considerate: M1=*AI_Coordination*; M2 = *Division_of_Labor*, M3= *Expertise*; M4= *Organizational_Impact*. È stato utilizzato il modello 6 per la mediazione seriale, con stima degli effetti indiretti tramite *bootstrapping* per calcolare gli intervalli di confidenza al 95% degli effetti mediati. In tal modo si verifica se *AI_Often* ha un effetto diretto significativo su Y e/o effetti indiretti significativi attraverso uno o più dei mediatori.

- **Analisi di mediazione-moderata:** per esplorare le differenze tra sottogruppi, si è testato se il genere, la generazione e altre variabili categoriali potessero moderare le relazioni chiave nel modello. Ad esempio, ci si è chiesti se la catena di mediazione identificata valesse in ugual modo per uomini e donne. A tal fine, è stato impiegato un modello di *moderated mediation* (PROCESS Macro, Modello 84) in cui il genere è inserito come variabile di moderazione sull'effetto della tratta del percorso mediato che è risultato significativo. In concreto, si è valutato se l'effetto indiretto delle catene differisse significativamente per i due generi o per le diverse generazioni, esaminando il cosiddetto indice di mediazione moderata fornito da PROCESS, il quale indica se la differenza tra gli effetti indiretti nei gruppi è statisticamente diversa da zero. Analogamente, è stata esplorata la moderazione della posizione lavorativa sui percorsi trovati, per verificare se i lavoratori di livello gerarchico diverso sperimentassero impatti differenti dall'adozione dell'IA. Infine, è stato considerato come moderatore anche il tempo di implementazione dell'IA in azienda, distinto in fasce temporali (Organ_AI 1 = < 6 mesi, Organ_AI 5= non ancora adottata).

3.3.2 Analisi descrittiva e correlazione

In un primo momento è stata svolta un'analisi descrittiva comprensiva di tutte le variabili del *dataset* per testare la normalità delle variabili, condizione necessaria per svolgere le analisi successive. Da tale analisi emerge che l'asimmetria mostra valori compresi tra -0.80 e +0.22, mentre la curtosi varia tra -0.90 e +0.30. Entrambe le statistiche rientrano nei *range* comunemente accettati per la normalità statistica (± 1), suggerendo una sufficiente approssimazione alla normalità per l'utilizzo di tecniche parametriche come la regressione lineare multipla. In particolare, dalle tabelle riportate in Appendice, le variabili *AI_Coordination*, *Division_of_Labor*, *Expertise* e *Organizational_Impact* mostrano distribuzioni bilanciate, con livelli di asimmetria e curtosi prossimi allo zero, e quindi compatibili con l'assunzione di normalità. Anche i costrutti legati al benessere organizzativo (*Commitment*, *Work_Identification*, *Job_Satisfaction*) rientrano nei limiti accettabili: pur mostrando lievi deviazioni dalla simmetria perfetta, non emergono distorsioni tali da compromettere l'utilizzo della regressione multipla.

Inoltre, è opportuno sottolineare che tutte le variabili coinvolte sono state trattate come scale *Likert* a intervalli regolari, e quindi assimilabili a variabili continue, condizione necessaria per Pearson. La numerosità del campione ($N = 105$) può essere considerata adeguata all'impiego di test statistici di tipo parametrico, come la correlazione di Pearson. Infatti, è ampiamente documentato che tali test risultano robusti rispetto a moderate violazioni dell'assunzione di normalità quando la dimensione campionaria supera le 30 unità. Pertanto, lievi scostamenti rispetto alla distribuzione normale non compromettono in modo sostanziale l'affidabilità dei risultati ottenuti.

L'analisi di correlazione consente di verificare se esistono delle relazioni lineari, statisticamente significative tra le variabili prese in considerazione. In linea con le convenzioni statistiche più diffuse, si prende in considerazione un livello di significatività pari a 0,05, corrispondente ad un intervallo di confidenza del 95%. La matrice di correlazione sotto riportata, sintetizza le relazioni bivariate fra tutte le variabili principali prese in considerazione:

Correlazioni

		AlCoord	DivLab	Expert	OrgImp	Commit	WorkId	JobSat	How often do you interact with AI-based tools?
AlCoord	Correlazione di Pearson	1	,721**	,650**	,580**	,648**	,282**	,188	,148
	Sign. (a due code)		<,001	<,001	<,001	<,001	,004	,055	,131
	N	105	105	105	105	105	105	105	105
DivLab	Correlazione di Pearson	,721**	1	,728**	,568**	,602**	,276**	,329**	,301**
	Sign. (a due code)	<,001		<,001	<,001	<,001	,004	<,001	,002
	N	105	105	105	105	105	105	105	105
Expert	Correlazione di Pearson	,650**	,728**	1	,595**	,573**	,252**	,214*	,235*
	Sign. (a due code)	<,001	<,001		<,001	<,001	,009	,028	,016
	N	105	105	105	105	105	105	105	105
OrgImp	Correlazione di Pearson	,580**	,568**	,595**	1	,622**	,267**	,334**	,225*
	Sign. (a due code)	<,001	<,001	<,001		<,001	,006	<,001	,021
	N	105	105	105	105	105	105	105	105
Commit	Correlazione di Pearson	,648**	,602**	,573**	,622**	1	,270**	,294**	,128
	Sign. (a due code)	<,001	<,001	<,001	<,001		,005	,002	,195
	N	105	105	105	105	105	105	105	105
WorkId	Correlazione di Pearson	,282**	,276**	,252**	,267**	,270**	1	,611**	,048
	Sign. (a due code)	,004	,004	,009	,006	,005		<,001	,625
	N	105	105	105	105	105	105	105	105
JobSat	Correlazione di Pearson	,188	,329**	,214*	,334**	,294**	,611**	1	,114
	Sign. (a due code)	,055	<,001	,028	<,001	,002	<,001		,246
	N	105	105	105	105	105	105	105	105
How often do you interact with AI-based tools?	Correlazione di Pearson	,148	,301**	,235*	,225*	,128	,048	,114	1
	Sign. (a due code)	,131	,002	,016	,021	,195	,625	,246	
	N	105	105	105	105	105	105	105	105

** La correlazione è significativa a livello 0,01 (a due code).

* La correlazione è significativa a livello 0,05 (a due code).

Tabella 5: Correlazioni tra tutte le variabili prese in esame, SPSS

Intervalli di confidenza

	Correlazione di Pearson	Sign. (a due code)	95% intervalli di confidenza (a 2 code) ^a	
			Inferiore	Superiore
AlCoord - DivLab	,721	<,001	,614	,802
AlCoord - Expert	,650	<,001	,524	,749
AlCoord - OrgImp	,580	<,001	,437	,695
AlCoord - Commit	,648	<,001	,522	,747
AlCoord - WorkId	,282	,004	,095	,449
AlCoord - JobSat	,188	,055	-,004	,366
AlCoord - How often do you interact with AI-based tools?	,148	,131	-,045	,330
DivLab - Expert	,728	<,001	,623	,807
DivLab - OrgImp	,568	<,001	,422	,685
DivLab - Commit	,602	<,001	,464	,712
DivLab - WorkId	,276	,004	,089	,444
DivLab - JobSat	,329	<,001	,146	,489
DivLab - How often do you interact with AI-based tools?	,301	,002	,116	,466
Expert - OrgImp	,595	<,001	,455	,706
Expert - Commit	,573	<,001	,429	,689
Expert - WorkId	,252	,009	,064	,424
Expert - JobSat	,214	,028	,023	,390
Expert - How often do you interact with AI-based tools?	,235	,016	,045	,408
OrgImp - Commit	,622	<,001	,488	,727
OrgImp - WorkId	,267	,006	,080	,437
OrgImp - JobSat	,334	<,001	,152	,494
OrgImp - How often do you interact with AI-based tools?	,225	,021	,035	,399
Commit - WorkId	,270	,005	,082	,439
Commit - JobSat	,294	,002	,108	,459
Commit - How often do you interact with AI-based tools?	,128	,195	-,066	,312
WorkId - JobSat	,611	<,001	,475	,719
WorkId - How often do you interact with AI-based tools?	,048	,625	-,145	,238
JobSat - How often do you interact with AI-based tools?	,114	,246	-,079	,299

a. La stima si basa sulla trasformazione da r a z di Fisher.

Tabella 6: Intervalli di confidenza tra tutte le variabili prese in esame, SPSS

L'analisi delle correlazioni tra la frequenza di interazione con strumenti di intelligenza artificiale- *AI Often* - e le variabili organizzative evidenzia alcune relazioni significative. In particolare, emerge una correlazione positiva moderata con la *Division of Labor* ($r = 0.301$, $p = 0.002$), indicando che un uso più frequente dell'AI si associa a una percezione più chiara e strutturata della suddivisione dei compiti all'interno dell'organizzazione. Inoltre, l'intervallo di confidenza non contiene lo zero, indicando una correlazione significativa. Analogamente, la frequenza di interazione con l'AI mostra una correlazione positiva, seppur più debole, con la percezione di *Expertise* ($r = 0.235$, $p = 0.016$) e con *Organizational Impact* ($r = 0.225$, $p = 0.021$). In entrambi i casi, l'intervallo di confidenza è più stretto e vicino allo zero, motivo per il quale la forza della relazione è più debole. Ciò suggerisce che, pur con un'intensità minore, un maggiore utilizzo degli strumenti di Intelligenza Artificiale si accompagna a una più elevata percezione dell'evoluzione delle competenze professionali e a una maggior consapevolezza dell'AI nei processi organizzativi e decisionali.

Nel complesso, questi risultati evidenziano come l'esposizione frequente all'AI non influisca uniformemente su tutte le dimensioni indagate, ma tenda a rafforzare soprattutto la percezione della strutturazione del lavoro e, in misura più contenuta, la trasformazione delle competenze individuali e l'impatto percepito sull'organizzazione.

Spostando l'attenzione sulle variabili organizzative, la relazione più forte che emerge è tra *AI_Coordination* e *Division_of_Labor* ($r = 0.721$, $p < 0.001$), indicando che un coordinamento efficace mediato dall'AI è associato a una divisione del lavoro più chiara. Relazioni altrettanto forti emergono tra *AI_Coordination* con *Expertise* ($r = 0.650$, $p < 0.001$) e *Commitment* ($r = 0.648$, $p < 0.001$), suggerendo che il coordinamento digitale contribuisce sia alla trasformazione delle competenze sia al rafforzamento del legame emotivo e motivazionale con l'organizzazione.

La *Division_of_Labor* è fortemente collegata all'*Expertise* ($r = 0.728$, $p < 0.001$) e mostra associazioni positive con *Organizational Impact* ($r = 0.568$, $p < 0.001$) e *Commitment* ($r = 0.602$, $p < 0.001$), evidenziando come la strutturazione dei compiti tramite l'AI favorisca lo sviluppo e la specializzazione di competenze; il cambiamento organizzativo e il legame verso l'azienda. Anche le correlazioni tra *Expertise*, *Organizational Impact* ($r = 0.595$, $p < 0.001$) e *Commitment* ($r = 0.573$, $p < 0.001$) confermano questo legame.

Sul versante del benessere, la variabile *Work Identification* si associa fortemente alla *Job Satisfaction* ($r = 0.611$, $p < 0.001$). Tale relazione è in linea con i modelli motivazionali: sentirsi

parte del proprio ruolo aumenta la soddisfazione. Invece, il legame tra *Commitment* e *Job_Satisfaction* esiste ma è meno forte ($r = 0,294$, $p = 0,002$).

Relazioni più deboli, ma significative, mostrano che un buon coordinamento algoritmico e una strutturazione del lavoro chiara mediata dall'Intelligenza Artificiale, influenzano positivamente l'identificazione del lavoratore con il proprio ruolo e la soddisfazione lavorativa, seppur in misura minore: *AI_Coordination* – *Work_Identification* ($r=0,282, p=0,004$) e *Division_of_Labor* – *Work_Identification* ($r = 0,276$, $p = 0,004$).

Infine, il cambiamento delle competenze a seguito di tale trasformazione digitale, ha un impatto modesto sull'identificazione con il proprio lavoro e una maggior soddisfazione lavorativa: *Expertise* – *Work_Identification* ($r = 0,252$, $p = 0,009$) e *Expertise* – *Job_Satisfaction* ($r = 0,214$, $p=0,028$). Questo suggerisce che i lavoratori che percepiscono un'evoluzione delle proprie competenze a seguito dell'introduzione dell'IA tendono anche a sentirsi leggermente più coinvolti nel proprio lavoro e più soddisfatti, forse perché avvertono accresciuta la propria professionalità e valore.

In sintesi, la matrice di correlazione evidenzia *pattern* coerenti con il modello teorico. L'adozione frequente dell'IA è associata a specifici cambiamenti positivi, i quali a loro volta sono fortemente intercorrelati e collegati con esiti di maggiore benessere lavorativo. Non emergono invece correlazioni dirette tra l'uso dei nuovi strumenti e il *welfare* lavorativo, il che lascia presagire la presenza di effetti indiretti. Le correlazioni forniscono dunque un primo supporto alle ipotesi: confermano che l'IA non lascia invariato il contesto organizzativo, bensì si accompagna a trasformazioni tangibili nella percezione dei ruoli e delle competenze, e che queste trasformazioni sono strettamente collegate al coinvolgimento e alla soddisfazione dei lavoratori.

3.3.3 Regressione multipla con la frequenza di utilizzo dell'AI come variabile dipendente

Per approfondire quali fattori spiegano la frequenza di utilizzo degli strumenti di IA da parte dei dipendenti, sono stati stimati due modelli di regressione multipla con *AI_Often* come variabile dipendente:

Modello 1

H₀: le variabili organizzative (*AI_Coord*, *DivLab*, *Expert*, *OrgImp*) non hanno alcun effetto significativo sulla frequenza di interazione con strumenti AI.

H₁: almeno una delle variabili organizzative ha un effetto significativo sulla frequenza di interazione con strumenti AI.

Modello 2

H₀: nessuna delle variabili organizzative né delle variabili attributive (genere, educazione, posizione, area) ha un effetto significativo sulla frequenza di interazione con strumenti AI.

H₁: almeno una tra le variabili organizzative o attributive ha un effetto significativo sulla frequenza di interazione con strumenti AI.

