

MASSIMO TAMIATTI

DALL'INVASIONE ALGORITMICA ALL'**AGENTIFICAZIONE** DEL LAVORO

Governance, etica e futuro nell'era dell'AI



Frontiere
Tecnologia
Lavoro

2026

1

Che cos'è? Si tratta di un sito di divulgazione scientifica, che studia l'impatto delle nuove tecnologie sul mondo del lavoro. I cambiamenti oggi sono molto veloci ed esponenziali. Le professioni non sono un qualcosa di statico ma di molto dinamico, quindi le evoluzioni in atto possono estinguere, trasformare o possono farne nascere di completamente nuove. Quali competenze saranno necessarie per affrontare tutto questo? Upskilling e Reskilling, ma non solo.

Perché? Dopo 25 anni di ricerche ho capito che per andare oltre i confini i numeri non bastano, gli approcci tradizionali sono obsoleti e legati al vecchio modo di fare ricerca. Oggi bisogna studiare i segnali deboli che sono sotto i nostri occhi e che guardiamo ma non vediamo. Soffermarsi su questi segnali e immaginare più futuri possibili sarà il compito di Frontiere. Le fonti tradizionali rispondono in tempo reale, ma fotografano solo il presente. Oggi però c'è la necessità di fare previsioni almeno a cinque anni, senza trascurare il passato e senza dimenticare che su ogni tema affrontato c'è una storia da considerare.

A chi serve? Troppo facile fare previsioni a lungo termine che non si possono verificare. Sono convinto che al mondo della ricerca serva un confronto immediato sia con il mondo della scuola e dell'orientamento e sia con i giovani, che rischiano di studiare per professioni che in pochi anni non ci saranno più. Saranno qui per scambiare opinioni con chi è indeciso, con chi è disorientato e con chi ha intenzione di ripartire. Ai giovani dico che nel mondo del lavoro esiste un mondo di sopra e un mondo di sotto. Un mondo di chi programma gli algoritmi(pochi e ben pagati) e di chi lavora per gli algoritmi(molti e mal pagati) e bisogna capire cosa fare per uscire da un mercato del lavoro che rischia di condannare i giovani a stipendi di 1.000 euro al mese per tutta la vita.

Com'è? Ogni tema viene trattato con pubblicazioni leggere e con una particolare attenzione all'aspetto della profondità temporale. Le pubblicazioni saranno accompagnate da brevi video(short), che vanno subito al punto. Frontieretecnologialavoro.com si posiziona come uno spazio di informazione, orientamento e confronto proprio su queste tematiche.



Linkedin: <https://it.linkedin.com/in/massimo-tamiatti-77605168>

Facebook: <https://www.facebook.com/profile.php?id=100006516969397>

Telegram: <https://t.me/frontieretecnologialavoro>

Instagram: <https://www.instagram.com/tamiattimassimo/>

PRESENTAZIONE

Dopo l'esplosione dell'AI generativa nel 2022-2023, una nuova frontiera tecnologica sta emergendo: gli agenti AI autonomi. A differenza dei chatbot che rispondono a richieste singole, questi sistemi percepiscono l'ambiente, pianificano strategie complesse e agiscono nel mondo reale senza supervisione costante. Questo saggio esplora tre dimensioni critiche dell'agentificazione del lavoro: (1) le implicazioni per la governance e la responsabilità giuridica in contesti dove le decisioni sono delegate a sistemi probabilistici; (2) la trasformazione radicale delle competenze richieste, dalla logica del "fare" alla logica dell'"essere"; (3) tre scenari plausibili per il mercato del lavoro italiano al 2030, tra opportunità di potenziamento umano e rischi di polarizzazione sociale.

Indice

Da assistenti a costruttori cosa cambia	pag. 6
Il ruolo della governance e del controllo umano	pag. 12
Dove vengono usati gli Agenti?	pag. 17
Quali sono i problemi irrisolti?	pag. 21
Chi serve per governare questi problemi?	pag. 22
Ripensare le competenze: dalla logica del fare alla logica dell'essere	pag. 25
Il ritorno dei filosofi: perché le competenze umanistiche sono la nuova ingegneria?	pag. 32
Il bivio del 2030: tre scenari per un futuro "agentificato"	pag. 34
Conclusioni	pag. 36

DALL'INVASIONE ALGORITMICA ALL'AGENTIFICAZIONE DEL MERCATO DEL LAVORO

Negli ultimi anni si è parlato tantissimo di Intelligenza Artificiale, ma oggi il nuovo grande tema sono gli agenti AI. Se ne discute ovunque: c'è "hype", entusiasmo, timori, promesse di rivoluzioni nel lavoro e nella produttività. Ma prima di lasciarci trasportare dall'entusiasmo, vale la pena fermarsi e capire meglio che cosa sono davvero questi agenti, cosa possono fare oggi e quali sono i loro limiti reali.

Da assistenti a costruttori: che cosa cambia?

Fino a poco tempo fa, quando pensavamo all'Intelligenza Artificiale nel lavoro, immaginavamo assistenti con cui parlare: ChatGPT che scrive la mail al tuo posto, un chatbot che risponde alle domande frequenti, un generatore di testi che produce contenuti su richiesta. Tutti strumenti che aspettano un nostro input, elaborano qualcosa nel loro "cervello digitale" e ci restituiscono un output: un testo, un'immagine, un'idea.

Ma mentre ChatGPT aspetta che tu gli chiedi di scrivere una mail, un AI agent può monitorare autonomamente la tua casella di posta, identificare messaggi urgenti, redigere bozze di risposta appropriate al contesto, consultare database aziendali per verificare informazioni e persino programmare riunioni, tutto senza che tu debba dare istruzioni specifiche per ogni passaggio.

Insomma stiamo passando dai sistemi che "suggeriscono" a sistemi che "agiscono" autonomamente. È un incredibile balzo in avanti. Secondo il McKinsey Global Survey 2025, il 23% delle organizzazioni

sta già scalando sistemi agentici¹ e un ulteriore 39% ha iniziato a sperimentarli.

Gli agenti AI rappresentano una frattura sostanziale rispetto a questo modello. Non si limitano a rispondere: agiscono. Ricevono un obiettivo complesso, lo scompongono in passaggi successivi, mobilitano diversi strumenti (API, database, servizi online), prendono decisioni in corso d'opera e arrivano a un risultato finale operativo, tutto senza chiedere conferma a ogni step.

Prendiamo piattaforme come Devin o Manus: l'utente dice "fammi un sito con queste caratteristiche" e in 40-50 minuti riceve un sito funzionante. Devin, in particolare, agisce come "full stack AI developer"² capace di scrivere codice, testarlo e distribuire l'applicazione, ruoli che normalmente richiederebbero persone diverse con competenze distinte.

Quindi non stiamo più parlando solo di "generare un testo" o "riassumere un documento": stiamo parlando di fare un lavoro completo, che può includere scrivere codice, testarlo, interrogare database, orchestrare altre entità digitali e negoziare tra strumenti e sistemi.

La vera promessa in ambito aziendale è diretta e provocatoria: in azienda potresti avere un agente che copre più figure professionali, con un costo inferiore rispetto a una o più risorse umane. Da questo punto in poi, la domanda che si pone naturale è: ma davvero è così semplice? Possiamo veramente delegare decisioni complesse a sistemi che, per quanto intelligenti, rimangono software?

¹ si parla di far crescere o adattare sistemi composti da agenti intelligenti in modo efficiente, mantenendo o migliorando prestazioni, affidabilità e controllo quando aumenta la complessità (più agenti, più utenti, più dati) o quando cambiano le condizioni operative.

² Un full stack AI developer è un professionista versatile che gestisce l'intero ciclo di sviluppo di applicazioni basate sull'intelligenza artificiale, coprendo sia gli aspetti tecnici che quelli operativi. Questa figura combina competenze di sviluppo software tradizionale con expertise specifica in AI, permettendo di realizzare soluzioni end-to-end senza dipendere da team specializzati.

L'evoluzione naturale dell'AI: dal pensiero all'azione. Per capire come siamo arrivati qui, è utile rivedere rapidamente il percorso dell'Intelligenza Artificiale negli ultimi anni.

Dall'esplosione di ChatGPT alla fine del 2022, l'AI generativa ha dominato la scena con i modelli linguistici di grandi dimensioni (LLM),³ addestrati su quantità enormi di dati non strutturati, capaci di rispondere a domande, scrivere testi, tradurre e codificare. Questi modelli hanno dimostrato una capacità straordinaria di "capire" il linguaggio naturale e di adattarsi a scenari su cui non erano stati esplicitamente addestrati.

Ma questa era ancora una forma di AI reattiva: ricevi una domanda e rispondi. L'algoritmo non ha scopo, né memoria a lungo termine o autonomia decisionale. È uno strumento che aspetta.