Riepilogo del modello									
Modello	R	R-quadro	R-quadro adattato	Errore std. della stima	Modifica R-quadro	Statistiche delle modifiche			Sign. Modifica F
						Modifica F	gl1	gl2	
1	,333 ^a	,111	,076	,838	,111	3,123	4	100	,018
2	,508 ^b	,258	,103	,826	,147	1,219	14	86	,277

a. Predittori: (costante), OrgImp, DivLab, AICoord, Expert

b. Predittori: (costante), OrgImp, DivLab, AICoord, Expert, CurrentPosition=Executives, Education=Bachelor's degree, BusinessFunction=Sustainability, BusinessFunction=Operations / Process Improvement, Education=PhD or equivalent, BusinessFunction=Strategy / Corporate Development, Education=Postgraduate course (, BusinessFunction=Sales & Commercial, BusinessFunction=Marketing / Customer Experience, CurrentPosition=Managers, BusinessFunction=Finance / Risk & Compliance, Gender=Male, BusinessFunction=Human Resources / People & Organization, Education=Master's degree

Tabella 7: Riepilogo modello di regressione multipla con *AI_Often*= variabile dipendente, SPSS

ANOVA ^a						
Modello		Somma dei quadrati	gl	Media quadratica	F	Sign.
1	Regressione	8,779	4	2,195	3,123	,018 ^b
	Residuo	70,268	100	,703		
	Totale	79,048	104			
2	Regressione	20,415	18	1,134	1,664	,062 ^c
	Residuo	58,633	86	,682		
	Totale	79,048	104			

a. Variabile dipendente: How often do you interact with AI-based tools?

b. Predittori: (costante), OrgImp, DivLab, AICoord, Expert

c. Predittori: (costante), OrgImp, DivLab, AICoord, Expert, CurrentPosition=Executives, Education=Bachelor's degree, BusinessFunction=Sustainability, BusinessFunction=Operations / Process Improvement, Education=PhD or equivalent, BusinessFunction=Strategy / Corporate Development, Education=Postgraduate course (, BusinessFunction=Sales & Commercial, BusinessFunction=Marketing / Customer Experience, CurrentPosition=Managers, BusinessFunction=Finance / Risk & Compliance, Gender=Male, BusinessFunction=Human Resources / People & Organization, Education=Master's degree

Tabella 8: Tabella ANOVA regressione multipla con *AI_Often*= variabile dipendente, SPSS

L'analisi di regressione è stata presentata attraverso due prospettive statistiche: il Riepilogo del modello e la tabella ANOVA. La prima tabella evidenzia la bontà del modello in termini di varianza spiegata (R^2 e R^2 aggiustato) e fornisce le statistiche di cambiamento, utili per verificare se l'aggiunta di ulteriori predittori migliora significativamente la capacità esplicativa del modello. In questo caso, il Modello 1, costruito con i soli predittori di natura organizzativa, risulta significativo ($F = 3,123$; $p = 0,018$), mentre il passaggio al Modello 2, che integra variabili sociodemografiche e di contesto, non produce un incremento significativo ($F = 1,219$; $p = 0,277$). La tabella ANOVA, invece, valuta la significatività complessiva di ciascun modello

rispetto a un modello nullo, evidenziando che il Modello 1 è globalmente significativo ($F = 3,123$; $p = 0,018$) per cui viene accettata l'ipotesi alternativa, mentre il Modello 2 non raggiunge la soglia di significatività statistica ($F = 1,664$; $p = 0,062$) per cui viene accettata l'ipotesi alternativa dal momento che l'aggiunta delle variabili attributive non produce un miglioramento significativo del modello. Nel complesso, i risultati confermano che i fattori organizzativi e legati all'*expertise* sono predittori significativi, mentre l'inclusione di ulteriori variabili individuali e di contesto non apporta un contributo statisticamente rilevante.

Coefficienti ^a								
		Coefficienti non standardizzati		Coefficienti standardizzati			Statistiche di collinearità	
Modello		B	Errore standard	Beta	t	Sign.	Tolleranza	VIF
1	(Costante)	3,010	,385		7,824	<,001		
	AICoord	-,160	,118	-,196	-1,348	,181	,422	2,368
	DivLab	,320	,145	,346	2,201	,030	,360	2,781
	Expert	,045	,166	,040	,270	,788	,407	2,457
	OrgImp	,136	,144	,118	,948	,345	,574	1,742
2	(Costante)	3,585	,515		6,963	<,001		
	AICoord	-,206	,143	-,253	-1,438	,154	,280	3,575
	DivLab	,403	,156	,436	2,589	,011	,304	3,290
	Expert	-,070	,176	-,062	-,395	,694	,349	2,861
	OrgImp	,140	,149	,121	,939	,350	,516	1,939
	Gender=Male	,073	,201	,042	,361	,719	,644	1,552
	Education=Bachelor's degree	-,298	,314	-,172	-,948	,346	,264	3,795
	Education=Master's degree	-,447	,321	-,246	-1,393	,167	,276	3,625
	Education=PhD or equivalent	-,454	,519	-,100	-,876	,384	,659	1,518
	Education=Postgraduate course (	-1,127	,500	-,277	-2,255	,027	,573	1,745
	CurrentPosition=Executives	,741	,296	,251	2,498	,014	,857	1,167
	CurrentPosition=Managers	,083	,190	,046	,438	,662	,782	1,279
	BusinessFunction=Finance / Risk & Compliance	,246	,360	,075	,682	,497	,712	1,405
	BusinessFunction=Human Resources / People & Organization	-,136	,277	-,058	-,490	,625	,624	1,603
	BusinessFunction=Marketing / Customer Experience	-,553	,303	-,203	-1,825	,071	,699	1,432
	BusinessFunction=Operations / Process Improvement	-,247	,288	-,100	-,857	,394	,638	1,568
	BusinessFunction=Sales & Commercial	-,351	,431	-,086	-,814	,418	,770	1,299
	BusinessFunction=Strategy / Corporate Development	-,004	,267	-,002	-,014	,988	,670	1,493
	BusinessFunction=Sustainability	-,127	,510	-,024	-,249	,804	,899	1,113

a. Variabile dipendente: How often do you interact with AI-based tools?

Tabella 9: Tabella dei coefficienti della regressione multipla con *AI_Often*= variabile dipendente, SPSS

Nel Modello 1, analizzando nello specifico i coefficienti, emerge come significativo soltanto *Division_of_Labor* ($\beta = 0,346$, $p = 0,030$), mentre tutte le altre variabili non raggiungono la significatività statistica. Questo risultato indica che già nel modello di base la strutturazione del lavoro rappresenta un predittore rilevante della frequenza di interazione con strumenti di Intelligenza Artificiale.

Nel Modello 2, che include sia le variabili organizzative sia quelle attributive, emergono come significativi alcuni predittori: *Division_of_Labor* ($\beta = 0,436$, $p = 0,011$) e la posizione lavorativa di *Executives* ($\beta = 0,251$, $p = 0,014$), che mostra un livello più alto di *AI adoption* rispetto agli *Operatives*. Inoltre, tra le variabili educative, si rileva un effetto negativo significativo per *Education = Postgraduate course* ($\beta = -0,277$, $p = 0,027$), questo implica che chi possiede un *Postgraduate course* interagisce con strumenti AI meno frequentemente rispetto alla categoria di riferimento (*Postgraduate diploma*). Tale quadro mette in evidenza che la frequenza di interazione con strumenti di Intelligenza Artificiale è spiegata principalmente dalla percezione di una chiara divisione dei compiti, dal ricoprire ruoli di vertice all'interno dell'organizzazione e, in misura minore, dal percorso formativo intrapreso.

Nel confronto tra le diverse aree operative, tenendo in considerazione la categoria *Technology/Digital* come *baseline*, non emergono differenze statisticamente significative, con l'eccezione di una tendenza nel *Marketing/Customer Experience* ($\beta = -0,203$, $p = 0,071$) verso livelli più bassi rispetto a *Technology*. Questo risultato, pur non superando la soglia di significatività convenzionale ($p < .05$), suggerisce che in alcune aree funzionali l'adozione dell'IA può essere meno radicata rispetto a quelle orientate al digitale.

Tutte le altre variabili risultano non significative: *Organizational_Impact* ($\beta = 0,121$, $p = 0,350$), *AI_Coordination* ($\beta = -0,253$, $p = 0,154$), *Expertise* ($\beta = -0,062$, $p = 0,694$) e anche le variabili attributive (genere, altri livelli di istruzione, altre aree funzionali). Questo suggerisce che, pur essendo concettualmente rilevanti, tali fattori non hanno un peso statistico sufficiente nel predire la frequenza di utilizzo degli strumenti AI.

In sintesi, i risultati indicano che l'interazione con l'AI non dipende genericamente da tutte le caratteristiche organizzative o personali, ma è associata in modo più netto a tre elementi: la strutturazione del lavoro tramite la divisione dei compiti, la posizione apicale rivestita dal lavoratore e, in senso negativo, il possesso di un *Postgraduate course* rispetto al *Postgraduate diploma*.

Analizzando le statistiche di collinearità, la tolleranza rappresenta la proporzione di varianza di una variabile indipendente che non è spiegata dalle altre variabili nel modello. Nel modello, i valori di tolleranza si pongono nel *range* da 0,264 a 0,899, per cui non scendono sotto la soglia critica dello 0,20. Tale situazione indica che non ci sono problemi seri di collinearità: ogni predittore conserva una quota di varianza unica sufficiente. Il VIF rappresenta l'inverso della tolleranza, ovvero indica in che misura la varianza stimata di un coefficiente di regressione sia

influenzata dalla presenza di correlazioni lineari tra i predittori inclusi nel modello. Valori oltre 10 - alcuni autori dicono già oltre 5- indicano multicollinearità preoccupante. Nel modello in analisi i valori oscillano da 1,113 a 3,795, per cui si trovano ben al di sotto della soglia critica.

		Variabili escluse ^a				Statistiche di collinearità		
Modello		Beta in	t	Sign.	Correlazione parziale	Tolleranza	VIF	Tolleranza minima
1	Gender=Male	,086 ^b	,819	,415	,082	,809	1,236	,352
	Education=Bachelor's degree	,078 ^b	,822	,413	,082	,991	1,009	,359
	Education=Master's degree	-,062 ^b	-,649	,518	-,065	,987	1,013	,359
	Education=PhD or equivalent	-,010 ^b	-,100	,921	-,010	,979	1,021	,356
	Education=Postgraduate course (	-,178 ^b	-1,887	,062	-,186	,974	1,027	,356
	CurrentPosition=Executives	,203 ^b	2,192	,031	,215	,998	1,002	,359
	CurrentPosition=Managers	-,072 ^b	-,745	,458	-,075	,947	1,056	,351
	BusinessFunction=Finance / Risk & Compliance	,119 ^b	1,176	,242	,117	,866	1,154	,335
	BusinessFunction=Human Resources / People & Organization	-,026 ^b	-,275	,784	-,028	,978	1,022	,352
	BusinessFunction=Marketing / Customer Experience	-,126 ^b	-1,311	,193	-,131	,954	1,048	,355
	BusinessFunction=Operations / Process Improvement	-,072 ^b	-,754	,452	-,076	,979	1,021	,356
	BusinessFunction=Sales & Commercial	-,069 ^b	-,706	,482	-,071	,924	1,082	,358
	BusinessFunction=Strategy / Corporate Development	,089 ^b	,928	,356	,093	,964	1,037	,359
	BusinessFunction=Sustainability	,001 ^b	,014	,989	,001	,979	1,022	,358

a. Variabile dipendente: How often do you interact with AI-based tools?

b. Predittori nel modello: (costante), OrgImp, DivLab, AICoord, Expert

Tabella 10: Tabella delle variabili della regressione multipla con *AI_Often*= variabile dipendente, SPSS

La tabella delle variabili escluse conferma che i fattori demografici e attributivi non hanno un contributo statisticamente significativo nella previsione della frequenza di interazione con strumenti di IA. I loro valori di Beta ipotetici sono molto bassi e non significativi, e le correlazioni parziali sono prossime allo zero. Inoltre, i valori di tolleranza elevati e i VIF prossimi all'unità dimostrano che non vi sono problemi di multicollinearità. L'esclusione di queste variabili indica che le differenze nella frequenza di interazione con l'IA sono spiegate più efficacemente da altre dimensioni organizzative e percettive già incluse nel modello.

3.3.4 Regressione multipla con *Job Satisfaction* come variabile dipendente

Un secondo insieme di modelli di regressione multipla è stato stimato considerando come variabile dipendente la soddisfazione lavorativa, principale indicatore di benessere organizzativo nello studio. Anche in questo caso si sono considerati due modelli gerarchici:

Modello 1

H₀: le variabili organizzative non spiegano in modo statisticamente rilevante i livelli di *Job Satisfaction*

H₁: almeno una delle variabili organizzative spiega in modo statisticamente rilevante i livelli di *Job Satisfaction*

Modello 2

H₀: nessuna delle variabili organizzative né delle variabili attributive contribuisce a spiegare in maniera significativa la variabilità della *Job Satisfaction*.

H₁: almeno una tra le variabili organizzative o attributive contribuisce a spiegare in maniera significativa la variabilità della *Job Satisfaction*.

Riepilogo del modello

Modello	R	R-quadrato	R-quadrato adattato	Errore std. della stima
1	,675 ^a	,455	,422	,742
2	,707 ^b	,500	,422	,741

a. Predittori: (costante), Workld, Expert, Commit, Orglmp, AlCoord, DivLab

b. Predittori: (costante), Workld, Expert, Commit, Orglmp, AlCoord, DivLab, CurrentPosition=Executives, Education=Bachelor's degree, Education=PhD or equivalent, Education=Postgraduate course (, CurrentPosition=Managers, How often do you interact with AI-based tools?, Gender=Male, Education=Master's degree

Tabella 11: Riepilogo del modello della regressione multipla con *Job Satisfaction*= variabile dipendente, SPSS

ANOVA^a

Modello		Somma dei quadrati	gl	Media quadratica	F	Sign.
1	Regressione	45,045	6	7,508	13,639	<,001 ^b
	Residuo	53,945	98	,550		
	Totale	98,990	104			
2	Regressione	49,507	14	3,536	6,432	<,001 ^c
	Residuo	49,483	90	,550		
	Totale	98,990	104			

a. Variabile dipendente: JobSat

b. Predittori: (costante), Workld, Expert, Commit, Orglmp, AlCoord, DivLab

c. Predittori: (costante), Workld, Expert, Commit, Orglmp, AlCoord, DivLab, CurrentPosition=Executives, Education=Bachelor's degree, Education=PhD or equivalent, Education=Postgraduate course (, CurrentPosition=Managers, How often do you interact with AI-based tools?, Gender=Male, Education=Master's degree

Tabella 12: Tabella ANOVA della regressione multipla con *Job Satisfaction*= variabile dipendente, SPSS

L'analisi di regressione multipla con *Job Satisfaction* come variabile dipendente evidenzia un modello complessivamente robusto. Il Modello 1 spiega circa il 45,5% della varianza della soddisfazione lavorativa ($R^2 = 0,455$; R^2 aggiustato = 0,422) ed è altamente significativo ($F = 13,639$; $p < 0,001$). L'aggiunta di variabili attributive e di contesto nel Modello 2 incrementa leggermente il valore di R^2 (0,500), ma non produce un miglioramento sostanziale della capacità esplicativa, come mostrato dal valore di R^2 aggiustato che resta invariato (0,422). Nonostante ciò, anche il secondo modello risulta complessivamente significativo ($F = 6,432$; $p < 0,001$). Nel complesso, i risultati indicano che la soddisfazione lavorativa è spiegata in misura consistente dalle dimensioni organizzative considerate, mentre le variabili sociodemografiche e di ruolo non apportano un contributo aggiuntivo rilevante.

Coefficienti ^a						
		Coefficienti non standardizzati		Coefficienti standardizzati		
Modello		B	Errore standard	Beta	t	Sign.
1	(Costante)	1,200	,399		3,011	,003
	AlCoord	-,249	,110	-,272	-2,270	,025
	DivLab	,318	,130	,307	2,449	,016
	Expert	-,181	,147	-,144	-1,232	,221
	OrgImp	,249	,135	,193	1,847	,068
	Commit	,078	,088	,097	,889	,376
	Workld	,541	,076	,562	7,131	<,001
2	(Costante)	,675	,567		1,190	,237
	AlCoord	-,313	,120	-,343	-2,607	,011
	DivLab	,271	,138	,262	1,966	,052
	Expert	-,096	,153	-,077	-,629	,531
	OrgImp	,302	,139	,234	2,179	,032
	Commit	,062	,089	,077	,697	,487
	Workld	,532	,077	,553	6,944	<,001
	Gender=Male	-,224	,164	-,115	-1,365	,176
	Education=Bachelor's degree	,302	,275	,155	1,097	,276
	Education=Master's degree	,542	,285	,267	1,904	,060
	Education=PhD or equivalent	,269	,459	,053	,586	,560
	Education=Postgraduate course (	,440	,446	,096	,987	,326
	CurrentPosition=Executives	,280	,266	,085	1,050	,297
	CurrentPosition=Managers	,162	,163	,080	,991	,324
	How often do you interact with AI-based tools?	,037	,094	,033	,397	,692

a. Variabile dipendente: JobSat

Tabella 13: Tabella dei coefficienti della regressione multipla con *Job Satisfaction*= variabile dipendente, SPSS

Nello specifico, analizzando i coefficienti emerge *Work_Identification* come il predittore più forte ($\beta = 0,562$, $p < 0,001$), seguito da *Division_of_Labor* con un effetto positivo significativo ($\beta = 0,307$, $p = 0,018$). Un risultato interessante è l'*AI_Coordination*, il quale mostra invece un effetto negativo e significativo ($\beta = -0,272$, $p = 0,025$), suggerendo che una percezione di coordinamento fortemente basata sull'IA possa ridurre la soddisfazione percepita. Tale risultato potrebbe derivare dal fatto che il coordinamento fortemente basato sull'Intelligenza Artificiale tende a ridurre la discrezionalità e il senso di autonomia del lavoratore, configurandosi come una forma di controllo esterno, e al contempo può contribuire a indebolire la dimensione relazionale del lavoro, con ulteriori ricadute negative sul benessere e sulla soddisfazione lavorativa. *Organizational_Impact* risulta al limite della significatività ($\beta = 0,193$, $p = 0,068$), segnalando una tendenza positiva non robusta, mentre *Expertise* e *Commitment* non risultano predittori rilevanti.

Per tale motivo, nel primo modello l'ipotesi nulla viene scartata, per cui la soddisfazione lavorativa si lega soprattutto al senso di identificazione con il proprio lavoro e alla chiarezza dei ruoli, mentre il coordinamento tramite l'IA appare addirittura ridurre la soddisfazione.

Il Modello 2, che include anche variabili attributive, rimane significativo e spiega circa il 50% della varianza ($R^2 = 0,500$), ma l'incremento rispetto al Modello 1 è modesto (+4,5%). Analizzando i coefficienti, *Work_Identification* si conferma come il predittore più rilevante ($\beta = 0,553$, $p < 0,001$), mentre restano significativi *AI_Coordination* ($\beta = -0,343$, $p = 0,011$) e *Division_of_Labor* ($\beta = 0,262$, $p = 0,052$). *Organizational_Impact* diventa positivo e significativo ($\beta = 0,234$, $p = 0,032$), rafforzando l'idea che percepire un impatto sugli obiettivi organizzativi accresca la soddisfazione. Le variabili attributive non raggiungono la significatività, ad eccezione di una vicinanza alla soglia per *Education = Master's degree* ($\beta = 0,267$, $p = 0,060$), che suggerisce un potenziale effetto positivo non pienamente robusto, mentre le altre qualifiche educative, di genere, così come le posizioni manageriali e dirigenziali non mostrano differenze rilevanti.

In conclusione, anche nel Modello 2 viene scartata l'ipotesi nulla e si può affermare che, analizzando il quadro generale, sono le percezioni soggettive legate all'identificazione con il lavoro e a come l'IA incida sulla chiarezza dei compiti, sul coordinamento e sul senso di impatto organizzativo a costituire i predittori più solidi della soddisfazione lavorativa, mentre le variabili demografiche e l'uso dell'AI hanno un peso marginale.

Variabili escluse ^a						
Modello		Beta in	t	Sign.	Correlazione parziale	Statistiche di collinearità Tolleranza
1	Gender=Male	-,090 ^b	-1,085	,280	-,110	,805
	Education=Bachelor's degree	-,049 ^b	-,648	,518	-,066	,987
	Education=Master's degree	,121 ^b	1,620	,108	,162	,984
	Education=PhD or equivalent	-,022 ^b	-,287	,774	-,029	,972
	Education=Postgraduate course (	,021 ^b	,278	,782	,028	,963
	CurrentPosition=Executives	,069 ^b	,916	,362	,093	,989
	CurrentPosition=Managers	,072 ^b	,934	,353	,094	,944
	How often do you interact with AI-based tools?	,015 ^b	,186	,853	,019	,883

a. Variabile dipendente: JobSat

b. Predittori nel modello: (costante), Workld, Expert, Commit, Orglmp, AlCoord, DivLab

Tabella 14: Variabili escluse della regressione multipla con *Job Satisfaction*= variabile dipendente, SPSS

Le variabili attributive escluse dal primo blocco non avrebbero migliorato il modello: tutte presentano valori $p > 0,05$, confermando che il loro contributo predittivo è limitato.

3.3.5 Analisi di mediazione

L'analisi di mediazione è stata condotta con l'obiettivo di comprendere se e in che modo la frequenza di utilizzo degli strumenti di Intelligenza Artificiale (*AI_Often*) influenzi il livello di impegno organizzativo- *Commitment* - e la soddisfazione lavorativa - *Job Satisfaction*- attraverso il coinvolgimento di variabili intermedie.

Sono state svolte due analisi distinte, nelle quali entrambe avevano come variabile indipendente $X = AI_Often$ e come variabili mediatrici $M1-M4 = AI_Coord, DivLab, Expert, OrgImp$. Nella prima analisi la variabile dipendente è stata $Y = Commitment$, mentre nella seconda analisi $Y = JobSatisfaction$.

***** DIRECT AND INDIRECT EFFECTS OF X ON Y *****

Direct effect of X on Y					
Effect	se	t	p	LLCI	ULCI
-,0777	,1006	-,7727	,4415	-,2774	,1219

Indirect effect(s) of X on Y:

	Effect	BootSE	BootLLCI	BootULCI
TOTAL	,2536	,1074	,0358	,4666
Ind1	,0618	,0523	-,0267	,1842
Ind2	,0433	,0414	-,0295	,1339
Ind3	,0044	,0141	-,0210	,0398
Ind4	,0331	,0395	-,0409	,1162
Ind5	,0223	,0290	-,0234	,0931
Ind6	,0047	,0076	-,0087	,0231
Ind7	,0186	,0187	-,0068	,0644
Ind8	,0122	,0162	-,0172	,0500
Ind9	,0101	,0125	-,0072	,0414
Ind10	,0050	,0103	-,0146	,0292
Ind11	,0063	,0117	-,0121	,0361
Ind12	,0052	,0075	-,0071	,0235
Ind13	,0053	,0051	-,0020	,0175
Ind14	,0140	,0088	,0005	,0342
Ind15	,0072	,0075	-,0027	,0265

Indirect effect key:

Ind1	AI_Often	->	AICoord	->	Commit						
Ind2	AI_Often	->	DivLab	->	Commit						
Ind3	AI_Often	->	Expert	->	Commit						
Ind4	AI_Often	->	OrgImp	->	Commit						
Ind5	AI_Often	->	AICoord	->	DivLab	->	Commit				
Ind6	AI_Often	->	AICoord	->	Expert	->	Commit				
Ind7	AI_Often	->	AICoord	->	OrgImp	->	Commit				
Ind8	AI_Often	->	DivLab	->	Expert	->	Commit				
Ind9	AI_Often	->	DivLab	->	OrgImp	->	Commit				
Ind10	AI_Often	->	Expert	->	OrgImp	->	Commit				
Ind11	AI_Often	->	AICoord	->	DivLab	->	Expert	->	Commit		
Ind12	AI_Often	->	AICoord	->	DivLab	->	OrgImp	->	Commit		
Ind13	AI_Often	->	AICoord	->	Expert	->	OrgImp	->	Commit		
Ind14	AI_Often	->	DivLab	->	Expert	->	OrgImp	->	Commit		
Ind15	AI_Often	->	AICoord	->	DivLab	->	Expert	->	OrgImp	->	Commit

Tabella 15: Modello di mediazione con $Y = Commitment$, SPSS

L'analisi di mediazione ha evidenziato che l'uso frequente dell'AI non ha un effetto diretto significativo sul *Commitment* ($\beta = -0,0777$, $p = 0,4415$). Tuttavia, l'effetto complessivo risulta interamente indiretto e si realizza attraverso un percorso specifico. Come si può osservare dai risultati dettagliati delle tabelle riportate in Appendice B, le analisi preliminari identificano un percorso significativo che collega $AI_Often \rightarrow Division_of_Labor$ ($\beta = 0,2148$, $p = 0,0036$), $Division_of_Labor \rightarrow Expertise$ ($\beta = 0,4323$, $p < 0,001$), $Expertise \rightarrow Organizational_Impact$ ($\beta = 0,2979$, $p = 0,0085$) e, infine, $Organizational_Impact \rightarrow Commitment$ ($\beta = 0,5066$, $p =$

0,0007). Questo percorso mediato (Ind14) risulta significativo poiché l'intervallo di confidenza non include lo zero [0,0005; 0,0342].

Nella seconda analisi di mediazione, con $Y=Job\ Satisfaction$, i risultati mostrano che l'effetto diretto della frequenza di utilizzo di strumenti di intelligenza artificiale sulla soddisfazione lavorativa non è significativo ($\beta = -0,0135$, $p = 0,9020$), suggerendo che la frequenza di utilizzo dell'AI non incrementa di per sé la soddisfazione. Tuttavia, dall'analisi emerge che tale effetto risulta indiretto e l'utilizzo dell'AI incide sulla soddisfazione tramite altre variabili.

***** DIRECT AND INDIRECT EFFECTS OF X ON Y *****

Direct effect of X on Y

Effect	se	t	p	LLCI	ULCI
-,0135	,1091	-,1235	,9020	-,2301	,2031

Indirect effect(s) of X on Y:

	Effect	BootSE	BootLLCI	BootULCI
TOTAL	,1412	,0675	,0188	,2815
Ind1	-,0293	,0371	-,1191	,0231
Ind2	,0843	,0532	,0004	,2049
Ind3	-,0050	,0148	-,0407	,0221
Ind4	,0249	,0299	-,0389	,0841
Ind5	,0435	,0361	-,0204	,1216
Ind6	-,0053	,0086	-,0266	,0089
Ind7	,0140	,0164	-,0045	,0582
Ind8	-,0139	,0179	-,0542	,0181
Ind9	,0076	,0110	-,0052	,0370
Ind10	,0038	,0081	-,0101	,0234
Ind11	-,0072	,0124	-,0381	,0115
Ind12	,0039	,0063	-,0048	,0203
Ind13	,0040	,0045	-,0013	,0154
Ind14	,0105	,0076	,0000	,0290
Ind15	,0054	,0064	-,0017	,0229

Indirect effect key:

Ind1	AI_Often	->	AICoord	->	JobSat						
Ind2	AI_Often	->	DivLab	->	JobSat						
Ind3	AI_Often	->	Expert	->	JobSat						
Ind4	AI_Often	->	OrgImp	->	JobSat						
Ind5	AI_Often	->	AICoord	->	DivLab	->	JobSat				
Ind6	AI_Often	->	AICoord	->	Expert	->	JobSat				
Ind7	AI_Often	->	AICoord	->	OrgImp	->	JobSat				
Ind8	AI_Often	->	DivLab	->	Expert	->	JobSat				
Ind9	AI_Often	->	DivLab	->	OrgImp	->	JobSat				
Ind10	AI_Often	->	Expert	->	OrgImp	->	JobSat				
Ind11	AI_Often	->	AICoord	->	DivLab	->	Expert	->	JobSat		
Ind12	AI_Often	->	AICoord	->	DivLab	->	OrgImp	->	JobSat		
Ind13	AI_Often	->	AICoord	->	Expert	->	OrgImp	->	JobSat		
Ind14	AI_Often	->	DivLab	->	Expert	->	OrgImp	->	JobSat		
Ind15	AI_Often	->	AICoord	->	DivLab	->	Expert	->	OrgImp	->	JobSat

Tabella 16: Modello di mediazione con $Y=Job\ Satisfaction$, SPSS

È opportuno sottolineare un aspetto particolarmente rilevante emerso dalle analisi che riguarda la stabilità della catena di mediazione. Infatti, sia considerando come variabile dipendente il *Commitment*, sia adottando la *Job Satisfaction* come *outcome*, il percorso significativo rimane invariato (Ind14): *AI_Often* → *Division_of_Labor* ($\beta = 0,2148$, $p = 0,0036$), *Division_of_Labor* → *Expertise* ($\beta = 0,4323$, $p < 0,001$), *Expertise* → *Organizational_Impact* ($\beta = 0,2979$, $p = 0,0085$), *Organizational_Impact* → *JobSat* ($\beta = 0,3806$, $p = 0,0175$).

Questo risultato suggerisce che la frequenza di utilizzo dell'AI non agisce in maniera diretta né sull'impegno né sulla soddisfazione lavorativa, ma influenza tali dimensioni attraverso un meccanismo indiretto e strutturato. In particolare, un uso più frequente dell'IA si associa a una percezione di maggiore chiarezza e strutturazione nella divisione del lavoro; tale chiarezza

favorisce il riconoscimento delle competenze richieste, che a loro volta alimentano la percezione di un impatto organizzativo significativo. È proprio quest'ultima dimensione a rappresentare il passaggio decisivo che, infine, rafforza sia l'impegno verso l'organizzazione sia la soddisfazione lavorativa.

La sovrapponibilità della catena nei due modelli indica quindi che, pur trattandosi di *outcome* distinti, essi condividono un comune meccanismo sottostante mediato dalle stesse variabili organizzative. Questo rafforza l'idea che l'adozione di strumenti di Intelligenza Artificiale non produca effetti immediati sul benessere o sul coinvolgimento dei lavoratori, ma operi attraverso un processo di riorganizzazione percepita del lavoro e di valorizzazione delle competenze.

3.3.6 Analisi di moderazione

Dopo aver identificato, attraverso le analisi di mediazione, un percorso significativo che collega la frequenza di interazione con strumenti di Intelligenza Artificiale alla soddisfazione lavorativa tramite una catena di variabili organizzative; si è ritenuto opportuno approfondire se tale relazione possa variare in funzione di caratteristiche socio-demografiche dei rispondenti. In questo senso, il passaggio dall'analisi di mediazione a quella di mediazione moderata permette di verificare se tali effetti differiscono tra gruppi, come ad esempio genere o generazione. L'obiettivo è quindi comprendere se il percorso individuato agisca in maniera uniforme su tutti i lavoratori oppure se alcune categorie ne siano maggiormente influenzate.

L'analisi di mediazione moderata ha preso in esame il ruolo del genere nella catena che collega la frequenza di interazione con strumenti di Intelligenza Artificiale (*AI_Often*) alla soddisfazione lavorativa (*JobSat*), passando attraverso la divisione del lavoro (*DivLab*), la percezione del cambiamento di *expertise* (*Expert*) e l'impatto organizzativo (*OrgImp*).