Gli agenti AI rappresentano l'evoluzione naturale di questa capacità. Combinano le abilità linguistiche e di ragionamento dei modelli generativi con una capacità di azione autonoma: percezione dell'ambiente, elaborazione di strategie, pianificazione di azioni e loro esecuzione per perseguire un obiettivo specifico senza supervisione costante.

Un agente percepisce i dati strutturati e non strutturati (quello che trova nel web, nei documenti, nei database), elabora strategie per raggiungere l'obiettivo, pianifica una sequenza di azioni e le esegue attraverso strumenti e API.⁴ A differenza dei sistemi agentivi tradizionali, che richiedevano programmazione laboriosa basata su regole rigide, gli agenti moderni costruiti su modelli di base generativi hanno il potenziale per adattarsi a scenari diversi: lo stesso agente può

³ Gli LLM (Large Language Models) o modelli linguistici di grandi dimensioni, sono sistemi avanzati di intelligenza artificiale in grado di comprendere, analizzare e generare testo in linguaggio naturale. Si chiamano "large" perché contengono miliardi di parametri e vengono addestrati su enormi quantità di dati testuali, come libri, articoli e contenuti web.

⁴ Un'API (acronimo di Application Programming Interface, ovvero "l'interfaccia di programmazione delle applicazioni") è un insieme di regole e protocolli che consente a diverse applicazioni software di comunicare tra loro per scambiare dati, funzionalità e servizi.

pianificare un itinerario di viaggio complesso, gestire la logistica su più piattaforme, collaborare con altri agenti e imparare dai feedback. Secondo McKinsey e altre fonti di ricerca, questa transizione "*dalla* *pensiero all'azione*" rappresenta la prossima frontiera dopo l'esplosione della generazione di contenuti. Il valore economico che se ne intuisce è enorme: McKinsey stima che gli agenti generativi potrebbero portare tra 2,6 e 4,4 trilioni di dollari di valore aggiunto all'anno nel mondo.

Gli agenti AI non sono dunque una tecnologia separata dai modelli linguistici: gli LLM costituiscono il nucleo centrale (il "cervello") degli agenti. La differenza sta nell'architettura che li circonda: mentre un LLM tradizionale risponde a singole richieste, un agente LLM-based aggiunge i componenti per pianificare, memorizzare e agire nel mondo reale.

Come funzionano: l'architettura della decisione. Un moderno agente di Intelligenza Artificiale può essere descritto come un sistema software che prende decisioni lungo un intero processo, partendo da un obiettivo iniziale e arrivando a un risultato operativo. Per farlo, combina diverse capacità. Anzitutto, ha una forma di "*percezione ambientale*": accede a dati strutturati e non strutturati, come documenti, database, informazioni sul web o tramite API, che gli permettono di costruire un quadro aggiornato del contesto in cui deve agire. Su questa base intervengono i modelli linguistici, che interpretano la richiesta, scompongono il compito, pianificano i passi successivi e modificano il piano in funzione del feedback che arriva man mano dall'ambiente o dall'utente.

Un elemento chiave è la persistenza di memoria: l'agente mantiene uno stato del progetto e della conversazione, ricordando le decisioni prese, i vincoli posti dall'utente, i tentativi già fatti e i risultati parziali, così da non ripartire ogni volta da zero. A questo si aggiunge l'integrazione con strumenti esterni: l'agente non si limita a produrre testo, ma è in grado di chiamare l'API, interrogare i database, usare

software e piattaforme web per eseguire concretamente le azioni decise in fase di pianificazione. Quando queste componenti sono ben orchestrate, l'agente può procedere in autonomia lungo più passi decisionali di seguito, senza dover chiedere conferma umana a ogni singola azione, pur rimanendo sotto supervisione complessiva.

Questa architettura diventa particolarmente potente quando è organizzata secondo un modello a orchestra. In questo caso esiste un "orchestratore" centrale che riceve l'input dall'utente e coordina un insieme di agenti specializzati, ognuno focalizzato su un compito specifico (per esempio analisi dei dati, scrittura di codice, ricerca documentale, controllo di qualità). L'orchestratore scompone la richiesta complessa in sotto-compiti, li assegna agli agenti più adatti, ne supervisiona l'esecuzione e infine ricompone i risultati in un output unico e coerente che l'utente vede come risposta finale.

Dal punto di vista dell'utente tutto questo appare molto semplice: si interagisce con un solo "fronte", si formula una richiesta e si riceve un risultato già rifinito, spesso sotto forma di documento, codice funzionante o procedura operativa pronta all'uso. Dietro le quinte, però, la dinamica è molto più articolata: decine di agenti dialogano tra loro, si scambiano dati intermedi, raffinano più volte i risultati e li verificano reciprocamente, realizzando in modo automatizzato quella che, in un'organizzazione tradizionale, sarebbe una vera e propria catena di lavoro distribuita.

Gli agenti nel mondo reale: dal browser al workflow aziendale. I browser⁵ rappresentano uno dei campi di sperimentazione più interessanti per gli agenti AI. Fino a poco tempo

⁵ Un browser è un programma che permette di accedere alle risorse del Web (pagine, immagini, video) e di navigare tra i siti Internet. Un browser, detto anche "navigatore", è un'applicazione installata su computer o dispositivi mobili che riceve contenuti da server in rete e li presenta in forma comprensibile all'utente. Usa protocolli come HTTP per richiedere le pagine e un motore di rendering per trasformare il codice (per esempio HTML, CSS, JavaScript) in ciò che si vede sullo schermo.

fa, gli sviluppatori integravano strumenti AI nei browser con risultati altalenanti. Ma di recente, con il lancio di ChatGPT Agent di OpenAI e Comet di Perplexity, l'idea di un browser potenziato da un assistente AI autonomo è tornata in auge.

I due sistemi funzionano in modo diverso. Comet è un browser indipendente: l'utente naviga normalmente sul web e poi "convoca" un assistente AI che può aiutarlo a scrivere un'email o completare un'attività di routine. OpenAI ha costruito il suo strumento all'interno di ChatGPT stesso: l'utente dialoga attraverso l'interfaccia web e assegna compiti al bot, che utilizza un browser virtuale per eseguirli. In entrambi i casi, gli agenti possono controllare cursori, inserire testo, cliccare sui link e compilare moduli online.

Se questa tendenza dovesse consolidarsi, il web potrebbe trasformarsi in una città fantasma dove imperversano gli agenti: mentre i software automatici interagiscono 24/7 con servizi, database e piattaforme, gli esseri umani si avventurano sempre più di rado in rete. Le implicazioni sono affascinanti e inquietanti allo stesso tempo.⁶

Nelle aziende, il discorso è ancora più pragmatico. Gli agenti vengono integrati nei flussi di lavoro reale per automatizzare processi che finora richiedevano un coordinamento umano: dall'onboarding⁷ di un nuovo dipendente (l'agente prepara documenti, invia email, configura account) alla gestione della fatturazione (legge documenti contabili, classifica spese, genera report). Il settore legale, fiscale e contabile sta già sperimentando agenti in grado di analizzare contratti e documenti normativi. Per evitare i rischi di frode si stanno già usando agenti per identificare pattern sospetti nei dati finanziari.

⁶ I dati del 2024-2025 confermano il sorpasso: i bot AI rappresentano ora il 51-52,3% del traffico web totale, superando per la prima volta in un decennio l'attività umana. Di questo traffico automatizzato, il 35% proviene da crawler per addestrare LLM (OpenAI, Anthropic, Google DeepMind), mentre il 17% è costituito da bot operativi che automatizzano email, calendari e web scraping. Meta da sola genera il 52% di tutto il traffico AI bot, più di Google (23%) e OpenAI (20%) messi insieme.

⁷ L'onboarding di un nuovo dipendente è il processo strutturato che un'azienda utilizza per accoglierlo, integrarlo e guidarlo nell'utilizzo di un prodotto o servizio.

La promessa economica e il paradosso. La promessa è semplice: più efficienza, meno costi, meno errori umani. Un agente può lavorare 24/7, non si stanca, non dimentica, può gestire decine di task contemporaneamente. A livello organizzativo, significa ridisegnare i ruoli: meno "esecutori di compiti ripetitivi", più "supervisori di agenti" e "strateghi che definiscono gli obiettivi".

Ma qui emerge un paradosso interessante. Da un lato, i professionisti che utilizzano l'AI nel loro lavoro sanno che questo cambiamento è in arrivo. Molti, specialmente nei settori legale, fiscale e contabile, stanno già familiarizzando su come la GenAI possa trasformare il loro lavoro quotidiano. Dall'altro lato, molti altri stanno appena iniziando a capire cosa significhi, per la loro professione, un futuro basato sull'AI e quanto gli agenti rappresentino un ulteriore passo nel buio.