Come si può osservare dalle tabelle sottostanti, i risultati mostrano che entrambe le catene indirette:

- $AI_Often \rightarrow DivLab \rightarrow OrgImp \rightarrow JobSat$ e
- $AI_Often \rightarrow DivLab \rightarrow Expert \rightarrow OrgImp \rightarrow JobSat$

risultano statisticamente significative per il gruppo maschile (*gender* = 1), mentre per il gruppo femminile (*gender* = 2) gli intervalli di confidenza includono lo zero, segnalando l'assenza di un effetto robusto. Tuttavia, l'indice di mediazione moderata non è risultato significativo, indicando che la differenza osservata tra uomini e donne non può essere considerata statisticamente solida. In altri termini, pur emergendo un effetto indiretto rilevante tra *AI_Often*

e *JobSat* per il sesso maschile, non vi sono evidenze sufficienti per affermare che il genere costituisca un moderatore sistematico della catena di mediazione.

```
***** DIRECT AND INDIRECT EFFECTS OF X ON Y *****

Direct effect of X on Y
      Effect      se      t      p      LLCI      ULCI
      ,0046      ,1085      ,0420      ,9666      -,2106      ,2197

Conditional indirect effects of X on Y:

INDIRECT EFFECT:
AI_Often  ->  DivLab  ->  JobSat

      Gender_g      Effect      BootSE      BootLLCI      BootULCI
      1,0000      ,1243      ,0682      -,0027      ,2632
      2,0000      ,0832      ,0764      -,0381      ,2628

Index of moderated mediation (difference between conditional indirect effects):
      Index      BootSE      BootLLCI      BootULCI
Gender_g      -,0411      ,0818      -,2104      ,1286

INDIRECT EFFECT:
AI_Often  ->  Expert  ->  JobSat

      Gender_g      Effect      BootSE      BootLLCI      BootULCI
      1,0000      ,0012      ,0226      -,0484      ,0498
      2,0000      -,0032      ,0263      -,0597      ,0527

Index of moderated mediation (difference between conditional indirect effects):
      Index      BootSE      BootLLCI      BootULCI
Gender_g      -,0043      ,0358      -,0823      ,0708

INDIRECT EFFECT:
AI_Often  ->  OrgImp  ->  JobSat

      Gender_g      Effect      BootSE      BootLLCI      BootULCI
      1,0000      -,0294      ,0438      -,1382      ,0342
      2,0000      ,0509      ,0408      -,0109      ,1470

Index of moderated mediation (difference between conditional indirect effects):
      Index      BootSE      BootLLCI      BootULCI
Gender_g      ,0803      ,0677      -,0095      ,2453

INDIRECT EFFECT:
AI_Often  ->  DivLab  ->  Expert  ->  JobSat

      Gender_g      Effect      BootSE      BootLLCI      BootULCI
      1,0000      -,0478      ,0484      -,1574      ,0353
      2,0000      -,0320      ,0412      -,1334      ,0325

Index of moderated mediation (difference between conditional indirect effects):
      Index      BootSE      BootLLCI      BootULCI
Gender_g      ,0158      ,0410      -,0526      ,1192

INDIRECT EFFECT:
AI_Often  ->  DivLab  ->  OrgImp  ->  JobSat

      Gender_g      Effect      BootSE      BootLLCI      BootULCI
      1,0000      ,0325      ,0275      -,0004      ,1029
      2,0000      ,0218      ,0248      -,0097      ,0879

Index of moderated mediation (difference between conditional indirect effects):
      Index      BootSE      BootLLCI      BootULCI
Gender_g      -,0107      ,0256      -,0753      ,0317

INDIRECT EFFECT:
AI_Often  ->  Expert  ->  OrgImp  ->  JobSat

      Gender_g      Effect      BootSE      BootLLCI      BootULCI
      1,0000      -,0008      ,0123      -,0275      ,0250
      2,0000      ,0020      ,0140      -,0258      ,0322

Index of moderated mediation (difference between conditional indirect effects):
      Index      BootSE      BootLLCI      BootULCI
Gender_g      ,0027      ,0188      -,0342      ,0439

INDIRECT EFFECT:
AI_Often  ->  DivLab  ->  Expert  ->  OrgImp  ->  JobSat

      Gender_g      Effect      BootSE      BootLLCI      BootULCI
      1,0000      ,0301      ,0203      ,0006      ,0764
      2,0000      ,0202      ,0202      -,0076      ,0690

Index of moderated mediation (difference between conditional indirect effects):
      Index      BootSE      BootLLCI      BootULCI
Gender_g      -,0100      ,0199      -,0534      ,0279
```

Tabella 17: Analisi di mediazione moderata con Y= *Job Satisfaction* e W= *Gender*, SPSS

Successivamente si è svolta una nuova analisi di mediazione moderata utilizzando come *outcome* il *Commitment* e mantenendo invariata la catena di mediazioni e $W = \text{Gender}$. Le tabelle riassuntive sono state inserite in Appendice E e mostrano che, anche in questo caso, il genere non modera significativamente gli effetti indiretti. In tutte le catene analizzate gli intervalli di confidenza degli indici di moderazione includono lo zero, indicando l'assenza di differenze statisticamente significative tra uomini e donne.

Dal punto di vista interpretativo, ciò significa che l'effetto dell'uso frequente dell'IA sul *Commitment* segue la stessa logica osservata per la *Job Satisfaction*, ossia attraverso catene di mediazione che coinvolgono la divisione del lavoro, l'*expertise* e la percezione di impatto organizzativo, ma senza variazioni sostanziali dovute al genere.

Oltre al genere, è stato esplorato come moderatore la generazione di appartenenza dei rispondenti. L'analisi sintetica delle due procedure di mediazione moderata indica che il percorso attraverso il quale la frequenza d'uso di strumenti di IA influenza gli esiti lavorativi è generalmente supportato dai dati, ma non risulta sistematicamente moderato dalla variabile generazionale. In particolare, alcune catene indirette emergono come rilevanti nel gruppo dei Millennials (Generazione 2), suggerendo che in questo gruppo l'uso frequente dell'IA si accompagna più chiaramente a una ridefinizione dei compiti e a un maggiore impatto organizzativo, i quali, a loro volta, influenzano positivamente la *Job Satisfaction* e il *Commitment* (vedi Appendice G e H). Tuttavia, gli indici di moderazione includono lo zero, per cui la differenza tra generazioni non è sufficientemente robusta per essere considerata significativa. Questo pattern implica che le differenze osservate tra coorti sono più descrittive che inferenziali: possono evidenziare tendenze reali, come una maggiore ricettività dei Millennials verso i benefici dell'IA, ma non superano il controllo statistico necessario a dichiarare un effetto di moderazione stabile. Tale esito può essere ricondotto sia alla dimensione campionaria relativamente contenuta e quindi alla potenza limitata dei test, sia all'omogeneità del contesto organizzativo indagato e alla parziale sovrapposizione delle percezioni intergenerazionali. In conclusione, mentre il meccanismo mediativo individuato sembra valido trasversalmente alle generazioni, la prova di una moderazione generazionale sistematica non è confermata dai dati raccolti.

È stata testata anche la posizione lavorativa inerente ai rispondenti come variabile moderatrice (vedi Appendice I e L). È interessante sottolineare come, cambiando l'*outcome*, il percorso rilevante rimanga sempre quello che collega la frequenza di utilizzo degli strumenti di intelligenza artificiale alla divisione del lavoro, quindi alla percezione di competenze e infine all'impatto organizzativo. Tale sequenza rimane il percorso più robusto attraverso il quale

l'Intelligenza Artificiale influisce sugli esiti psicologici ed organizzativi. L'introduzione della posizione lavorativa come variabile moderatrice ha messo in evidenza che, per i soggetti collocati in ruoli più elevati (Posizione 3= *Executives*) , alcuni effetti indiretti risultano positivi e talvolta significativi, soprattutto nelle catene che passano per la divisione del lavoro; tuttavia, l'indice di mediazione moderata non risulta mai significativo, né per *Commitment* né per *Job Satisfaction*. Ciò indica che le differenze osservate tra posizioni organizzative non sono statisticamente robuste e non possono essere generalizzate.

In ultima analisi, anche la variabile *Organization_AI*, che indica il tempo trascorso dall'adozione di sistemi AI in azienda , non condiziona in maniera significativa le catene di mediazione considerate. Nonostante ciò, osservando le medie degli effetti indiretti (vedi Appendice M e N) emergono alcune tendenze interessanti: le organizzazioni che hanno implementato sistemi di IA da uno a tre anni (valore 3), il percorso *AI_Often* → *Division of Labor* → *Commitment* assume un effetto indiretto positivo e più marcato rispetto ad altre categorie, sebbene non statisticamente significativo. Analogamente, la catena più complessa *AI_Often* → *Division of Labor* → *Expert* → *Commitment* tende ad assumere valori più elevati nelle organizzazioni con adozione intermedia, suggerendo che un tempo maggiore di familiarizzazione con i sistemi di IA potrebbe favorire processi più chiari di divisione del lavoro e riconoscimento delle competenze, i quali a loro volta contribuiscono all'impegno verso l'organizzazione e alla soddisfazione lavorativa. Questi risultati, pur non significativi, indicano una dinamica temporale potenzialmente rilevante: l'impatto dell'IA sui processi organizzativi non sembra essere immediato, ma può rafforzarsi man mano che l'organizzazione sviluppa pratiche consolidate e *routine* di utilizzo.

In sintesi, i risultati dimostrano che, considerato come *outcome* sia il *Commitment* che la *Job Satisfaction*, il meccanismo principale attraverso cui l'AI influenza atteggiamenti e percezioni organizzative è sempre lo stesso, mentre il ruolo moderatore della posizione appare debole e non supportato da evidenza solida.

3.4. Limiti della ricerca quantitativa

Prima di procedere alle conclusioni, è fondamentale riconoscere i limiti della presente ricerca, così da contestualizzare correttamente i risultati e tracciare possibili sviluppi futuri. In primo luogo, è importante sottolineare la dimensione e composizione del campione, infatti 105 rispondenti sono un numero discreto ma non elevato e ciò può incidere sulla potenza statistica delle analisi. Tale quadro potrebbe aver limitato le analisi di mediazione moderata, in quanto il numero relativamente piccolo di partecipanti per sottogruppo riduce la capacità di rilevare le differenze significative. Dunque, non si può affermare con certezza assoluta che non esistano differenze di genere o di generazione, ma possiamo sottolineare che non ne sono emerse di statisticamente robuste nei dati. Inoltre, il campione è relativamente omogeneo poiché tutti i dipendenti lavorano in aziende di consulenza *knowledge-based* e ciò limita la generalizzazione dei risultati ad altre realtà organizzative per cui non è noto se le dinamiche sarebbero le stesse in settori manifatturieri, nella pubblica amministrazione o in altri contesti. Ad esempio, l'effetto negativo del coordinamento algoritmico sulla soddisfazione potrebbe essere amplificato in ambienti con culture del lavoro differenti, o al contrario attenuato in imprese *tech-native*. In sintesi, la trasferibilità delle conclusioni va circoscritta: esse valgono per aziende *knowledge-intensive* con personale qualificato, ma richiedono conferma in altri contesti.

È importante sottolineare che tutte le variabili sono state rilevate tramite *self-report* dei partecipanti, il che introduce noti problemi di *bias* da risposta soggettiva. Sebbene si siano seguite accortezze, come l'ordine logico delle domande, non è escluso che parte delle associazioni trovate siano influenzate da questi *bias*. Ad esempio, un rispondente ottimista potrebbe aver dato valutazioni alte sia sull'uso dell'IA sia sulla soddisfazione, amplificando la correlazione tra le due. Infine, il modello testato, seppur articolato, è inevitabilmente una semplificazione della realtà. Ci potrebbero essere altre variabili, non considerate, che influenzano sia l'uso dell'IA sia il clima di soddisfazione lavorativa come la formazione ricevuta in ambito tecnico.

Alla luce di questi limiti, i risultati vanno interpretati con prudenza. Nonostante ciò, la ricerca offre indicazioni preziose e coerenti, che mantengono validità interna per il contesto studiato e che gettano le basi per ulteriori approfondimenti.

Discussione dei risultati finali della ricerca

La presente ricerca ha esplorato in che modo l'adozione di strumenti di Intelligenza Artificiale stia trasformando la progettazione organizzativa nelle aziende *knowledge-based*, con particolare attenzione ai meccanismi di coordinamento, alla divisione del lavoro, alle competenze richieste e agli esiti per i lavoratori. I risultati quantitativi ottenuti evidenziano un quadro organico e coerente: l'uso frequente dell'IA nel lavoro quotidiano non impatta direttamente sul benessere organizzativo dei dipendenti, ma agisce indirettamente attraverso una serie di cambiamenti strutturali e percepiti nei processi organizzativi. In altre parole, l'IA funge da agente di trasformazione del contesto di lavoro, ridefinendo il modo in cui i compiti sono coordinati e suddivisi, le competenze sono impiegate e valorizzate, e in ultima istanza come i lavoratori si rapportano alla propria organizzazione. Tali trasformazioni mediano l'effetto della nuova tecnologia sugli atteggiamenti organizzativi, confermando le ipotesi teoriche iniziali e offrendo interessanti spunti di riflessione sia pratici sia teorici.

Dall'analisi delle correlazioni tra le variabili principali emerge un *pattern* significativo e in linea con il modello proposto. La frequenza di interazione con strumenti di IA mostra associazioni positive con diverse dimensioni organizzative chiave. In particolare, un uso frequente dell'IA è moderatamente correlato ad una maggiore chiarezza e strutturazione nella divisione del lavoro ($r \approx 0,30, p < 0,01$). Questa correlazione riflette il ruolo dell'IA come nuovo meccanismo di coordinamento: sistemi algoritmici che distribuiscono automaticamente incarichi, impostano priorità e forniscono raccomandazioni rendono più esplicito chi fa cosa e in che ordine, riducendo la necessità di coordinamento diretto tradizionale attraverso riunioni e istruzioni manageriali. La forte correlazione positiva osservata tra la percezione di un coordinamento mediato dall'IA e *Division_of_Labor* ($r > 0,70$) conferma che un coordinamento efficace basato su algoritmi si associa effettivamente a una più chiara suddivisione del lavoro e ad una riorganizzazione delle attività all'interno dei *team*. In altre parole, laddove i dipendenti percepiscono che l'IA contribuisce attivamente al coordinamento – tramite l'assegnazione di compiti o la sincronizzazione delle attività - essi riportano anche una migliore definizione dei ruoli e dei compiti. Ciò rappresenta come l'IA può fungere da collante organizzativo integrando le attività dei membri all'interno dei *team*.

Parallelamente, l'interazione frequente con l'IA risulta positivamente correlata, seppur con coefficienti più contenuti, con la percezione di cambiamenti nelle competenze richieste (*Expertise*, $r \approx 0,23, p < 0,05$) e con la percezione di un impatto positivo dell'IA

sull'organizzazione (*Organizational_Impact*, $r \approx 0,22$, $p < 0,05$). Queste associazioni indicano che chi utilizza spesso strumenti di IA tende anche a riconoscere più nettamente i cambiamenti nel proprio lavoro in termini di nuove competenze da sviluppare e di influenza sui risultati organizzativi. Ad esempio, tali dipendenti possono percepire che il proprio lavoro richiede l'apprendimento di abilità legate all'interpretazione degli *output* dell'IA come il *data analysis* o lo sviluppo di capacità legate all'interazione con sistemi intelligenti. Al contempo, i lavoratori notano come essa consenta miglioramenti in termini di efficienza, qualità o rapidità nell'ottenere risultati: decisioni più informate, riduzione degli errori, maggior velocità nell'esecuzione di compiti complessi. Questo quadro iniziale di correlazioni fornisce un primo supporto alle ipotesi: l'IA non lascia invariato il contesto organizzativo, bensì si accompagna a trasformazioni positive in aspetti organizzativi specifici, i quali a loro volta sono collegati con esiti favorevoli per i lavoratori. Coerentemente, non si riscontrano correlazioni dirette significative tra la semplice frequenza d'uso dell'IA e gli indicatori di benessere lavorativo e ciò suggerisce che l'impatto della nuova tecnologia sui lavoratori opera tramite effetti indiretti, mediati dai cambiamenti organizzativi percepiti. In sintesi, queste evidenze correlazionali dipingono l'IA come catalizzatore di un processo di cambiamento: l'adozione frequente dell'IA si associa a specifiche modifiche nella struttura del lavoro che, interconnesse tra loro, preparano il terreno per miglioramenti negli atteggiamenti e nel coinvolgimento dei dipendenti.

Per verificare formalmente tale meccanismo, si è ricorso a modelli di regressione e di mediazione seriale. Le analisi di regressione multipla hanno confermato in parte il ruolo delle variabili organizzative considerate. La *Division_of_Labor* mostra un effetto positivo sulla soddisfazione (β positivo, $p \approx 0,05$), così come *Organizational_Impact* ($\beta \approx 0,23$, $p < 0,05$), mentre *AI_Coordination* presenta un coefficiente negativo significativo ($\beta \approx -0,34$, $p < 0,05$). Questo risultato, apparentemente controintuitivo, suggerisce che una percezione di coordinamento eccessivamente basato sull'IA possa ridurre la soddisfazione lavorativa. In altre parole, se i lavoratori avvertono che il controllo e l'assegnazione delle attività sono fortemente automatizzati dall'algoritmo, la loro soddisfazione tende a diminuire. Ciò potrebbe derivare da un sentimento di perdita di autonomia o di eccessiva sorveglianza: essere coordinati da una “mano invisibile” algoritmica può far sentire il dipendente meno padrone del proprio lavoro e più vincolato a istruzioni impersonali, generando potenzialmente alienazione o stress. Questo dato si collega direttamente alle riflessioni teoriche sull'equilibrio tra controllo algoritmico e autonomia umana. Come discusso nei capitoli precedenti, l'IA introduce forme di coordinamento altamente efficienti ma anche intrusive, capaci di monitorare istantaneamente

le *performance* e indirizzare i comportamenti, ad esempio attraverso raccomandazioni o valutazioni automatizzate. Se da un lato tali meccanismi aumentano l'istantaneità del controllo manageriale - si pensi alla raccolta massiva di dati sulle attività e al *feedback* immediato -, dall'altro la loro opacità e pervasività possono compromettere la percezione di autonomia e di significato del lavoro da parte dell'individuo. La letteratura recente *sull'algorithmic management* evidenzia come queste dinamiche offrono coordinamento e ottimizzazione senza precedenti, ma rischiano di imprigionare i lavoratori in una sorta di "gabbia invisibile" fatta di monitoraggio continuo e regole implicite dettate dalle macchine (Kellogg, Valentine & Christin, 2020). I nostri risultati quantitativi riflettono questo duplice volto dell'IA. La chiarezza organizzativa dovuta a compiti maggiormente definiti e una riduzione delle ambiguità sulle responsabilità, aumenta il benessere lavorativo. Al contempo, una percezione di controllo algocentrico eccessivo si associa ad un lieve calo di soddisfazione, segnalando la presenza di possibili resistenze o effetti psicologici negativi.

Le analisi di mediazione seriale hanno permesso di svelare il percorso attraverso cui l'IA influenza le variabili legate al benessere lavorativo e i risultati sono estremamente informativi. Confermando quanto già indicato dalle correlazioni, *AI_Often* non predice direttamente né il *Ccommitment* né la soddisfazione lavorativa; tuttavia, l'IA incrementa tali *outcome* attraverso una catena di mediatori ben definita. Nello specifico, si è identificato un percorso significativo comune a entrambi gli *outcome*: *AI_Often* → *Division_of_Labor* → *Expertise* → *Organizational_Impact* → *Outcome*. In termini pratici, il lavoratore che utilizza tali strumenti intelligenti tende a percepire il proprio lavoro riorganizzato in modo più efficiente e trasparente, identificando le abilità cruciali da possedere e coinvolgendolo in un processo di adeguamento delle competenze. L'effetto di queste conseguenze si manifesta attraverso la conoscenza che il suo contributo ha un impatto tangibile sui risultati dell'organizzazione e questa consapevolezza alimenta il suo attaccamento all'azienda e la soddisfazione nel ruolo. È un processo virtuoso che spiega come l'IA possa tradursi in benefici per i lavoratori solo se accompagnata da adeguati cambiamenti organizzativi. È notevole osservare che la medesima catena di mediazione risulta significativa per due esiti diversi, suggerendo l'esistenza di un meccanismo sottostante comune al benessere organizzativo dei lavoratori nell'era dell'IA. *Commitment* organizzativo e *Job Satisfaction*, pur essendo costrutti distinti, sembrano entrambi dipendere da come il lavoro viene ridisegnato e vissuto con l'introduzione dell'IA. Questo risultato si collega ai concetti di automazione vs. aumentazione discussi in letteratura. L'Intelligenza Artificiale funge da leva di "*augmented management*", ossia un supporto tecnologico che, se utilizzato

per liberare il potenziale umano- riducendo compiti ripetitivi, chiarendo obiettivi, migliorando l’allocazione delle risorse e l’uso delle competenze specialistiche-, può generare maggior coinvolgimento e appagamento nel lavoro. In linea con Davenport e Kirby (2016), il valore umano rimane centrale a condizione di evolvere verso una collaborazione consapevole con le macchine. I dati mostrano proprio questo: i lavoratori traggono beneficio dall’IA quando questa arricchisce il contesto di lavoro, creando chiarezza, promuovendo l’apprendimento e fornendo strumenti per avere successo; mentre non si rilevano benefici diretti dall’IA in assenza di tali condizioni. Tale messaggio è molto forte, infatti l’introduzione dell’IA deve essere accompagnata da interventi organizzativi mirati, altrimenti il cambiamento tecnologico rischia di rimanere sterile o addirittura di generare disagi.

Spostando l’attenzione verso l’analisi delle risposte aperte del questionario, è interessante sottolineare che i partecipanti evidenziano anche alcuni timori riguardo la perdita di controllo da parte dei lavoratori. Un rispondente ha affermato di sentirsi più motivato ma al contempo intimidito dal confronto costante con le prestazioni dell’IA, temendo di essere “oscurato” da quest’ultima e tale insicurezza genera un disagio latente sul posto di lavoro. Questo sentimento rispecchia il concetto del paradosso dell’autonomia sorvegliata: il lavoratore rimane formalmente libero di agire, ma sa di essere continuamente tracciato e misurato da sistemi intelligenti, il che può farlo sentire inadeguato o sotto pressione. D’altro canto, come emerge da diverse risposte aperte dei partecipanti, l’IA è percepita prevalentemente come un alleato. Molti riportano che grazie all’IA riescono a risparmiare tempo, ad essere più veloci e accurati, a dedicarsi a compiti di livello più elevato. Un partecipante ha sintetizzato efficacemente: *“All around we save time because of AI. Every month AI gets better and helps the flow in certain areas of the business”*, enfatizzando il miglioramento continuo e il beneficio operativo tangibile. Un altro ha osservato: *“AI has improved the efficiency of my work and freed up time for higher-level tasks”*, evidenziando il concetto di aumentazione: l’IA automatizza le mansioni più ripetitive, permettendo al professionista di concentrarsi su attività a maggior valore aggiunto. Questi riscontri qualitativi si intrecciano perfettamente con i dati quantitativi: la percezione di efficienza e risparmio di tempo riportata dai lavoratori trova riscontro nell’effetto indiretto dell’IA sulla soddisfazione attraverso l’aumento di chiarezza e impatto. Un ulteriore tema emerso è la natura complementare dell’IA: *“It’s a help, not a replacement, so far”* scrive un rispondente, mentre un altro aggiunge che l’IA *“definitely facilitates some aspects of work but we still need critical thinking when using it and can’t just take everything it gives us as 100% correct”*.

Queste affermazioni riflettono la consapevolezza che l'Intelligenza Artificiale viene vista come uno strumento di supporto e non come una sostituzione totale dell'uomo. Permane l'idea, quindi, che il giudizio critico è indispensabile per validare e contestualizzare gli output algoritmici. Tale visione è in linea con l'approccio *Human-in-the-loop*, secondo cui la collaborazione uomo-macchina è ottimale quando l'essere umano rimane coinvolto nel ciclo decisionale, supervisionando e intervenendo in modo critico. In letteratura, Davenport & Kirby (2016) sostengono un approccio che non si traduce in “uomini contro macchine”, bensì “uomini con macchine”. I partecipanti sembrano aver interiorizzato questo principio, percependo l'IA come un amplificatore delle proprie capacità professionali più che una minaccia diretta al proprio ruolo.

Nonostante la predominanza di opinioni positive, non mancano nelle risposte aperte spunti critici che meritano attenzione e trovano riscontro concettuale. Infatti, una partecipante ha espresso cautela verso l'IA, consapevole che se non progettata con criteri inclusivi può rinforzare *bias* esistenti e per tale motivo, invoca trasparenza, inclusività e *accountability* nei sistemi intelligenti adottati. Questo commento ribadisce un richiamo dalla letteratura: algoritmi non neutralizzati dai pregiudizi umani rischiano di perpetuare o amplificare disuguaglianze, minando la fiducia dei lavoratori nei risultati prodotti dall'IA come valutazioni distorte o decisioni di supporto non eque. Pertanto, integrare l'IA in modo etico e trasparente, garantendo coerenza con i valori organizzativi e sociali, diventa un imperativo categorico affinché non si vanifichino i suoi potenziali benefici. Inoltre, oltre alla già citata preoccupazione per i possibili *bias* algoritmici e al timore di essere surclassati dall'IA, un partecipante nota come “*repetitive tasks are now automated, leaving more space for complex activities*” ma aggiunge che i neoassunti faticano di più ad apprendere sul campo a causa di questa automazione. Ciò fa eco a studi come Kellogg et al. (2020), che avvertono del rischio di *deskilling* o di opportunità ridotte di apprendimento a causa della capacità dell'IA di svolgere tutte le attività di base. Tale situazione richiama anche il concetto di *moral crumple zone* (Elish, 2019) che disegna l'uomo come un supervisore passivo che perde dimestichezza operativa, pur restando formalmente responsabile degli errori. In conclusione, l'introduzione dell'IA deve essere accompagnata da iniziative che consentano al personale di comprendere i sistemi, sviluppare competenze digitali e mantenere una visione d'insieme, affinché l'esperienza non venga completamente esternalizzata alla macchina. Fortunatamente, altri studi suggeriscono che tali sfide sono superabili, ad esempio, ricerche su programmi intensivi di apprendimento mostrano che anche individui senza *background* tecnico possono acquisire competenze significative in poco tempo (Kaynak, 2023). Inoltre, concetti come il *vicarious learning* (Rostain & Huisin, 2023)

evidenziano che i lavoratori possono imparare osservando gli algoritmi in azione o collaborando con specialisti IT senza dover diventare essi stessi programmatori. Ciò significa che, con adeguate strategie di formazione continua l'adozione dell'IA può diventare occasione di crescita trasversale.

In ultima analisi, un elemento di particolare interesse emerso dallo studio empirico è la sostanziale uniformità con cui i benefici dell'IA si distribuiscono attraverso diversi sottogruppi di lavoratori. Si era ipotizzato che variabili attributive come il genere, la generazione di appartenenza, la posizione organizzativa ricoperta o il tempo di adozione di tali strumenti tecnologici all'interno dell'azienda, potessero moderare gli effetti dell'IA, privilegiando o penalizzando alcune categorie. Ad esempio, ci si chiedeva se i lavoratori più giovani mostrassero un adattamento più rapido e quindi beneficiassero maggiormente dall'IA rispetto ai colleghi più anziani, oppure se gli uomini traessero più vantaggi rispetto alle donne nell'integrare l'IA, o ancora se i dirigenti percepissero cambiamenti diversi rispetto ai funzionari operativi. Le analisi di mediazione moderata condotte hanno tuttavia evidenziato che nessuna di queste variabili categoriali influenza in modo statisticamente significativo il percorso mediativo individuato. Il modello delineato sembra valere trasversalmente alle varie categorie di dipendenti analizzate, indicando che – almeno nel contesto abbastanza omogeneo della ricerca – l'impatto dell'IA segue un meccanismo comune e condiviso. Questo è un dato incoraggiante: suggerisce che le trasformazioni positive abilitate dall'IA possono potenzialmente giovare a tutti i lavoratori, senza creare necessariamente disparità marcate tra gruppi. Tuttavia, è importante interpretare tale risultato con cautela. Da un lato, la mancanza di differenze significative potrebbe riflettere una reale trasversalità dei fenomeni in atto, ovvero l'IA, essendo una tecnologia pervasiva, ridisegna processi che coinvolgono l'intera organizzazione, e i benefici, così come le sfide, investono indistintamente uomini e donne, giovani e meno giovani, manager e subalterni, purché tutti immersi nello stesso contesto operativo. Dall'altro lato, non si può escludere che l'assenza di effetti differenziali sia dovuta ai limiti del campione. Pertanto, il messaggio che deriva è duplice: da un lato l'implementazione dell'IA, se ben gestita, può avere un impatto positivo diffuso, senza necessariamente ampliare i *gap* tra categorie di lavoratori; dall'altro, resta fondamentale monitorare e garantire che tutti i gruppi abbiano pari opportunità di adattamento e successo nell'era digitale. In particolare, le organizzazioni dovrebbero curare la formazione inclusiva e il supporto mirato laddove necessario, come *mentorship* tecnologiche per chi ha più difficoltà iniziali, in modo che nessuno rimanga indietro nella trasformazione.

Conclusioni

I risultati di questo studio confermano che l'Intelligenza Artificiale, nelle organizzazioni *knowledge-based* contemporanee, si configura come un fattore abilitante del cambiamento organizzativo. Tale tecnologia agisce sui meccanismi di coordinamento rendendoli algoritmici e influisce sulla struttura del lavoro tramite l'automazione di alcune attività e mediante nuove competenze. Tale scenario si ripercuote sugli individui in termini di percezioni e atteggiamenti. I benefici non sono automatici né garantiti, infatti dipendono criticamente da come l'Intelligenza Artificiale viene adottata e integrata. Quando le condizioni sono favorevoli tramite una maggior chiarezza dei ruoli, la formazione del personale e la trasparenza degli algoritmi, i sistemi di intelligenza generativa possono influenzare positivamente i lavoratori al fine di essere coinvolti e allineati con gli obiettivi aziendali. Allo stesso tempo, emergono sfide che le organizzazioni devono affrontare: evitare che il coordinamento algoritmico diventi disumanizzante, preservare spazi di autonomia e creatività, assicurare equità e inclusione algoritmica e prevenire la perdita di competenze critiche nel lungo periodo.

In secondo luogo, il lavoro offre un contributo pratico-manageriale importante poiché i risultati fungono da linee guida per le imprese che adottano, o intendono adottare, soluzioni di IA. Dai *pattern* emersi, i dirigenti e i *manager* possono trarre indicazioni concrete:

- l'Intelligenza Artificiale va accompagnata da una riallocazione chiara dei compiti e da una comunicazione trasparente su nuovi ruoli e responsabilità. Solo in tale modo i dipendenti ne capiscono il senso e vivono questa trasformazione come miglioramento organizzativo e non come caos o minaccia;
- è fondamentale investire in formazione del personale per valorizzare le competenze, in modo che i lavoratori possano essere capaci di interagire con l'IA e accrescere la propria professionalità. È opportuno formare il personale all'uso dei dati sviluppati dall'intelligenza generativa, ad interpretare gli *output* degli algoritmi ed a sviluppare *soft skills* complementari alla tecnologia;

- è opportuno misurare i risultati positivi derivanti dall'IA come l'efficienza, la qualità, i nuovi servizi, condividendoli con i dipendenti in modo che essi percepiscano l'impatto organizzativo del loro lavoro aumentato dall'IA in modo da motivarli e renderli partecipi del successo;
- sul versante opposto, occorre prevenire gli effetti collaterali negativi. È importante stabilire politiche chiare per evitare un uso dell'Intelligenza Artificiale che violi la *privacy* o crei un clima di sorveglianza opprimente. Il controllo deve essere mantenuto dall'essere umano, il quale possiede sempre margine decisionale, anche in presenza di raccomandazioni algoritmiche. Inoltre, è fondamentale garantire trasparenza, quindi spiegare in termini comprensibili come funzionano i sistemi e le motivazioni che determinano l'automatizzazione delle attività.

Queste implicazioni sono direttamente derivate dall'interpretazione contestuale dei risultati quantitativi e qualitativi della tesi e rappresentano un valore applicativo aggiunto: possono orientare la progettazione di interventi di *change management* nelle fasi di *digital transformation*, promuovendo l'accettazione e l'efficacia dell'Intelligenza Artificiale in azienda.