La realtà che ci troviamo di fronte è qualcosa di profondamente interconnesso con il nostro mondo, per le sue dinamiche e i suoi problemi. Non è una visione fantascientifica: è l'automazione di processi reali che già oggi hanno conseguenze economiche e sociali concrete.

Il ruolo della governance e del controllo umano

Qui entrano in gioco la governance, la regolazione e il controllo umano. L'Europa ha tentato di intervenire con l'AI Act definendo i livelli di rischio e gli obblighi di trasparenza. Ma integrare davvero i principi etici dentro questi modelli è complesso, perché l'etica non è un parametro tecnico: è un asset intangibile che cambia con la cultura, col contesto e con i valori di riferimento.

Per uscire dalle secche del dibattito teorico, le organizzazioni hanno bisogno di strumenti pragmatici per decidere dove e come impiegare gli agenti. Non tutti i task sono uguali e non tutti richiedono lo stesso livello di controllo umano.

Un approccio efficace è mappare ogni processo su una Matrice Rischio-Autonomia, incrociando la criticità dell'attività con il grado di libertà concesso all'agente. Questa matrice offre una regola aurea per la governance: all'aumentare del rischio, l'autonomia dell'agente deve diminuire proporzionalmente. L'errore più comune oggi è il disallineamento: aziende che applicano livelli di autonomia da "fascia bassa" a processi di "fascia alta" (es. chatbot che gestiscono reclami sensibili senza supervisione) o viceversa quelle che sprecano risorse umane per controllare task banali che l'AI potrebbe gestire da sola. Adottare questo schema significa accettare che l'efficienza non sia l'unico valore. In fascia alta, l'obiettivo non è "fare prima", ma "decidere meglio". L'agente qui non sostituisce l'esperto, ma ne potenzia la capacità di analisi, lasciando intatta la prerogativa del giudizio.

Oltre la "compliance": tre principi di progettazione responsabile. Adottare una matrice di rischio è il primo passo, ma non basta. Per governare davvero un'architettura ad agenti, le organizzazioni devono integrare tre principi non negoziabili direttamente nel design dei processi, *be* prima che il software venga acceso.

Innanzitutto la tracciabilità per default (Auditability by design). In un sistema dove dieci agenti collaborano per produrre un risultato, "chi ha fatto cosa" diventa rapidamente un mistero. La tracciabilità non può essere un pensiero successivo. Ogni azione critica di un agente – ogni query a un database, ogni testo generato, ogni decisione logica – deve essere registrata in un log immutabile e leggibile dagli umani. Se un agente medico suggerisce una terapia sbagliata, dobbiamo poter ricostruire quale informazione (o allucinazione) ha innescato quell'errore. Senza questa scatola nera accessibile, non esiste governance, ma solo fiducia cieca.

Poi la supervisione significativa (Meaningful Human Control). La legge e il buon senso richiedono la supervisione umana, ma c'è il rischio che diventi una farsa: l'operatore che clicca "Approva" cento volte all'ora senza leggere davvero. Il principio della "supervisione significativa" impone di progettare i flussi di lavoro in modo che l'umano abbia il *tempo*, le *informazioni* e la *competenza* reale per poter contestare la macchina. Se il ritmo imposto dall'agente è troppo alto per permettere un dubbio ragionato, allora non c'è controllo umano, c'è solo un umano usato come timbro di gomma per la responsabilità legale.

Infine il diritto all'intervento (The Right to Override). Nessun processo gestito da agenti deve essere un vicolo cieco. Deve sempre esistere un "pulsante rosso", una procedura chiara e accessibile che permetta all'operatore umano di interrompere l'automazione, revocare una decisione presa dall'agente e prendere il controllo manuale della situazione. Questo diritto all'intervento non è un'emergenza, è una caratteristica di sicurezza standard. Un sistema che non può essere fermato o corretto in corsa da un umano autorizzato non è un sistema autonomo: è un sistema fuori controllo.

Inoltre, la realtà competitiva è spietata. La regolazione arriva sempre dopo la tecnologia. L'AI è già diventata un driver strategico per le aziende e un fattore che influenza fortemente i mercati finanziari. C'è una forte spinta competitiva globale, nonostante non si siano ancora poste delle regole chiare, sottoscritte da tutti gli attori interessati a cominciare dagli Stati Uniti d'America. Il risultato è che gli agenti entrano nelle aziende anche se sappiamo che il loro utilizzo ha ancora limiti enormi e comporta possibili rischi.

Per questo motivo, l'elemento fondamentale di qualsiasi implementazione di agenti AI deve essere questo: la supervisione umana rimane indispensabile. Non nel senso di "controllare ogni

singolo click", sarebbe impossibile e neutralizzerebbe i vantaggi degli agenti, ma nel senso di:

- Avere persone competenti che capiscano il dominio (legale, fiscale, medico) e che sappiano riconoscere quando un output dell'agente è insensato o pericoloso.
- Definire chiaramente i confini dell'autonomia dell'agente: in quali situazioni possa decidere da solo e in quali debba chiedere la conferma umana.
- Mantenere una traccia (audit trail) di quello che l'agente ha fatto, così in caso di problema ci sia responsabilità e tracciabilità.
- Essere consapevoli che gli agenti sono oggi "prematuri" se immaginati in completa autonomia. Mancano di consapevolezza di sé, del senso del limite e della comprensione profonda delle conseguenze.

Il rebus della responsabilità: di chi è la colpa quando l'orchestra stona? Immaginiamo uno scenario ormai plausibile: un'azienda utilizza un "Agente Orchestratore" per gestire la catena di fornitura. Questo agente, autonomamente, ne incarica un secondo (specializzato in logistica) di prenotare una spedizione urgente e un terzo (specializzato in pagamenti) di saldare il conto. Se l'agente logistico sbaglia rotta causando il deperimento della merce o quello dei pagamenti invia fondi a un IBAN fraudolento intercettato online, chi paga?

Nel vecchio mondo del software deterministico, la risposta era spesso legata al "bug": se il codice era scritto male, la colpa era del fornitore. Nel mondo degli agenti AI, dove il comportamento è probabilistico e non deterministico, entriamo in una "terra di nessuno" giuridica.

La frammentazione della responsabilità (The Liability Fragmentation). Il problema principale delle architetture ad agenti è la non linearità. Se l'Agente A (orchestratore) dà un comando

ambiguo all'Agente B (esecutore) e B lo interpreta nel modo peggiore possibile causando un danno, la colpa è dell'ambiguità di A o della mancanza di prudenza di B? Le dottrine giuridiche emergenti, in linea con l'AI Act europeo, tendono a spostare il focus dal *produttore del software* (Provider) all'*utilizzatore professionale* (Deployer).

Il principio della "Culpa in eligendo" e della "Culpa in vigilando". Le prime analisi legali suggeriscono che le aziende saranno chiamate a rispondere secondo due antichi principi del diritto, riadattati all'era digitale: Culpa in eligendo (Colpa nella scelta): se un'azienda decide di affidare un processo critico (come i pagamenti) a un sistema di agenti autonomi, si assume il rischio d'impresa intrinseco a quella scelta. Non potrà difendersi dicendo "è stato il software". Aver scelto uno strumento probabilistico per un task che non tollera errori è, di per sé, una negligenza manageriale e si tratta di Culpa in vigilando (Colpa nella sorveglianza): qui torna in gioco il tema della governance. Se l'azienda non ha predisposto sistemi di monitoraggio adeguati (*Human-in-the-loop*)⁸, è responsabile per non aver fermato l'agente in tempo. La responsabilità non è dell'agente che sbaglia, ma dell'umano che non lo stava guardando.

Il fornitore vs. l'utilizzatore: il caso del "SaaS che agisce"⁹. Un punto critico riguarda i termini di servizio (ToS). I grandi provider di LLM (come OpenAI, Google, Anthropic) nei loro contratti specificano chiaramente che l'output dei modelli è fornito "as is", ciò significa che i risultati generati dai modelli vengono consegnati nello stato grezzo in cui sono prodotti, senza garanzie, modifiche, verifiche o supporto

⁸ Human-in-the-loop (HITL) è un approccio in cui gli esseri umani intervengono attivamente nei processi di intelligenza artificiale e machine learning per supervisionare, correggere o migliorare le decisioni automatizzate. Questo metodo integra il feedback umano nel ciclo di addestramento e funzionamento dei modelli AI, garantendo maggiore accuratezza, etica e affidabilità.