Appendice A

Analisi descrittive delle variabili prese in esame

Descrittive

			Statistica	Errore std.
AI_Coordination	Medio		2,766	,1042
	95% di intervallo di confidenza per la media	Limite inferiore	2,559	
		Limite superiore	2,972	
	Media ritagliata al 5%		2,743	
	Mediana		2,800	
	Varianza		1,140	
	Deviazione std.		1,0675	
	Minimo		1,0	
	Massimo		5,0	
	Intervallo		4,0	
	Intervallo interquartile		1,5	
	Asimmetria		,224	,236
	Curtosi		-,683	,467
Division_of_Labor	Medio		3,027	,0920
	95% di intervallo di confidenza per la media	Limite inferiore	2,844	
		Limite superiore	3,209	
	Media ritagliata al 5%		3,028	
	Mediana		3,000	
	Varianza		,889	
	Deviazione std.		,9426	
	Minimo		1,0	
	Massimo		5,0	
	Intervallo		4,0	
	Intervallo interquartile		1,4	
	Asimmetria		-,064	,236
	Curtosi		-,386	,467
Expertise	Medio		3,086	,0759
	95% di intervallo di confidenza per la media	Limite inferiore	2,935	
		Limite superiore	3,236	
	Media ritagliata al 5%		3,102	
	Mediana		3,200	
	Varianza		,605	
	Deviazione std.		,7775	
	Minimo		1,0	
	Massimo		5,0	
	Intervallo		4,0	
	Intervallo interquartile		1,0	
	Asimmetria		-,318	,236
	Curtosi		,123	,467
Organizational_Impact	Medio		3,0857	,07372
	95% di intervallo di confidenza per la media	Limite inferiore	2,9395	
		Limite superiore	3,2319	
	Media ritagliata al 5%		3,0847	
	Mediana		3,0000	
	Varianza		,571	
	Deviazione std.		,75545	
	Minimo		1,25	
	Massimo		5,00	
	Intervallo		3,75	
	Intervallo interquartile		1,00	
	Asimmetria		-,014	,236
	Curtosi		,303	,467

Commitment	Medio		2,81	,117
	95% di intervallo di confidenza per la media	Limite inferiore	2,58	
		Limite superiore	3,04	
	Media ritagliata al 5%		2,79	
	Mediana		3,00	
	Varianza		1,444	
	Deviazione std.		1,202	
	Minimo		1	
	Massimo		5	
	Intervallo		4	
	Intervallo interquartile		2	
	Asimmetria		,105	,236
	Curtosi		-,904	,467
Work_Identification	Medio		3,86	,099
	95% di intervallo di confidenza per la media	Limite inferiore	3,66	
		Limite superiore	4,05	
	Media ritagliata al 5%		3,92	
	Mediana		4,00	
	Varianza		1,027	
	Deviazione std.		1,014	
	Minimo		1	
	Massimo		5	
	Intervallo		4	
	Intervallo interquartile		2	
	Asimmetria		-,724	,236
	Curtosi		-,040	,467
Job_Satisfaction	Medio		3,99	,095
	95% di intervallo di confidenza per la media	Limite inferiore	3,80	
		Limite superiore	4,18	
	Media ritagliata al 5%		4,06	
	Mediana		4,00	
	Varianza		,952	
	Deviazione std.		,976	
	Minimo		1	
	Massimo		5	
	Intervallo		4	
	Intervallo interquartile		2	
	Asimmetria		-,804	,236
	Curtosi		,015	,467

Test di normalità

	Kolmogorov-Smirnov ^a			Shapiro-Wilk		
	Statistica	gl	Sign.	Statistica	gl	Sign.
Al_Coordination	,078	105	,128	,971	105	,021
Division_of_Labor	,072	105	,200*	,984	105	,234
Expertise	,095	105	,021	,982	105	,158
Organizational_Impact	,099	105	,013	,978	105	,082
Commitment	,169	105	<,001	,913	105	<,001
Work_Identification	,251	105	<,001	,860	105	<,001
Job_Satisfaction	,247	105	<,001	,839	105	<,001

*. Questo è un limite inferiore della significatività effettiva.

a. Correzione di significatività di Lilliefors

Appendice B

Analisi di mediazione Y= *Commitment*

```

Model: 6
Y: Commit
X: AI_Often
M1: AICoord
M2: DivLab
M3: Expert
M4: OrgImp

Sample
Size: 105

*****

OUTCOME VARIABLE:
AICoord

Model Summary
      R      R-sq      MSE      F      df1      df2      p
,1482    ,0220    1,1254    2,3125    1,0000    103,0000    ,1314

Model
      coeff      se      t      p      LLCI      ULCI
constant    2,0227    ,4995    4,0495    ,0001    1,0320    3,0133
AI_Often    ,1814    ,1193    1,5207    ,1314    -,0552    ,4181

*****

OUTCOME VARIABLE:
DivLab

Model Summary
      R      R-sq      MSE      F      df1      df2      p
,7470    ,5579    ,4005    64,3693    2,0000    102,0000    ,0000

Model
      coeff      se      t      p      LLCI      ULCI
constant    ,4590    ,3208    1,4308    ,1556    -,1773    1,0953
AI_Often    ,2148    ,0720    2,9845    ,0036    ,0720    ,3576
AICoord    ,6103    ,0588    10,3837    ,0000    ,4938    ,7269

*****

OUTCOME VARIABLE:
Expert

Model Summary
      R      R-sq      MSE      F      df1      df2      p
,7510    ,5641    ,2713    43,5632    3,0000    101,0000    ,0000

Model
      coeff      se      t      p      LLCI      ULCI
constant    1,1029    ,2667    4,1352    ,0001    ,5738    1,6320
AI_Often    ,0332    ,0618    ,5380    ,5918    -,0893    ,1558
AICoord    ,1947    ,0694    2,8055    ,0060    ,0570    ,3323
DivLab    ,4323    ,0815    5,3035    ,0000    ,2706    ,5939

*****

OUTCOME VARIABLE:
OrgImp

Model Summary
      R      R-sq      MSE      F      df1      df2      p
,6566    ,4311    ,3377    18,9413    4,0000    100,0000    ,0000

Model
      coeff      se      t      p      LLCI      ULCI
constant    1,0577    ,3217    3,2874    ,0014    ,4194    1,6960
AI_Often    ,0654    ,0690    ,9479    ,3455    -,0715    ,2023
AICoord    ,2024    ,0804    2,5182    ,0134    ,0429    ,3619
DivLab    ,0929    ,1028    ,9039    ,3682    -,1110    ,2969
Expert    ,2979    ,1110    2,6834    ,0085    ,0776    ,5181

*****

OUTCOME VARIABLE:
Commit

Model Summary
      R      R-sq      MSE      F      df1      df2      p
,7288    ,5312    ,7112    22,4332    5,0000    99,0000    ,0000

Model
      coeff      se      t      p      LLCI      ULCI
constant    -,3942    ,4915    -,8021    ,4244    -1,3695    ,5810
AI_Often    -,0777    ,1006    -,7727    ,4415    -,2774    ,1219
AICoord    ,3408    ,1203    2,8329    ,0056    ,1021    ,5794
DivLab    ,2014    ,1498    1,3447    ,1818    -,0958    ,4987
Expert    ,1318    ,1668    ,7903    ,4313    -,1991    ,4628
OrgImp    ,5066    ,1451    3,4907    ,0007    ,2186    ,7946

```

Appendice C

Analisi di mediazione $Y=Job\ Satisfaction$

```

Model: 6
Y: JobSat
X: AI_Often
M1: AICoord
M2: DivLab
M3: Expert
M4: OrgImp

Sample
Size: 105

*****

OUTCOME VARIABLE:
AICoord

Model Summary
      R      R-sq      MSE      F      df1      df2      p
    ,1482    ,0220    1,1254    2,3125    1,0000    103,0000    ,1314

Model
      coeff      se      t      p      LLCI      ULCI
constant    2,0227    ,4995    4,0495    ,0001    1,0320    3,0133
AI_Often    ,1814    ,1193    1,5207    ,1314    -,0552    ,4181

*****

OUTCOME VARIABLE:
DivLab

Model Summary
      R      R-sq      MSE      F      df1      df2      p
    ,7470    ,5579    ,4005    64,3693    2,0000    102,0000    ,0000

Model
      coeff      se      t      p      LLCI      ULCI
constant    ,4590    ,3208    1,4308    ,1556    -,1773    1,0953
AI_Often    ,2148    ,0720    2,9845    ,0036    ,0720    ,3576
AICoord     ,6103    ,0588    10,3837    ,0000    ,4938    ,7269

*****

OUTCOME VARIABLE:
Expert

Model Summary
      R      R-sq      MSE      F      df1      df2      p
    ,7510    ,5641    ,2713    43,5632    3,0000    101,0000    ,0000

Model
      coeff      se      t      p      LLCI      ULCI
constant    1,1029    ,2667    4,1352    ,0001    ,5738    1,6320
AI_Often    ,0332    ,0618    ,5380    ,5918    -,0893    ,1558
AICoord     ,1947    ,0694    2,8055    ,0060    ,0570    ,3323
DivLab      ,4323    ,0815    5,3035    ,0000    ,2706    ,5939

*****

OUTCOME VARIABLE:
OrgImp

Model Summary
      R      R-sq      MSE      F      df1      df2      p
    ,6566    ,4311    ,3377    18,9413    4,0000    100,0000    ,0000

Model
      coeff      se      t      p      LLCI      ULCI
constant    1,0577    ,3217    3,2874    ,0014    ,4194    1,6960
AI_Often    ,0654    ,0690    ,9479    ,3455    -,0715    ,2023
AICoord     ,2024    ,0804    2,5182    ,0134    ,0429    ,3619
DivLab      ,0929    ,1028    ,9039    ,3682    -,1110    ,2969
Expert      ,2979    ,1110    2,6834    ,0085    ,0776    ,5181

*****

OUTCOME VARIABLE:
JobSat

Model Summary
      R      R-sq      MSE      F      df1      df2      p
    ,4035    ,1628    ,8371    3,8498    5,0000    99,0000    ,0031

Model
      coeff      se      t      p      LLCI      ULCI
constant    2,5926    ,5332    4,8621    ,0000    1,5346    3,6506
AI_Often    -,0135    ,1091    -,1235    ,9020    -,2301    ,2031
AICoord     -,1615    ,1305    -1,2375    ,2188    -,4204    ,0974
DivLab      ,3924    ,1625    2,4141    ,0176    ,0699    ,7148
Expert      -,1498    ,1810    -,8280    ,4097    -,5089    ,2092
OrgImp      ,3806    ,1574    2,4175    ,0175    ,0682    ,6930

```

Appendice D

Analisi di moderazione $W=Gender$ e $Y=Job Satisfaction$

Model: 84							
Y: JobSat							
X: AI_Often							
M1: DivLab							
M2: Expert							
M3: OrgImp							
W: Gender_g							
Sample Size: 105							

OUTCOME VARIABLE: DivLab							
Model Summary							
	R	R-sq	MSE	F	df1	df2	p
	,3383	,1145	,8102	4,3522	3,0000	101,0000	,0063
Model							
	coeff	se	t	p	LLCI	ULCI	
constant	,3756	1,3833	,2716	,7865	-2,3684	3,1197	
AI_Often	,5465	,3263	1,6750	,0970	-,1007	1,1938	
Gender_g	,8232	,8537	,9642	,3372	-,8703	2,5167	
Int_1	-,1357	,2036	-,6668	,5064	-,5396	,2681	
Product terms key:							
Int_1 : AI_Often x Gender_g							
Test(s) of highest order unconditional interaction(s):							
	R2-chng	F	df1	df2	p		
X*W	,0039	,4446	1,0000	101,0000	,5064		

OUTCOME VARIABLE: Expert							
Model Summary							
	R	R-sq	MSE	F	df1	df2	p
	,7319	,5357	,2919	28,8452	4,0000	100,0000	,0000
Model							
	coeff	se	t	p	LLCI	ULCI	
constant	1,5468	,8306	1,8622	,0655	-,1011	3,1947	
AI_Often	-,0289	,1986	-,1454	,8847	-,4228	,3651	
DivLab	,6064	,0597	10,1538	,0000	,4880	,7249	
Gender_g	-,2092	,5148	-,4064	,6853	-1,2305	,8121	
Int_1	,0226	,1225	,1849	,8537	-,2203	,2656	
Product terms key:							
Int_1 : AI_Often x Gender_g							
Test(s) of highest order unconditional interaction(s):							
	R2-chng	F	df1	df2	p		
X*W	,0002	,0342	1,0000	100,0000	,8537		

OUTCOME VARIABLE: OrgImp							
Model Summary							
	R	R-sq	MSE	F	df1	df2	p
	,6448	,4157	,3503	14,0893	5,0000	99,0000	,0000
Model							
	coeff	se	t	p	LLCI	ULCI	
constant	2,7085	,9255	2,9264	,0043	,8720	4,5450	
AI_Often	-,3295	,2175	-1,5146	,1331	-,7611	,1022	
DivLab	,2375	,0932	2,5470	,0124	,0525	,4225	
Expert	,3631	,1095	3,3144	,0013	,1457	,5804	
Gender_g	-1,0554	,5644	-1,8700	,0644	-2,1753	,0645	
Int_1	,2412	,1342	1,7978	,0753	-,0250	,5075	
Product terms key:							
Int_1 : AI_Often x Gender_g							
Test(s) of highest order unconditional interaction(s):							
	R2-chng	F	df1	df2	p		
X*W	,0191	3,2321	1,0000	99,0000	,0753		

OUTCOME VARIABLE: JobSat							
Model Summary							
	R	R-sq	MSE	F	df1	df2	p
	,3871	,1498	,8416	4,4060	4,0000	100,0000	,0025
Model							
	coeff	se	t	p	LLCI	ULCI	
constant	2,6210	,5341	4,9070	,0000	1,5613	3,6808	
AI_Often	,0046	,1085	,0420	,9666	-,2106	,2197	
DivLab	,3025	,1458	2,0749	,0406	,0133	,5918	
Expert	-,1920	,1782	-1,0778	,2837	-,5456	,1615	
OrgImp	,3330	,1531	2,1755	,0319	,0293	,6368	

Appendice E

Analisi di moderazione $W=Gender$ e $Y=Commitment$

```
***** DIRECT AND INDIRECT EFFECTS OF X ON Y *****

Direct effect of X on Y
      Effect      se      t      p      LLCI      ULCI
      -,1158      ,1031     -1,1227     ,2643     -,3204     ,0888

Conditional indirect effects of X on Y:

INDIRECT EFFECT:
AI_Often  ->  DivLab  ->  Commit

      Gender_g      Effect      BootSE      BootLLCI      BootULCI
      1,0000      ,1606      ,0870      ,0136      ,3526
      2,0000      ,1075      ,0917      -,0559      ,3055

Index of moderated mediation (difference between conditional indirect effects):
      Index      BootSE      BootLLCI      BootULCI
Gender_g      -,0531      ,1095      -,3132      ,1426

INDIRECT EFFECT:
AI_Often  ->  Expert  ->  Commit

      Gender_g      Effect      BootSE      BootLLCI      BootULCI
      1,0000      -,0014      ,0241      -,0558      ,0497
      2,0000      ,0036      ,0285      -,0568      ,0661

Index of moderated mediation (difference between conditional indirect effects):
      Index      BootSE      BootLLCI      BootULCI
Gender_g      ,0050      ,0382      -,0759      ,0886

INDIRECT EFFECT:
AI_Often  ->  OrgImp  ->  Commit

      Gender_g      Effect      BootSE      BootLLCI      BootULCI
      1,0000      -,0536      ,0686      -,1845      ,0960
      2,0000      ,0929      ,0609      -,0160      ,2219

Index of moderated mediation (difference between conditional indirect effects):
      Index      BootSE      BootLLCI      BootULCI
Gender_g      ,1464      ,0922      -,0261      ,3358

INDIRECT EFFECT:
AI_Often  ->  DivLab  ->  Expert  ->  Commit

      Gender_g      Effect      BootSE      BootLLCI      BootULCI
      1,0000      ,0550      ,0470      -,0255      ,1586
      2,0000      ,0368      ,0396      -,0301      ,1283

Index of moderated mediation (difference between conditional indirect effects):
      Index      BootSE      BootLLCI      BootULCI
Gender_g      -,0182      ,0431      -,1287      ,0509

INDIRECT EFFECT:
AI_Often  ->  DivLab  ->  OrgImp  ->  Commit

      Gender_g      Effect      BootSE      BootLLCI      BootULCI
      1,0000      ,0592      ,0359      ,0092      ,1480
      2,0000      ,0397      ,0362      -,0182      ,1240

Index of moderated mediation (difference between conditional indirect effects):
      Index      BootSE      BootLLCI      BootULCI
Gender_g      -,0196      ,0413      -,1183      ,0550

INDIRECT EFFECT:
AI_Often  ->  Expert  ->  OrgImp  ->  Commit

      Gender_g      Effect      BootSE      BootLLCI      BootULCI
      1,0000      -,0014      ,0215      -,0522      ,0375
      2,0000      ,0036      ,0244      -,0446      ,0567

Index of moderated mediation (difference between conditional indirect effects):
      Index      BootSE      BootLLCI      BootULCI
Gender_g      ,0050      ,0332      -,0569      ,0826

INDIRECT EFFECT:
AI_Often  ->  DivLab  ->  Expert  ->  OrgImp  ->  Commit

      Gender_g      Effect      BootSE      BootLLCI      BootULCI
      1,0000      ,0549      ,0290      ,0092      ,1213
      2,0000      ,0368      ,0315      -,0134      ,1086

Index of moderated mediation (difference between conditional indirect effects):
      Index      BootSE      BootLLCI      BootULCI
Gender_g      -,0181      ,0340      -,0900      ,0508
```