⁹ Software as a Service (SaaS)(software come servizio) è un modello di distribuzione del software in cui le applicazioni vengono fornite da un provider tramite Internet, accessibili via browser senza installazione locale. Il provider gestisce infrastruttura, aggiornamenti e manutenzione, mentre gli utenti pagano tipicamente su abbonamento.

ulteriore da parte del fornitore e che quindi la verifica spetta all'utente. Tuttavia, la nuova Direttiva UE sulla responsabilità per danno da prodotti difettosi sta estendendo il concetto di "prodotto" anche al software e all'AI. Se un agente viene venduto come "autonomo" ed è capace di svolgere un task specifico e poi fallisce clamorosamente non per un errore dell'utente, ma per un difetto di addestramento, il produttore potrebbe non potersi più nascondere dietro le clausole di esonero di responsabilità.

Verso una "Responsabilità oggettiva" per l'AI? Per i settori ad alto rischio (sanità, infrastrutture critiche), la tendenza dottrinale è quella di muoversi verso una responsabilità oggettiva (*strict liability*), simile a quella che si applica alle attività pericolose (art. 2050 del Codice Civile italiano). Chi trae profitto economico dall'uso di agenti autonomi deve rispondere dei danni che questi provocano, a prescindere dalla colpa o dal dolo specifico, semplicemente come costo del rischio introdotto nel sistema.

In sintesi: l'autonomia dell'agente non è un'attenuante per l'umano, ma un'aggravante. Più autonomia diamo al sistema, più stringente diventa il nostro dovere di garanzia verso terzi.

Dove vengono usati gli Agenti?

Se l'impatto su settori come quello legale, fiscale e contabile appare ormai una traiettoria definita, è guardando ad altri tre ambiti cruciali (sanità, pubblica amministrazione e marketing) che si coglie la reale portata trasformativa degli agenti AI. La promessa è sempre la stessa: più efficienza, automazione dei processi e liberazione di risorse umane qualificate. Ma è proprio nelle specificità di questi mondi che emergono le implicazioni più profonde sul lavoro e sulla responsabilità.

Sanità: l'agente come "co-pilota" clinico. Nel settore sanitario, gli agenti si stanno configurando come potenti "co-piloti" per medici e personale infermieristico. Il loro campo d'azione ideale è l'orchestrazione di compiti complessi con una forte componente procedurale: Gestione delle liste d'attesa: un agente può incrociare in tempo reale le disponibilità di diverse strutture, la priorità clinica dei pazienti (definita da un medico) e i percorsi diagnostici richiesti, ottimizzando le agende in modo dinamico. Monitoraggio dei pazienti cronici: un agente collegato a dispositivi indossabili può analizzare i dati vitali di un paziente diabetico o iperteso, segnalando al medico solo le anomalie che superano una soglia di rischio predefinita e richiedono un intervento umano.

Il vantaggio è evidente: liberare tempo prezioso, riducendo il carico burocratico e permettendo al personale sanitario di concentrarsi sulla clinica e sulla relazione con il paziente. Ma è proprio questo vantaggio a definire, per contrasto, i confini delle competenze non delegabili.

Un agente può identificare un pattern statistico di rischio, ma non può assumersi la responsabilità di: formulare una diagnosi complessa, dove l'intuizione clinica, l'esperienza e la comprensione del non-detto del paziente sono fondamentali; ad esempio, nel distinguere tra sintomi apparentemente simili, ma che siano riconducibili a patologie molto diverse; comunicare una prognosi infausta: un compito che richiede empatia, sensibilità e una capacità di adattamento al contesto emotivo della persona, che nessun algoritmo può replicare; prendere una decisione etica: come allocare una risorsa scarsa (un posto in terapia intensiva, un farmaco sperimentale) quando le linee guida non sono sufficienti e subentra il giudizio umano sul "caso per caso".

Il rischio più insidioso, qui, non è l'errore tecnico dell'algoritmo, ma la sua cecità al contesto. Un agente addestrato a massimizzare l'efficienza potrebbe, in teoria, suggerire di depriorizzare un paziente

anziano con comorbidità a favore di uno più giovane con maggiori probabilità di recupero rapido. È il medico, con il suo bagaglio di etica e responsabilità professionale a dover porre un veto, affermando un principio di cura che trascende la mera ottimizzazione statistica. Il ruolo umano si sposta così dalla compilazione di dati all'interpretazione critica degli stessi, diventando il garante ultimo del processo di cura.

Pubblica Amministrazione: da esecutore a garante di equità

Nella Pubblica Amministrazione, gli agenti promettono di smantellare le lunghe code e la burocrazia fine a sé stessa. Un "funzionario virtuale" può già oggi:

- Guidare l'utente nella compilazione di un'istanza: ad esempio, per un bonus edilizio, verificando in tempo reale la correttezza dei dati catastali tramite API e segnalando i documenti mancanti prima ancora dell'invio.

- Orchestrare un processo inter-ufficio: come il rilascio di un'autorizzazione commerciale, gestendo in parallelo le verifiche presso l'ufficio tecnico, l'ASL e la Camera di commercio.

Il ruolo del funzionario umano si evolve: non più un mero esecutore di procedure, ma un validatore del processo e un garante di equità. Il rischio, infatti, è sostituire il "muro di gomma" dello sportello fisico con l'opacità di un "filtro algoritmico", che applica le regole in modo ferreo ma cieco.

Le competenze umane non delegabili diventano quindi: la gestione dell'eccezione: riconoscere quando un caso individuale, pur non rientrando perfettamente nei parametri, ha diritto a essere trattato. Pensiamo a un cittadino anziano senza SPID che deve accedere a un servizio essenziale.; la tutela contro i bias: vigilare affinché l'agente, magari addestrato su dati storici distorti, non discrimini

involontariamente determinate categorie di cittadini, perpetuando disuguaglianze esistenti. Il funzionario diventa colui che non solo controlla la correttezza formale dell'output dell'agente, ma ne valuta l'impatto sostanziale, assicurando che l'automazione non diventi un alibi per l'ingiustizia.

Marketing: lo stratega umano contro l'autocrazia delle metriche. Nel marketing, la trasformazione è già in atto. Gli agenti non si limitano più a generare contenuti, ma orchestrano intere campagne. Un agente può analizzare i trend di mercato, definire il target di una nuova linea di prodotti, generare creatività, pianificare la spesa pubblicitaria su più canali e aggiustare la strategia in tempo reale sulla base dei tassi di conversione. La creatività umana sembra ridotta a un ruolo di supervisione. Ma è proprio qui che emerge il ruolo insostituibile dello stratega. L'agente è un ottimizzatore per eccellenza, ovvero data una metrica (es. massimizzare i click), troverà sempre la via più breve per raggiungerla. Questo può portare a un appiattimento della comunicazione, dove tutti i brand convergono verso le stesse soluzioni "che funzionano", perdendo identità. Lo stratega umano diventa il custode della visione di lungo periodo e dei valori del brand. Il suo compito non è scegliere l'immagine con il CTR¹⁰ più alto, ma decidere se quell'immagine è coerente con la storia e il posizionamento del marchio. È la persona che pone il veto a una campagna potenzialmente virale, ma eticamente discutibile, che sacrifica un guadagno a breve termine per proteggere la reputazione. L'agente vince la battaglia tattica; lo stratega umano vince la guerra di posizionamento, assicurando che il brand non diventi schiavo dell'autocrazia delle metriche.

¹⁰ CTR sta per Click-Through Rate, ovvero il tasso di clic calcolato come $(\text{clic} / \text{impressioni}) \times 100$, che misura l'attrattiva di un elemento visivo come un'immagine in ads o email. Un CTR alto indica che l'immagine cattura meglio l'attenzione del pubblico target.

Quali sono i problemi irrisolti?

Ma prima di concludere che gli agenti AI risolveranno tutto, occorre essere cauti poiché gli agenti, per quanto sofisticati, ereditano tutti i problemi dell'AI generativa che già conosciamo e che in alcuni casi amplificano.

Le allucinazioni rimangono il primo problema. I modelli linguistici moderni continuano a produrre risposte errate, ma espresse in modo estremamente convincente. Non è un difetto marginale: test recenti hanno mostrato che modelli celebrati per il loro "ragionamento avanzato" hanno tassi di allucinazione anche doppi rispetto ai modelli precedenti. Testi ben scritti, ma semanticamente scorretti. E se questo era già rischioso per un chatbot che scrive un'email, diventa potenzialmente pericoloso quando l'agente agisce sulla base di informazioni allucinate: interroga un database con dati inventati, invia richieste fondate su "fatti" che non esistono, prende decisioni su premesse false.