Appendice F

Analisi di moderazione con W= *Generation* e Y= *Job Satisfaction*

***** DIRECT AND INDIRECT EFFECTS OF X ON Y *****

Effect	se	t	p	LLCI	ULCI
,0046	,1085	,0420	,9666	-,2106	,2197

Conditional indirect effects of X on Y:

INDIRECT EFFECT:

AI_Often -> DivLab -> JobSat

Group_ge	Effect	BootSE	BootLLCI	BootULCI
2,0000	,1066	,0622	-,0007	,2413
2,0000	,1066	,0622	-,0007	,2413
3,0000	,0843	,0693	-,0251	,2470

Index of moderated mediation:

Group_ge	Index	BootSE	BootLLCI	BootULCI
	-,0223	,0600	-,1456	,1033

INDIRECT EFFECT:

AI_Often -> Expert -> JobSat

Group_ge	Effect	BootSE	BootLLCI	BootULCI
2,0000	-,0090	,0200	-,0530	,0307
2,0000	-,0090	,0200	-,0530	,0307
3,0000	,0050	,0202	-,0255	,0601

Index of moderated mediation:

Group_ge	Index	BootSE	BootLLCI	BootULCI
	,0140	,0232	-,0216	,0721

INDIRECT EFFECT:

AI_Often -> OrgImp -> JobSat

Group_ge	Effect	BootSE	BootLLCI	BootULCI
2,0000	,0090	,0280	-,0555	,0608
2,0000	,0090	,0280	-,0555	,0608
3,0000	,0245	,0323	-,0510	,0814

Index of moderated mediation:

Group_ge	Index	BootSE	BootLLCI	BootULCI
	,0155	,0299	-,0448	,0778

INDIRECT EFFECT:

AI_Often -> DivLab -> Expert -> JobSat

Group_ge	Effect	BootSE	BootLLCI	BootULCI
2,0000	-,0403	,0392	-,1273	,0282
2,0000	-,0403	,0392	-,1273	,0282
3,0000	-,0319	,0425	-,1381	,0297

Index of moderated mediation:

Group_ge	Index	BootSE	BootLLCI	BootULCI
	,0084	,0273	-,0555	,0638

INDIRECT EFFECT:

AI_Often -> DivLab -> OrgImp -> JobSat

Group_ge	Effect	BootSE	BootLLCI	BootULCI
2,0000	,0260	,0209	,0002	,0811
2,0000	,0260	,0209	,0002	,0811
3,0000	,0205	,0258	-,0034	,0947

Index of moderated mediation:

Group_ge	Index	BootSE	BootLLCI	BootULCI
	-,0054	,0166	-,0347	,0363

INDIRECT EFFECT:

AI_Often -> Expert -> OrgImp -> JobSat

Group_ge	Effect	BootSE	BootLLCI	BootULCI
2,0000	,0059	,0113	-,0152	,0330
2,0000	,0059	,0113	-,0152	,0330
3,0000	-,0032	,0104	-,0296	,0138

Index of moderated mediation:

Group_ge	Index	BootSE	BootLLCI	BootULCI
	-,0091	,0118	-,0383	,0078

INDIRECT EFFECT:

AI_Often -> DivLab -> Expert -> OrgImp -> JobSat

Group_ge	Effect	BootSE	BootLLCI	BootULCI
2,0000	,0262	,0178	,0005	,0692
2,0000	,0262	,0178	,0005	,0692
3,0000	,0207	,0207	-,0035	,0751

Index of moderated mediation:

Group_ge	Index	BootSE	BootLLCI	BootULCI
	-,0055	,0146	-,0323	,0293

Appendice G

Analisi di moderazione con W= *Generation* e Y= *Commitment*

```

***** DIRECT AND INDIRECT EFFECTS OF X ON Y *****

Direct effect of X on Y
Effect      se      t      p      LLCI      ULCI
-,1158      ,1031    -1,1227    ,2643    -,3204      ,0888

Conditional indirect effects of X on Y:

INDIRECT EFFECT:
AI_Often    ->    DivLab    ->    Commit

Group_ge    Effect    BootSE    BootLLCI    BootULCI
2,0000      ,1378      ,0777      ,0098      ,3098
2,0000      ,1378      ,0777      ,0098      ,3098
3,0000      ,1090      ,0876      -,0219      ,3208

Index of moderated mediation:
Index      BootSE    BootLLCI    BootULCI
Group_ge    -,0288      ,0772      -,1792      ,1324

INDIRECT EFFECT:
AI_Often    ->    Expert    ->    Commit

Group_ge    Effect    BootSE    BootLLCI    BootULCI
2,0000      ,0104      ,0230      -,0361      ,0603
2,0000      ,0104      ,0230      -,0361      ,0603
3,0000      -,0057      ,0207      -,0616      ,0256

Index of moderated mediation:
Index      BootSE    BootLLCI    BootULCI
Group_ge    -,0161      ,0233      -,0736      ,0191

INDIRECT EFFECT:
AI_Often    ->    OrgImp    ->    Commit

Group_ge    Effect    BootSE    BootLLCI    BootULCI
2,0000      ,0164      ,0527      -,0824      ,1342
2,0000      ,0164      ,0527      -,0824      ,1342
3,0000      ,0446      ,0572      -,0805      ,1462

Index of moderated mediation:
Index      BootSE    BootLLCI    BootULCI
Group_ge    ,0282      ,0509      -,0934      ,1127

INDIRECT EFFECT:
AI_Often    ->    DivLab    ->    Expert    ->    Commit

Group_ge    Effect    BootSE    BootLLCI    BootULCI
2,0000      ,0464      ,0392      -,0205      ,1338
2,0000      ,0464      ,0392      -,0205      ,1338
3,0000      ,0367      ,0417      -,0221      ,1423

Index of moderated mediation:
Index      BootSE    BootLLCI    BootULCI
Group_ge    -,0097      ,0294      -,0712      ,0564

INDIRECT EFFECT:
AI_Often    ->    DivLab    ->    OrgImp    ->    Commit

Group_ge    Effect    BootSE    BootLLCI    BootULCI
2,0000      ,0473      ,0297      ,0057      ,1208
2,0000      ,0473      ,0297      ,0057      ,1208
3,0000      ,0374      ,0392      -,0052      ,1424

Index of moderated mediation:
Index      BootSE    BootLLCI    BootULCI
Group_ge    -,0099      ,0282      -,0583      ,0561

INDIRECT EFFECT:
AI_Often    ->    Expert    ->    OrgImp    ->    Commit

Group_ge    Effect    BootSE    BootLLCI    BootULCI
2,0000      ,0107      ,0196      -,0264      ,0539
2,0000      ,0107      ,0196      -,0264      ,0539
3,0000      -,0059      ,0182      -,0499      ,0244

Index of moderated mediation:
Index      BootSE    BootLLCI    BootULCI
Group_ge    -,0166      ,0191      -,0611      ,0136

INDIRECT EFFECT:
AI_Often    ->    DivLab    ->    Expert    ->    OrgImp    ->    Commit

Group_ge    Effect    BootSE    BootLLCI    BootULCI
2,0000      ,0478      ,0259      ,0064      ,1069
2,0000      ,0478      ,0259      ,0064      ,1069
3,0000      ,0378      ,0319      -,0049      ,1149

Index of moderated mediation:
Index      BootSE    BootLLCI    BootULCI
Group_ge    -,0100      ,0252      -,0553      ,0458

```

Appendice H

Analisi di moderazione con W= *Generation* e Y= *Job Satisfaction*

```
***** DIRECT AND INDIRECT EFFECTS OF X ON Y *****

Direct effect of X on Y
      Effect      se      t      p      LLCI      ULCI
      ,0046      ,1085      ,0420      ,9666      -,2106      ,2197

Conditional indirect effects of X on Y:

INDIRECT EFFECT:
AI_Often  ->  DivLab      ->  JobSat

      Group_ge      Effect      BootSE      BootLLCI      BootULCI
      2,0000      ,1066      ,0622      -,0007      ,2413
      2,0000      ,1066      ,0622      -,0007      ,2413
      3,0000      ,0843      ,0693      -,0251      ,2470

      Index of moderated mediation:
      Index      BootSE      BootLLCI      BootULCI
      Group_ge      -,0223      ,0600      -,1456      ,1033

INDIRECT EFFECT:
AI_Often  ->  Expert      ->  JobSat

      Group_ge      Effect      BootSE      BootLLCI      BootULCI
      2,0000      -,0090      ,0200      -,0530      ,0307
      2,0000      -,0090      ,0200      -,0530      ,0307
      3,0000      ,0050      ,0202      -,0255      ,0601

      Index of moderated mediation:
      Index      BootSE      BootLLCI      BootULCI
      Group_ge      ,0140      ,0232      -,0216      ,0721

INDIRECT EFFECT:
AI_Often  ->  OrgImp      ->  JobSat

      Group_ge      Effect      BootSE      BootLLCI      BootULCI
      2,0000      ,0090      ,0280      -,0555      ,0608
      2,0000      ,0090      ,0280      -,0555      ,0608
      3,0000      ,0245      ,0323      -,0510      ,0814

      Index of moderated mediation:
      Index      BootSE      BootLLCI      BootULCI
      Group_ge      ,0155      ,0299      -,0448      ,0778

INDIRECT EFFECT:
AI_Often  ->  DivLab      ->  Expert      ->  JobSat

      Group_ge      Effect      BootSE      BootLLCI      BootULCI
      2,0000      -,0403      ,0392      -,1273      ,0282
      2,0000      -,0403      ,0392      -,1273      ,0282
      3,0000      -,0319      ,0425      -,1381      ,0297

      Index of moderated mediation:
      Index      BootSE      BootLLCI      BootULCI
      Group_ge      ,0084      ,0273      -,0555      ,0638

INDIRECT EFFECT:
AI_Often  ->  DivLab      ->  OrgImp      ->  JobSat

      Group_ge      Effect      BootSE      BootLLCI      BootULCI
      2,0000      ,0260      ,0209      ,0002      ,0811
      2,0000      ,0260      ,0209      ,0002      ,0811
      3,0000      ,0205      ,0258      -,0034      ,0947

      Index of moderated mediation:
      Index      BootSE      BootLLCI      BootULCI
      Group_ge      -,0054      ,0166      -,0347      ,0363

INDIRECT EFFECT:
AI_Often  ->  Expert      ->  OrgImp      ->  JobSat

      Group_ge      Effect      BootSE      BootLLCI      BootULCI
      2,0000      ,0059      ,0113      -,0152      ,0330
      2,0000      ,0059      ,0113      -,0152      ,0330
      3,0000      -,0032      ,0104      -,0296      ,0138

      Index of moderated mediation:
      Index      BootSE      BootLLCI      BootULCI
      Group_ge      -,0091      ,0118      -,0383      ,0078

INDIRECT EFFECT:
AI_Often  ->  DivLab      ->  Expert      ->  OrgImp      ->  JobSat

      Group_ge      Effect      BootSE      BootLLCI      BootULCI
      2,0000      ,0262      ,0178      ,0005      ,0692
      2,0000      ,0262      ,0178      ,0005      ,0692
      3,0000      ,0207      ,0207      -,0035      ,0751

      Index of moderated mediation:
      Index      BootSE      BootLLCI      BootULCI
      Group_ge      -,0055      ,0146      -,0323      ,0293
```

Appendice I

Analisi di moderazione con W=Posizione lavorativa e Y= *Commitment*

***** DIRECT AND INDIRECT EFFECTS OF X ON Y *****

Direct effect of X on Y

Effect	se	t	p	LLCI	ULCI
-,1158	,1031	-1,1227	,2643	-,3204	,0888

Conditional indirect effects of X on Y:

INDIRECT EFFECT:

AI_Often -> DivLab -> Commit

Position	Effect	BootSE	BootLLCI	BootULCI
2,0000	,0786	,0766	-,0630	,2488
3,0000	,1535	,0820	,0139	,3303
3,0000	,1535	,0820	,0139	,3303

Index of moderated mediation:

Position	Index	BootSE	BootLLCI	BootULCI
	,0750	,0858	-,0591	,2814

INDIRECT EFFECT:

AI_Often -> Expert -> Commit

Position	Effect	BootSE	BootLLCI	BootULCI
2,0000	,0082	,0227	-,0381	,0580
3,0000	-,0002	,0259	-,0635	,0500
3,0000	-,0002	,0259	-,0635	,0500

Index of moderated mediation:

Position	Index	BootSE	BootLLCI	BootULCI
	-,0084	,0321	-,0822	,0521

INDIRECT EFFECT:

AI_Often -> OrgImp -> Commit

Position	Effect	BootSE	BootLLCI	BootULCI
2,0000	,0684	,0539	-,0396	,1768
3,0000	,0032	,0573	-,1018	,1284
3,0000	,0032	,0573	-,1018	,1284

Index of moderated mediation:

Position	Index	BootSE	BootLLCI	BootULCI
	-,0653	,0631	-,1896	,0661

INDIRECT EFFECT:

AI_Often -> DivLab -> Expert -> Commit

Position	Effect	BootSE	BootLLCI	BootULCI
2,0000	,0266	,0324	-,0257	,1019
3,0000	,0520	,0437	-,0224	,1477
3,0000	,0520	,0437	-,0224	,1477

Index of moderated mediation:

Position	Index	BootSE	BootLLCI	BootULCI
	,0254	,0369	-,0258	,1184

INDIRECT EFFECT:

AI_Often -> DivLab -> OrgImp -> Commit

Position	Effect	BootSE	BootLLCI	BootULCI
2,0000	,0283	,0302	-,0187	,1015
3,0000	,0552	,0331	,0079	,1373
3,0000	,0552	,0331	,0079	,1373

Index of moderated mediation:

Position	Index	BootSE	BootLLCI	BootULCI
	,0270	,0305	-,0240	,1001

INDIRECT EFFECT:

AI_Often -> Expert -> OrgImp -> Commit

Position	Effect	BootSE	BootLLCI	BootULCI
2,0000	,0083	,0193	-,0330	,0472
3,0000	-,0002	,0224	-,0478	,0456
3,0000	-,0002	,0224	-,0478	,0456

Index of moderated mediation:

Position	Index	BootSE	BootLLCI	BootULCI
	-,0085	,0275	-,0617	,0496

INDIRECT EFFECT:

AI_Often -> DivLab -> Expert -> OrgImp -> Commit

Position	Effect	BootSE	BootLLCI	BootULCI
2,0000	,0269	,0272	-,0151	,0910
3,0000	,0526	,0283	,0082	,1174
3,0000	,0526	,0283	,0082	,1174

Index of moderated mediation:

Position	Index	BootSE	BootLLCI	BootULCI
	,0257	,0262	-,0224	,0829

Appendice L

Analisi di moderazione con W=Posizione lavorativa e Y= *Job Satisfaction*

```
***** DIRECT AND INDIRECT EFFECTS OF X ON Y *****

Direct effect of X on Y
      Effect      se      t      p      LLCI      ULCI
      ,0046      ,1085      ,0420      ,9666      -,2106      ,2197

Conditional indirect effects of X on Y:

INDIRECT EFFECT:
AI_Often  ->  DivLab      ->  JobSat

      Position      Effect      BootSE      BootLLCI      BootULCI
      2,0000      ,0608      ,0596      -,0478      ,1939
      3,0000      ,1188      ,0677      -,0021      ,2576
      3,0000      ,1188      ,0677      -,0021      ,2576

      Index of moderated mediation:
      Index      BootSE      BootLLCI      BootULCI
Position      ,0580      ,0691      -,0480      ,2239

INDIRECT EFFECT:
AI_Often  ->  Expert      ->  JobSat

      Position      Effect      BootSE      BootLLCI      BootULCI
      2,0000      -,0071      ,0209      -,0515      ,0365
      3,0000      ,0002      ,0220      -,0465      ,0472
      3,0000      ,0002      ,0220      -,0465      ,0472

      Index of moderated mediation:
      Index      BootSE      BootLLCI      BootULCI
Position      ,0073      ,0280      -,0545      ,0663

INDIRECT EFFECT:
AI_Often  ->  OrgImp      ->  JobSat

      Position      Effect      BootSE      BootLLCI      BootULCI
      2,0000      ,0376      ,0341      -,0249      ,1087
      3,0000      ,0017      ,0319      -,0718      ,0576
      3,0000      ,0017      ,0319      -,0718      ,0576

      Index of moderated mediation:
      Index      BootSE      BootLLCI      BootULCI
Position      -,0358      ,0421      -,1348      ,0310

INDIRECT EFFECT:
AI_Often  ->  DivLab      ->  Expert      ->  JobSat

      Position      Effect      BootSE      BootLLCI      BootULCI
      2,0000      -,0231      ,0323      -,1004      ,0300
      3,0000      -,0452      ,0443      -,1424      ,0362
      3,0000      -,0452      ,0443      -,1424      ,0362

      Index of moderated mediation:
      Index      BootSE      BootLLCI      BootULCI
Position      -,0221      ,0351      -,1097      ,0287

INDIRECT EFFECT:
AI_Often  ->  DivLab      ->  OrgImp      ->  JobSat

      Position      Effect      BootSE      BootLLCI      BootULCI
      2,0000      ,0155      ,0198      -,0108      ,0682
      3,0000      ,0303      ,0244      ,0002      ,0934
      3,0000      ,0303      ,0244      ,0002      ,0934

      Index of moderated mediation:
      Index      BootSE      BootLLCI      BootULCI
Position      ,0148      ,0202      -,0121      ,0669

INDIRECT EFFECT:
AI_Often  ->  Expert      ->  OrgImp      ->  JobSat

      Position      Effect      BootSE      BootLLCI      BootULCI
      2,0000      ,0046      ,0116      -,0189      ,0303
      3,0000      -,0001      ,0119      -,0241      ,0261
      3,0000      -,0001      ,0119      -,0241      ,0261

      Index of moderated mediation:
      Index      BootSE      BootLLCI      BootULCI
Position      -,0047      ,0152      -,0346      ,0290

INDIRECT EFFECT:
AI_Often  ->  DivLab      ->  Expert      ->  OrgImp      ->  JobSat

      Position      Effect      BootSE      BootLLCI      BootULCI
      2,0000      ,0148      ,0166      -,0082      ,0566
      3,0000      ,0289      ,0198      ,0004      ,0756
      3,0000      ,0289      ,0198      ,0004      ,0756

      Index of moderated mediation:
      Index      BootSE      BootLLCI      BootULCI
Position      ,0141      ,0161      -,0113      ,0535
```

Appendice M

Analisi di moderazione con $W = AI_Organization$ e $Y = Commitment$

***** DIRECT AND INDIRECT EFFECTS OF X ON Y *****

Effect	se	t	p	LLCI	ULCI
Direct effect of X on Y					
Effect	se	t	p	LLCI	ULCI
-,1158	,1031	-1,1227	,2643	-,3204	,0888

Conditional indirect effects of X on Y:

INDIRECT EFFECT:

AI_Often → DivLab → Commit

Organ_AI	Effect	BootSE	BootLLCI	BootULCI
2,0000	,1354	,0786	,0009	,3051
3,0000	,0985	,0738	-,0293	,2620
4,0000	,0616	,1073	-,1580	,2906

Index of moderated mediation:

Organ_AI	Index	BootSE	BootLLCI	BootULCI
Organ_AI	-,0369	,0583	-,1658	,0705

INDIRECT EFFECT:

AI_Often → Expert → Commit

Organ_AI	Effect	BootSE	BootLLCI	BootULCI
2,0000	-,0083	,0237	-,0657	,0337
3,0000	,0138	,0196	-,0158	,0634
4,0000	,0359	,0384	-,0186	,1313

Index of moderated mediation:

Organ_AI	Index	BootSE	BootLLCI	BootULCI
Organ_AI	,0221	,0252	-,0111	,0841

INDIRECT EFFECT:

AI_Often → OrgImp → Commit

Organ_AI	Effect	BootSE	BootLLCI	BootULCI
2,0000	,0341	,0512	-,0631	,1439
3,0000	,0106	,0522	-,0919	,1214
4,0000	-,0129	,0743	-,1597	,1339

Index of moderated mediation:

Organ_AI	Index	BootSE	BootLLCI	BootULCI
Organ_AI	-,0235	,0365	-,0976	,0484

INDIRECT EFFECT:

AI_Often → DivLab → Expert → Commit

Organ_AI	Effect	BootSE	BootLLCI	BootULCI
2,0000	,0469	,0414	-,0234	,1392
3,0000	,0341	,0343	-,0221	,1138
4,0000	,0213	,0428	-,0643	,1183

Index of moderated mediation:

Organ_AI	Index	BootSE	BootLLCI	BootULCI
Organ_AI	-,0128	,0244	-,0726	,0256

INDIRECT EFFECT:

AI_Often → DivLab → OrgImp → Commit

Organ_AI	Effect	BootSE	BootLLCI	BootULCI
2,0000	,0399	,0284	-,0003	,1075
3,0000	,0290	,0265	-,0070	,0946
4,0000	,0181	,0351	-,0425	,1022

Index of moderated mediation:

Organ_AI	Index	BootSE	BootLLCI	BootULCI
Organ_AI	-,0109	,0178	-,0499	,0238

INDIRECT EFFECT:

AI_Often → Expert → OrgImp → Commit

Organ_AI	Effect	BootSE	BootLLCI	BootULCI
2,0000	-,0089	,0205	-,0517	,0315
3,0000	,0149	,0167	-,0125	,0547
4,0000	,0388	,0282	-,0040	,1057

Index of moderated mediation:

Organ_AI	Index	BootSE	BootLLCI	BootULCI
Organ_AI	,0239	,0182	-,0043	,0657

INDIRECT EFFECT:

AI_Often → DivLab → Expert → OrgImp → Commit

Organ_AI	Effect	BootSE	BootLLCI	BootULCI
2,0000	,0507	,0286	,0038	,1131
3,0000	,0369	,0272	-,0084	,0988
4,0000	,0231	,0378	-,0500	,1050

Index of moderated mediation:

Organ_AI	Index	BootSE	BootLLCI	BootULCI
Organ_AI	-,0138	,0196	-,0574	,0237

Appendice N

Analisi di moderazione con $W=AI_Organization$ e $Y=Job\ Satisfaction$

***** DIRECT AND INDIRECT EFFECTS OF X ON Y *****

Direct effect of X on Y

Effect	se	t	p	LLCI	ULCI
,0046	,1085	,0420	,9666	-,2106	,2197

Conditional indirect effects of X on Y:

INDIRECT EFFECT:

AI_Often	->	DivLab	->	JobSat
Organ_AI	Effect	BootSE	BootLLCI	BootULCI
2,0000	,1048	,0610	-,0092	,2318
3,0000	,0762	,0575	-,0245	,2027
4,0000	,0477	,0802	-,1046	,2240

Index of moderated mediation:

Organ_AI	Index	BootSE	BootLLCI	BootULCI
	-,0285	,0421	-,1216	,0521

INDIRECT EFFECT:

AI_Often	->	Expert	->	JobSat
Organ_AI	Effect	BootSE	BootLLCI	BootULCI
2,0000	,0072	,0212	-,0275	,0607
3,0000	-,0120	,0187	-,0592	,0146
4,0000	-,0312	,0381	-,1285	,0195

Index of moderated mediation:

Organ_AI	Index	BootSE	BootLLCI	BootULCI
	-,0192	,0245	-,0845	,0113

INDIRECT EFFECT:

AI_Often	->	OrgImp	->	JobSat
Organ_AI	Effect	BootSE	BootLLCI	BootULCI
2,0000	,0187	,0289	-,0433	,0761
3,0000	,0058	,0295	-,0646	,0606
4,0000	-,0071	,0434	-,1079	,0705

Index of moderated mediation:

Organ_AI	Index	BootSE	BootLLCI	BootULCI
	-,0129	,0221	-,0632	,0276

INDIRECT EFFECT:

AI_Often	->	DivLab	->	Expert	->	JobSat
Organ_AI	Effect	BootSE	BootLLCI	BootULCI		
2,0000	-,0408	,0407	-,1308	,0317		
3,0000	-,0297	,0338	-,1101	,0244		
4,0000	-,0185	,0389	-,1153	,0455		

Index of moderated mediation:

Organ_AI	Index	BootSE	BootLLCI	BootULCI
	,0111	,0210	-,0277	,0589

INDIRECT EFFECT:

AI_Often	->	DivLab	->	OrgImp	->	JobSat
Organ_AI	Effect	BootSE	BootLLCI	BootULCI		
2,0000	,0219	,0195	-,0011	,0740		
3,0000	,0159	,0177	-,0030	,0642		
4,0000	,0100	,0215	-,0207	,0657		

Index of moderated mediation:

Organ_AI	Index	BootSE	BootLLCI	BootULCI
	-,0060	,0104	-,0292	,0129

INDIRECT EFFECT:

AI_Often	->	Expert	->	OrgImp	->	JobSat
Organ_AI	Effect	BootSE	BootLLCI	BootULCI		
2,0000	-,0049	,0119	-,0312	,0180		
3,0000	,0082	,0104	-,0058	,0339		
4,0000	,0213	,0191	-,0024	,0683		

Index of moderated mediation:

Organ_AI	Index	BootSE	BootLLCI	BootULCI
	,0131	,0120	-,0027	,0428

INDIRECT EFFECT:

AI_Often	->	DivLab	->	Expert	->	OrgImp	->	JobSat
Organ_AI	Effect	BootSE	BootLLCI	BootULCI				
2,0000	,0278	,0204	-,0003	,0764				
3,0000	,0202	,0179	-,0034	,0652				
4,0000	,0127	,0222	-,0246	,0682				

Index of moderated mediation:

Organ_AI	Index	BootSE	BootLLCI	BootULCI
	-,0076	,0116	-,0350	,0139

Bibliografia

- Bailey, D. E., & Barley, S. R. (2020). *Beyond design and use: How scholars should study intelligent technologies and work*. Information and Organization, 30(2), 100286.
- Barati, M., & Ansari, B. (2022). *Effects of algorithmic control on power asymmetry and inequality within organizations*. Journal of Management Control, 33(4).
- Bernstein, E., Bunch, J., Canner, N., & Lee, M. (2016). *Beyond the Holacracy Hype: The Overwrought Claims—and Actual Promise—of the Next Generation of Self-Managed Teams*. Harvard Business Review, 94(7/8)
- Brynjolfsson, E., & McAfee, A. (2014). *The Second Machine Age: Work, Progress, and Prosperity in a Time of Brilliant Technologies*. W. W. Norton & Company.
- Daugherty, P. R., & Wilson, H. J. (2018). *Human + Machine: Reimagining work in the age of AI*. Harvard Business Press.
- Davenport, T. H., & Kirby, J. (2016). *Only Humans Need Apply: Winners and Losers in the Age of Smart Machines*. Harper Business.
- Davenport, T. H., & Ronanki, R. (2018). *Artificial intelligence for the real world*. Harvard Business Review, 96(1).
- Elish, M.C. (2019). *Moral Crumple Zones: Cautionary Tales in Human-Robot Interaction*. Engaging Science, Technology, and Society, 5
- Faraj, S., Jarvenpaa, S. L., & Majchrzak, A. (2011). *Knowledge collaboration in online communities*. Organization Science, 22(5).
- Freidson, E. (1970). *Professional dominance: The social structure of medical care*. New York: Atherton Press.
- Gasser, U. & Almeida, V.A.F. (2017). *A Layered Model for AI Governance*. IEEE Internet Computing, 21(6).

- Kaynak, E. (2023). *Leveraging learning collectives: How novice outsiders break into an occupation*. *Organization Science*, 35(3).
- Lee, M. K., Kusbit, D., Metsky, E., & Dabbish, L. (2015). *Working with Machines: The Impact of Algorithmic and Data-Driven Management on Human Workers*. Proceedings of the 33rd Annual ACM Conference on Human Factors in Computing Systems (CHI '15).
- Levina, N., & Vaast, E. (2005). *The emergence of boundary spanning competence in practice: Implications for implementation and use of information systems*. *MIS Quarterly*, 29(2).
- Mintzberg, H. (1979). *The Structuring of Organizations: A Synthesis of the Research*. Englewood Cliffs, NJ: Prentice-Hall.
- Newell, S. (2015). *Managing knowledge and managing knowledge work: What we know and what the future holds*. *Journal of Information Technology*, 30(1).
- Newell, S., & Marabelli, M. (2015). *Strategic opportunities (and challenges) of algorithmic decision-making: A call for action on the long-term societal effects of 'datification'*. *Journal of Strategic Information Systems*, 24(1).
- Nicolini, D., Mengis, J., & Swan, J. (2017). *Understanding the role of objects in organizing knowledge in organizations: A practice-based view*.
- Pakarinen, M., & Huising, R. (2023). *The Relational Work of Expertise: Reconfiguring the Social Fabric of Knowledge in the Era of Artificial Intelligence*. *Organization Science*, 34(1).
- Pakarinen, J./M., & Huising, R. (2023). *Relational expertise and the challenge of AI integration in professional work*. Working paper
- Puranam, P., Alexy, O., & Reitzig, M. (2019). *What's "new" about new forms of organizing?* *Academy of Management Review*, 49(2).
- Raisch, S., & Krakowski, S. (2021). *Artificial intelligence and management: The automation–augmentation paradox in decision making*. *Academy of Management Review*, 46(1).

- Rostain, M. & Huising, R. (2023). *Vicarious Coding: Breaching Computational Opacity in the Digital Era*. Academy of Management Journal.
- Russell, S. J., & Norvig, P. (2021). *Artificial intelligence: A modern approach* (4th ed.). Pearson.
- Schad, J., Lewis, M. W., Raisch, S., & Smith, W. K. (2016). *Paradox Research in Management Science: Looking Back to Move Forward*. Academy of Management Annals, 10(1).
- Slayton, R., & Clark-Ginsberg, A. (2018). *Beyond regulatory capture: Coproducing expertise for critical infrastructure protection*. Regulation & Governance, 12(1).
- Valentine, M. A., Kellogg, K. C., & Christin, A. (2020). *Algorithms at work: The new contested terrain of control*. Academy of Management Annals, 14(1).
- Van den Broek, E., Sergeeva, A., & Huysman, M. (2022). *When the machine meets the expert: A field study of AI integration in HR decision-making*.