Il problema della scatola nera si aggrava. Già sappiamo poco come i modelli generativi arrivino a una conclusione. Con gli agenti, questa opacità si moltiplica. Se un agente è composto da dieci agenti specializzati che conversano tra loro, coordinati da un orchestratore, chi può dire davvero cosa è accaduto nel momento in cui è stata presa una decisione sbagliata? L'utente vede input e output, ma il processo interno rimane nascosto. Quindi il rischio di "explainability washing" (fingere di essere trasparenti mostrando solo quello che conviene) è tutt'altro che remoto.

Gli agenti non hanno etica intrinseca. Questo è il punto forse più importante. I sistemi di AI non hanno coscienza, autocoscienza, senso del limite, comprensione delle conseguenze etiche delle proprie azioni. Se l'etica non è stata progettata e implementata dagli umani fin dall'inizio (quella che gli esperti chiamano "etica by design"), il sistema

non potrà seguirla. Quando un agente ignora completamente gli impatti etici di una decisione, non è perché è "cattivo": è perché è stato costruito senza quei vincoli.

Chi serve per governare questi problemi?

Se gli agenti AI promettono di trasformare processi e organizzazioni, la conseguenza quasi inevitabile è l'emergere di nuove figure professionali. Non basta "usare" gli agenti: occorre progettarli, coordinarli, valutarli, incardinarli in cornici etiche e giuridiche. Invece di limitarci alla formula generica dei "supervisor di agenti", possiamo iniziare a dare un nome e un perimetro più preciso ad alcuni ruoli che, nei prossimi anni, potrebbero diventare centrali.

Di seguito tre profili emergenti che segnano bene il passaggio da un lavoro puramente esecutivo a un lavoro di governo dell'automazione.

L'"Orchestratore di Agenti AI": Se l'architettura a "orchestra" degli agenti si diffonderà davvero, servirà qualcuno che di quell'orchestra sia il direttore. L'Orchestratore di Agenti AI è la figura responsabile di progettare, coordinare e monitorare il lavoro di più agenti specializzati all'interno di un'organizzazione. Il suo compito non è scrivere il codice di ogni agente, ma quello di tradurre gli obiettivi strategici dell'azienda (o dell'ente pubblico) in task agentizzabili, decidere quali agenti attivare, con quali permessi, sulla base di quali dati; definire le regole di ingaggio tra agenti diversi (chi fa cosa, in che ordine, con quali alternative); sorvegliare le prestazioni complessive del sistema, identificando colli di bottiglia, ridondanze e conflitti tra output.

È una figura ibrida che deve capire di tecnologia quanto basta per non essere ostaggio dei tecnici, ma deve anche conoscere in profondità il dominio applicativo (finanza, sanità, risorse umane, pubblica amministrazione). In un certo senso, prende il posto di molti "middle manager" che oggi coordinano team umani, ma con una differenza fondamentale: il suo oggetto di lavoro non sono persone, bensì ecologie di agenti. Eppure, le ricadute sono profondamente umane, perché da come orchestra l'automazione dipendono i carichi di lavoro, i tempi e la qualità dei servizi offerti alle persone.

Il "Validatore" di processo automatizzato: Mentre gli agenti entrano nei processi decisionali, siano sanitari, amministrativi o finanziari, diventa indispensabile una figura che abbia il compito esplicito di mettere alla prova l'automazione. Il Validatore di Processo Automatizzato è colui che non si limita a verificare se l'output di un agente "funziona", ma si chiede se: è corretto nel merito; è robusto rispetto a scenari diversi da quelli di addestramento; è compatibile con le norme, i regolamenti e i diritti delle persone coinvolte. Nella pratica, questo ruolo può assumere forme diverse: nel settore sanitario, può essere il clinico che valuta sistematicamente come l'output dell'agente di "triage" si confronta con il giudizio medico, identificando pattern di errori ricorrenti; nella PA, può essere il funzionario che analizza un campione di pratiche gestite dall'agente, verificando se l'algoritmo ha trattato in modo equo situazioni simili o se emergono bias sistematici verso certi gruppi di utenti; in banca o in assicurazione, può essere l'analista che controlla se un sistema di scoring automatizzato esclude sistematicamente alcune categorie di richiedenti. È una professione che richiede competenze metodologiche (statistica, valutazione, audit), conoscenza del dominio e, soprattutto, indipendenza di giudizio: il validatore deve poter dire "no" all'automazione quando questa produce effetti non accettabili, anche se "sui numeri" sembra efficiente. In un mondo dove l'AI Act europeo chiede trasparenza, tracciabilità, gestione del rischio, il Validatore di Processo

Automatizzato diventa un tassello naturale delle nuove catene del valore.

L'Eticista degli algoritmi: Infine, c'è una figura che in parte esiste già, ma che l'avvento degli agenti rende ancora più indispensabile: l'Eticista degli Algoritmi. Non si occupa solo di redigere codici etici o linee guida da seguire, ma di tradurre principi etici astratti in vincoli concreti di progetto e di utilizzo. Il suo lavoro si colloca in almeno tre livelli: a monte, nella fase di progettazione nel definire quali dati possano essere utilizzati, quali obiettivi di ottimizzazione siano accettabili e quali no (massimizzare il profitto? ridurre il tempo medio di attesa? garantire priorità a soggetti vulnerabili?); nel mezzo, durante lo sviluppo, nel collaborare con data scientist, giuristi e decisori per identificare i punti in cui l'automazione può entrare in conflitto con i diritti fondamentali (privacy, non discriminazione, accesso ai servizi essenziali) e nel proporre salvaguardie; a valle, nella valutazione in esercizio, nel monitorare l'impatto reale degli agenti su persone e gruppi sociali, raccogliendo segnalazioni, casi problematici, effetti inattesi e aggiornando di conseguenza le policy. In questa prospettiva, la figura dell'Eticista degli Algoritmi è tutt'altro che ornamentale. È qualcuno che porta dentro l'organizzazione competenze filosofiche, giuridiche, sociologiche, storiche e le usa per opporre resistenza argomentata a una visione puramente efficientista dell'AI. È di fatto un mediatore tra diversi mondi, tra chi costruisce i sistemi, chi li governa e chi li subisce. Non è un caso che questo profilo, in molte realtà, nasca dall'incontro tra competenze tecnico-scientifiche e una forte formazione umanistica. In un contesto in cui gli agenti AI promettono di automatizzare sempre più decisioni, la domanda non è solo "che cosa possono fare?", ma "fino a dove vogliamo lasciarli arrivare?". L'Eticista degli Algoritmi aiuta a dare una risposta non improvvisata a questa domanda, portando dentro il processo di innovazione l'elemento che la tecnologia, da sola, non potrà mai fornire: un criterio che abbia un senso.

Ripensare le competenze: dalla logica del fare alla logica dell'essere

Se i nuovi ruoli organizzativi – l'Orchestratore di Agenti AI, il Validatore di processo automatizzato, l'Eticista degli algoritmi – descrivono che cosa le persone faranno nel prossimo decennio, è altrettanto importante comprendere come dovranno trasformarsi per abitare questi spazi professionali. Non basta semplicemente aggiungere "AI Literacy"¹¹ al bagaglio di competenze: occorre un ripensamento più radicale di ciò che significa acquisire competenze nell'era degli agenti, un ripensamento che tocca le fondamenta stesse dell'identità professionale e umana.

Gli Agenti AI rappresentano una rivoluzione industriale accelerata, con tempi di implementazione molto più rapidi rispetto alle precedenti ondate tecnologiche. La rottura innovativa è inevitabile, ma chi si adatta proattivamente, anziché subire passivamente il cambiamento, può trasformare questa sfida in un'opportunità genuina di crescita professionale ed economica. La questione non è se cambieremo, ma come cambieremo e che senso daremo a questo cambiamento.

Una disparità di genere che rivela strutture di lavoro obsolete

La ricerca dell'ILO (Organizzazione Internazionale del Lavoro) evidenzia un dato allarmante: significative differenze di genere nell'esposizione all'AI Generativa. Nei paesi ad alto reddito, l'8,5% dell'occupazione femminile si trova in lavori ad alto potenziale di automazione, rispetto al 3,9% di quella maschile. Questo pattern non è casuale, bensì il riflesso di una maggiore concentrazione delle donne in ruoli amministrativi e di supporto, particolarmente esposti alla tecnologia. Tuttavia, l'effetto "augmentation" – la capacità degli agenti di potenziare le competenze umane piuttosto che sostituirle – mostra

¹¹ L'AI Literacy rappresenta un insieme di competenze che include conoscere i concetti base dell'AI, applicarne gli strumenti in modo appropriato, valutarne i limiti e riflettere sulle implicazioni etiche, sociali e legali.

tendenze più equilibrate, con opportunità significative per entrambi i generi nei settori ad alta intensità di conoscenza. La vera sfida allora, non è tecnica ma politica: garantire parità di accesso alla formazione e riqualificazione professionale, affinché l'agentificazione del lavoro non diventi un nuovo meccanismo di disuguaglianza, ma un'occasione di riequilibrio e inclusione.

Competenze immediate e strategie di lungo periodo: cosa fare adesso? Per i professionisti, l'imperativo è duplice e richiede un'azione simultanea su due livelli temporali. Nel breve termine, ovvero nei prossimi 12-18 mesi, è indispensabile acquisire competenze fondamentali. In primo luogo, una AI Literacy di base: non si tratta di diventare tecnici o esperti di machine learning, bensì di comprendere le potenzialità e i limiti dell'AI Generativa come utenti consapevoli, capaci di riconoscere sia le promesse che i rischi. In secondo luogo, la sperimentazione diretta con strumenti come ChatGPT, Copilot e Claude per attività quotidiane, trasformando la paura iniziale in familiarità operativa. Parallelamente, imparare il Prompt Engineering, ovvero l'arte di comunicare efficacemente con sistemi di AI, permette di scoprire che la chiarezza linguistica e la precisione comunicativa diventano risorse professionali preziose, non solo nell'interazione con le macchine, ma anche nella collaborazione umana. Infine, investire nel potenziamento delle competenze relazionali – intelligenza emotiva, leadership, capacità di negoziazione – diventa strategico proprio perché rappresentano il confine invalicabile dell'automazione. Nel lungo periodo, invece, emerge la necessità di tre strategie di posizionamento più profonde. La prima è specializzarsi in aree dove l'AI è complementare piuttosto che sostitutiva. Non si tratta di evitare i settori interessati dall'automazione – sarebbe illusorio – bensì di identificare consapevolmente i ruoli dove la tecnologia amplifica il valore umano piuttosto che eroderlo e investire lì. La seconda è costruire competenze ibride: combinare l'expertise tecnica (come leggere un output di AI, valutarne l'affidabilità, e identificarne i bias)

con le capacità umane distintive come empatia, giudizio etico e visione strategica di lungo periodo. La terza è mantenere un atteggiamento di apprendimento continuo e adattabilità, consapevoli che la velocità di trasformazione tecnologica implica che nessun bagaglio formativo rimanga valido più di 3-5 anni senza un significativo aggiornamento. Per le organizzazioni, la sfida è strutturalmente simile, ma su scala collettiva richiede una trasformazione culturale oltre che tecnica. Un'adozione graduale e strategica degli agenti passa innanzitutto attraverso la formazione diffusa, non limitata alle figure di primo contatto, ma estesa a tutti i dipendenti, per creare un punto di riferimento culturale condiviso attorno ai temi dell'AI. Secondo elemento cruciale è la sperimentazione circoscritta: iniziare con progetti pilota su processi non critici, imparare dai fallimenti, identificare i pattern di successo per poi estendere progressivamente il campo d'azione. Parimenti importante è la gestione etica, cioè stabilire linee guida chiare per l'uso responsabile dell'AI prima che il sistema venga acceso, non come ripensamento successivo quando i danni sono già stati fatti. Infine, la riqualificazione proattiva non insegna semplicemente ai dipendenti come usare gli agenti, bensì come trasformarsi per lavorare con loro, rimanendo centrali nelle decisioni e nel controllo qualitativo.

Una gestione del cambiamento che realmente funziona coinvolge i lavoratori nelle decisioni di implementazione dell'AI, promuove una cultura della sperimentazione e innovazione senza punire il fallimento calcolato e soprattutto si concentra sul potenziamento delle capacità umane piuttosto che sulla mera automazione e sulla riduzione dei costi. In questo senso, le organizzazioni più intelligenti vedranno gli agenti non come sostituti, ma come alleati che liberano tempo e risorse mentali per attività a maggior valore aggiunto.

Le prospettive per l'Italia: una finestra di opportunità che si sta restringendo. L'Italia si trova in una posizione paradossale e insieme straordinaria. Da un lato, il suo DNA industriale e creativo – il

Made in Italy, le filiere manifatturiere di qualità, i distretti creativi e artigianali – potrebbe essere significativamente potenziato dagli Agenti AI, trasformando i vincoli di scala e velocità, che storicamente hanno limitato la competitività globale del tessuto PMI italiano. Secondo lo studio di Implement Consulting Group per Google, l'adozione diffusa dell'AI Generativa potrebbe contribuire a un aumento del PIL italiano dell'8% in dieci anni, equivalente a 150-170 miliardi di euro di valore aggiunto annuale. È una cifra che non riflette la semplice crescita economica, bensì una trasformazione strutturale della capacità produttiva del paese.

Dall'altro lato, tuttavia, permangono sfide ostative che rischiano di impedire all'Italia di cogliere questa opportunità. Innanzitutto, il gap infrastrutturale e di competenze: malgrado i progressi recenti, la carenza di infrastrutture digitali di qualità e la mancanza diffusa di competenze AI di base rimangono colli di bottiglia significativi, particolarmente nelle regioni del Mezzogiorno. In secondo luogo, l'adozione tra le PMI è ancora marginale: solo il 7% delle piccole e il 15% delle medie imprese ha progetti AI attivi, quando proprio queste categorie potrebbero trarne maggior vantaggio in termini di efficienza operativa e accesso a mercati globali. Infine, gli istituti di istruzione e formazione professionale non sono ancora sufficientemente allineati sulle competenze richieste dall'era degli agenti, producendo un disallineamento tra ciò che le scuole insegnano e ciò che il mercato chiede.

La finestra di opportunità rimane aperta, ma si sta restringendo visibilmente. I paesi che investiranno massicciamente da subito in formazione, infrastruttura digitale e sperimentazione organizzata avranno un vantaggio competitivo decisivo tra 5-7 anni. Per l'Italia, questo significa che le scelte di investimento che facciamo nel 2025-2026 determineranno la posizione competitiva del paese nel 2032-2033. L'urgenza, quindi, non è retorica: è strutturale.

Oltre le skill: l'essere e il fare nell'era dell'AI. Tuttavia, al di sotto di queste considerazioni pragmatiche e macroeconomiche, emerge una questione più profonda, una che trascende la mera accumulazione di competenze "tecniche" e tocca il significato stesso del lavoro umano. Non stiamo semplicemente aggiungendo una nuova materia al curriculum scolastico o al programma di upskilling aziendale. Stiamo affrontando una trasformazione ontologica – cioè, una trasformazione della nostra stessa natura e identità – di ciò che significa lavorare, creare valore, contribuire a una comunità, essere presenti nel mondo. La distinzione tra competenze dell'essere e competenze del fare diventa strategica proprio in questo contesto. La trasformazione tecnologica non è neutrale dal punto di vista esistenziale. Come insegna la fenomenologia, ogni tecnologia modifica il nostro modo di essere nel mondo: non è uno strumento che usiamo rimanendo immutati, bensì una forza che ci trasforma mentre la usiamo. L'AI Generativa non è semplicemente un software intelligente, bensì una trasformazione dell'esistere che richiede una ridefinizione radicale delle competenze in chiave ontologica.

Da un punto di vista pratico, poi, un buon sistema di orientamento formativo dovrebbe distinguere quali competenze sviluppare prioritariamente in base alla loro natura ontologica, sapendo che competenze diverse richiedono strategie di sviluppo diverse. Purtroppo, gli attuali sistemi di orientamento lasciano molto a desiderare, spesso limitandosi a mappature meccaniche di skill o a semplici proiezioni occupazionali. Eppure, avere un orientamento corretto è fondamentale per giungere a una strategia professionale che consenta ai professionisti di investire consapevolmente in competenze "uniquely human" – irriducibili alla logica computazionale e al calcolo algoritmico – che si arricchiscono e si trasformano attraverso l'ibridazione con l'AI.

Da un punto di vista antropologico, infine, la distinzione essere/fare nell'era dell'AI tocca una domanda fondamentale esistenziale: che cosa significa essere umani quando le macchine diventano

"intelligenti"? Le competenze dell'essere ci ricordano che l'umano non si definisce per ciò che fa – per i compiti che svolge, i prodotti che genera, i risultati che ottiene – bensì per come esiste, per la sua capacità irrinunciabile di interrogare il senso delle cose, di aprirsi all'alterità dell'altro, di creare significato attraverso la narrazione e il dialogo. Queste sono le competenze che rimangono intatte, anzi, che diventano ancora più preziose, proprio perché irriducibili all'automazione, perché non codificabili in un algoritmo.

Un nuovo paradigma pedagogico: educazione come vocazione ontologica. La distinzione essere/fare richiede un ripensamento radicale dei percorsi formativi, un ripensamento che tocca scuole, università, e organizzazioni. In primo luogo, è necessario sviluppare e proteggere la capacità riflessiva, l'intelligenza emotiva, la creatività come dimensioni costitutive della persona, non come "soft skills" marginali o supplementari, bensì come il nucleo stesso di una formazione umana autentica. In secondo luogo, occorre padroneggiare gli strumenti AI, non come sostituti delle capacità umane, bensì come estensioni intelligenti di quelle capacità che ne amplificano la portata. In terzo luogo, si tratta di integrare essere e fare in modalità creativa e responsabile, sapendo che il valore aggiunto umano nel prossimo decennio risiederà nella capacità di governare l'automazione, di dare significato ai dati e di prendersi cura dell'impatto umano e sociale delle scelte tecnologiche.

Come suggerisce il pedagogista brasiliano Paulo Freire, siamo "esseri che stanno essendo", ovvero mai completati, sempre in costante processo di formazione e trasformazione. Nell'era dell'AI, l'apprendimento continuo non è solo una necessità professionale, una skill da aggiornare periodicamente, ma una vocazione ontologica: il modo in cui realizziamo la nostra umanità in dialogo creativo con l'AI, il modo in cui rimaniamo vivi, consapevoli e capaci di crescere.

Verso un umanesimo tecnologico. L'analisi ontologica delle competenze nell'era dell'AI Generativa rivela che non stiamo assistendo a una sostituzione dell'umano – l'apocalittica narrativa del robot che prende il tuo lavoro – bensì a una sua ridefinizione profonda. Le competenze dell'essere diventano ancora più preziose proprio perché irriducibili alla logica computazionale, perché risiedono nel regno del significato, della relazione, della responsabilità. Le competenze del fare si evolvono in direzione collaborativa e sinergica, dove il lavoro non è più esercizio solitario di un'abilità, bensì orchestrazione intelligente tra intuizione umana e capacità di calcolo della macchina. Le competenze ibride aprono nuovi orizzonti di possibilità esistenziali, permettendo ai professionisti di operare in spazi che prima erano inaccessibili per mancanza di tempo, energia o risorse.

Questa prospettiva ci invita a pensare il futuro del lavoro non come una competizione uomo-macchina, una battaglia dove uno vince e l'altro perde, bensì come l'emergere di una nuova forma di umanesimo tecnologico, dove la tecnica è al servizio della realizzazione delle potenzialità umane più autentiche e più profonde. In questo scenario evolutivo, investire nelle competenze dell'essere – nella capacità di pensare criticamente, di relazionarsi empaticamente, di creare significato – non è nostalgico conservatorismo o una difesa romantica di un passato irrecuperabile, ma la via più radicale e necessaria per preparare un futuro in cui la tecnologia amplifichi ciò che di più prezioso abbiamo come specie: la nostra capacità di essere umani nel mondo, di costruire comunità, di dare senso al nostro vivere insieme.

Il ritorno dei filosofi: perché le competenze umanistiche sono la nuova ingegneria?

Per anni ci siamo sentiti ripetere un mantra: "servono più STEM" (Science, Technology, Engineering, Mathematics). In un mondo in cui bisognava costruire l'infrastruttura digitale, era vero. Ma nell'era degli agenti AI, il paradigma si sta ribaltando in modo inaspettato. Se strumenti come Manus¹² scrivono codice, testano software e gestiscono database meglio e più velocemente di un junior developer, il valore aggiunto umano non risiede più nella capacità di parlare la lingua della macchina (il codice), ma nella capacità di insegnare alla macchina a comprendere il mondo umano. Le competenze umanistiche – filosofia, storia, sociologia, linguistica – smettono di essere un ornamento culturale e diventano il "sistema operativo" necessario per governare gli agenti.

Dall'ingegneria del prompt all'ingegneria del contesto (Ermeneutica). Interagire con un'architettura di agenti non significa dare un comando ("fai questo"), ma definire un contesto di azione. Richiede la capacità di anticipare le ambiguità, di chiarire le intenzioni e di interpretare i risultati parziali. Questa è, in termini tecnici, l'ermeneutica: l'arte dell'interpretazione. Un ingegnere può ottimizzare la velocità di un agente, ma serve una mente allenata alla complessità linguistica e logica per accorgersi che l'agente ha "frinteso" una sfumatura culturale in una campagna di marketing o in una diagnosi. Saper porre la domanda giusta diventa più difficile – e più prezioso – che saper trovare la risposta.

La storia contro il bias dei dati. Gli agenti AI sono addestrati sul passato (i dati storici). Per definizione, tendono a replicare il passato.

¹² Manus AI è un "agente generale" che pianifica ed esegue task autonomamente, come analisi dati, programmazione o prenotazioni, navigando il web senza supervisione umana costante.

Se i dati storici sulle assunzioni in un'azienda sono viziati da pregiudizi di genere, l'agente HR (Risorse Umane) tenderà a replicarli, scambiandoli per "efficienza". Qui entra in gioco la coscienza storica. Solo chi ha studiato le evoluzioni sociali sa distinguere tra una "correlazione statistica" e una "causa storica". Lo storico sa che il passato non è una guida infallibile per il futuro, ma un archivio di scelte da contestualizzare. In azienda, questo significa avere persone capaci di dire all'agente: "Ignora questo pattern (modello) nei dati perché è frutto di un contesto sociale che non accettiamo più".

Etica applicata: oltre il "carrello ferroviario". Fino a ieri, l'etica dell'AI era un esercizio teorico (il famoso "dilemma del carrello"). Con gli agenti autonomi, diventa pratica quotidiana. Un agente di pricing deve decidere quanto aggressivamente alzare i prezzi durante un picco di domanda. Massimizzare il profitto è l'istruzione di default. Ma quando questo diventa sciacallaggio? La capacità di navigare queste zone grigie non si impara su GitHub¹³. Richiede una formazione in filosofia morale, capace di bilanciare utilitarismo (il bene per i più) e deontologia (ci sono regole che non si violano mai). Le aziende avranno bisogno di "Chief Ethics Officer" non per fare bella figura, ma per evitare disastri reputazionali causati da agenti troppo zelanti.

In sintesi. Paradossalmente, più l'AI diventa sofisticata e "umana" nelle sue interazioni, più abbiamo bisogno di umanisti. Non per rifiutare la tecnologia, ma per fornirle quella visione d'insieme, quel pensiero critico e quella profondità etica che nessun algoritmo, per quanto avanzato, può calcolare.

¹³ GitHub è una piattaforma online per sviluppatori che permette di ospitare, gestire e condividere codice sorgente, consentendo la collaborazione in tempo reale su progetti software.

Il bivio del 2030: tre scenari per un futuro "agentificato"

Se l'agentificazione è il trend tecnologico, il suo impatto sociale non è deterministico. Dipende da come incroceremo le variabili della regolamentazione, dell'accesso alla tecnologia e della cultura aziendale. Applicando le tecniche di previsione strategica, possiamo delineare tre possibili orizzonti per il mercato del lavoro alla fine del decennio.

Scenario 1: L'ecosistema collaborativo (Il modello "Centauro"). È lo scenario high-trust / high-regulation. Qui, la governance ha funzionato. Le aziende hanno adottato il principio della "supervisione significativa". Nel 2030, il lavoratore tipico opera come un "Centauro": una testa umana su un corpo digitale potente. Non passiamo più tempo a compilare Excel o a cercare file: il 70% del lavoro operativo è delegato a una flotta personale di agenti. L'umano si concentra esclusivamente su tre attività ad alto valore aggiunto: la definizione della strategia (dire agli agenti cosa cercare), la gestione delle eccezioni (risolvere ciò che gli agenti non capiscono) e la cura delle relazioni (empatia, negoziazione, leadership). In questo scenario, la produttività è esplosa, ma invece di tagliare posti di lavoro, le aziende hanno ridotto l'orario lavorativo o reinvestito in progetti di innovazione, prima impossibili per mancanza di tempo. Gli agenti sono "colleghi trasparenti": sappiamo sempre perché hanno preso una decisione.

Scenario 2: La rete opaca (La "Black Box Society"). È lo scenario della complessità "non governata". Abbiamo lanciato milioni di agenti autonomi senza standard di interoperabilità e controllo comuni. Il web è diventato quella "città fantasma" che temiamo: agenti di marketing parlano con agenti d'acquisto in un loop infinito di

transazioni velocissime e incomprensibili per l'umano. Nel mercato del lavoro, questo crea una crisi di fiducia. I manager non si fidano più dei report, perché non sanno se sono stati generati da un'allucinazione a cascata di dieci agenti diversi. Nascono i lavori di "Bonifica Digitale": squadre di umani pagati per disattivare automazioni impazzite, ripulire i dataset corrotti dai contenuti sintetici e ripristinare processi manuali, quando i sistemi digitali collassano sotto il peso della loro stessa complessità. È un mondo efficiente in apparenza, ma fragile e nevrotico nella sostanza.

Scenario 3: Il feudalesimo digitale (La concentrazione del potere). È lo scenario critico, dove la tecnologia è matura ma l'accesso è disuguale. Pochi giganti tecnologici possiedono gli "Agenti Orchestratori" più potenti (i "Super-Manager AI"). Le piccole imprese e i professionisti indipendenti non possono permettersi questi strumenti premium e rimangono tagliati fuori dal mercato, incapaci di competere in velocità e costi. Il lavoro si polarizza: da una parte un'élite di "Feudatari dell'AI" che governano imperi automatizzati con pochissimi dipendenti; dall'altra una massa di "digital gig-workers" che svolgono compiti che agli agenti costano troppo in termini di energia computazionale o che richiedono una responsabilità legale che le macchine non possono assumersi. In questo scenario, l'agente non è un collaboratore, ma un caporale algoritmico che detta ritmi e condizioni di lavoro a umani sempre più intercambiabili.

Il segnale debole. Quale di questi tre scenari prevarrà? Dipende dalle scelte di governance che facciamo oggi. Se trattiamo l'agente come una "black box" magica (Scenario 2) o come una proprietà esclusiva (Scenario 3), il futuro sarà distopico. Se lo progettiamo come un'infrastruttura trasparente al servizio dell'integrazione umana (Scenario 1), l'agentificazione potrà essere la più grande liberazione dal lavoro ripetitivo della storia recente.

Conclusioni

Quando abbiamo iniziato questo viaggio attraverso l'agentificazione del lavoro, partivamo dall'entusiasmo quasi febbrile del 2025 – l'anno in cui ogni settimana nasceva un nuovo strumento che prometteva rivoluzioni di produttività, efficienza miracolosa e trasformazioni radicali. Ci aspettavamo di parlare solo di tecnologia, di algoritmi, di architetture software. Invece ci siamo trovati, inaspettatamente, a confrontarci con questioni di filosofia morale, di etica applicata, di responsabilità umana. Questo non è un caso, né una deviazione dal percorso. È esattamente ciò che accade quando una tecnologia smette di essere uno strumento passivo nelle nostre mani e diventa un attore autonomo nel nostro ecosistema economico e sociale.

Gli agenti AI rappresentano qualcosa di profondamente diverso da tutto ciò che abbiamo conosciuto finora. Non sono semplicemente un miglioramento incrementale del software che utilizzavamo ieri. Sono una nuova specie di entità digitali che abitano il nostro mondo: hanno la capacità di percepire l'ambiente circostante, di prendere decisioni complesse scomponendo obiettivi in sotto-obiettivi, di agire in autonomia senza chiedere conferma a ogni passaggio, e – cosa forse più inquietante – di sbagliare in modi inattesi e di amplificare i nostri pregiudizi inconsci su scala industriale, trasformando “bias” individuali in discriminazioni sistemiche.

Se continuiamo a trattare questi sistemi come semplici utensili – come fossero martelli digitali più sofisticati – rischiamo di scivolare inesorabilmente verso gli scenari più cupi che abbiamo delineato. Pensiamo alla Rete Opaca, dove milioni di agenti interagiscono in loop incomprensibili, prendendo decisioni che nessun essere umano riesce più a ricostruire o a contestare, creando una sorta di burocrazia algoritmica infinitamente più frustrante di quella umana. Oppure al Feudalesimo Digitale, dove pochi giganti tecnologici controllano gli

agenti orchestratori più potenti, mentre la maggioranza dei lavoratori e delle piccole imprese viene relegata ai margini, ridotta a svolgere i compiti residuali che agli algoritmi costano troppo o per i quali serve un corpo fisico e una responsabilità legale che le macchine non possono assumere.

In entrambi questi scenari distopici, l'efficienza – quella promessa iniziale, quel mantra ripetuto ossessivamente – diventa in realtà una forma sottile ma pervasiva di tirannia. L'automazione non libera tempo per attività creative e significative, ma accelera il ritmo fino a renderlo insostenibile. La responsabilità non viene chiarita, ma si dissolve e si frammenta in una nebbia di codici, di log incomprensibili, di clausole contrattuali scritte per esonerare i fornitori da ogni obbligo.

La sfida della governance: rallentare per accelerare meglio

La vera sfida che sta davanti a chi guida organizzazioni, a chi prende decisioni politiche, a chi progetta i sistemi di domani non è di natura tecnologica. Non si tratta di imparare a programmare gli agenti, di scrivere codice migliore, di addestrare modelli più potenti. La sfida è imparare a governare questi sistemi – a stabilire confini chiari, a definire zone rosse dove l'automazione non può entrare, a costruire meccanismi efficaci di supervisione e controllo che non siano solo formalità burocratiche.

Questo richiede, innanzitutto, il coraggio di rallentare l'automazione quando il rischio è troppo elevato. Pensiamo al settore sanitario: quando un agente AI suggerisce una terapia, non possiamo permetterci di fidarci ciecamente solo perché "lo dice l'algoritmo". Serve un medico che comprenda davvero il contesto clinico del paziente, che sappia riconoscere quando un output apparentemente ragionevole è in realtà una correlazione statistica priva di senso

causale, che abbia l'autorità e la competenza per dire "no, in questo caso l'agente sta sbagliando, e questo è il motivo...".

Serve anche la lucidità di investire in quelle figure professionali di cui abbiamo parlato – gli Orchestratori, i Validatori, gli Eticisti degli Algoritmi – che non costruiscono gli agenti ma li governano, li mettono alla prova, li contestano quando necessario. Queste non sono posizioni ornamentali, create per dare un'impressione di responsabilità senza sostanza. Sono ruoli strategici essenziali, che fanno la differenza tra un'organizzazione che usa l'AI per potenziare le capacità umane e un'organizzazione che si fa travolgere dall'automazione senza controllo. Questo richiede soprattutto, l'intelligenza di capire che l'etica non è un freno all'innovazione o un ostacolo burocratico da superare il più rapidamente possibile. L'etica è l'unico sistema di navigazione sicuro quando ci avventuriamo in un territorio inesplorato. Senza principi etici chiari, codificati fin dalla fase di progettazione e poi verificati continuamente durante l'esecuzione, stiamo semplicemente navigando a vista in mare aperto, sperando di non schiantarci contro scogli che non possiamo vedere.

Progettare il futuro, non subirlo

L'invito finale – e questo è forse il messaggio più importante di tutto il saggio – è di non subire passivamente l'agentificazione, ma di progettarela attivamente. Il futuro del lavoro non sarà determinato da una competizione biologica tra esseri umani e macchine intelligenti, come se fossimo in una sorta di selezione naturale tecnologica dove vince il più efficiente. Sarà invece il risultato di una competizione tra visioni diverse dell'automazione: da una parte, chi usa le macchine per potenziare l'umano, per liberare tempo ed energia mentale da dedicare a ciò che davvero conta – la creatività, la relazione, la cura, il pensiero critico; dall'altra, chi le usa per sostituire l'umano, per

ridurre i costi a breve termine e per massimizzare profitti senza considerare le conseguenze sociali.

In questa partita cruciale, le competenze tecniche sono fondamentali, poiché sono loro che costruiscono il motore e che rendono possibile l'automazione. Ma sono le competenze umanistiche – la capacità di contestualizzare storicamente, di ragionare filosoficamente, di comprendere le dinamiche sociali, di applicare principi giuridici ed etici a situazioni concrete – che tengono saldamente il volante nelle nostre mani. Senza quella direzione, senza quella bussola morale, il motore più potente ci porterà semplicemente nella direzione sbagliata, più velocemente.

L'umanesimo tecnologico: una sintesi possibile

Ciò che emerge da questa analisi non è né un pessimismo apocalittico né un ottimismo ingenuo. È invece la possibilità concreta di un umanesimo tecnologico, un futuro in cui la tecnica è messa al servizio della realizzazione delle potenzialità umane più autentiche e profonde, non in competizione con esse.

In questo scenario, investire nelle competenze dell'essere, nella capacità di pensare criticamente, di relazionarsi empaticamente, di creare significato e dare senso alle esperienze e di assumersi responsabilità morali, non è un gesto di conservatorismo nostalgico e neanche un tentativo romantico di difendere un passato ormai irrecuperabile. È invece la via più radicale e necessaria per costruire un futuro in cui la tecnologia amplifichi ciò che di più prezioso abbiamo come specie: la nostra capacità di essere pienamente umani nel mondo, di costruire comunità solidali, di dare senso e direzione al nostro vivere insieme.

Gli agenti AI ci stanno offrendo un'opportunità straordinaria: quella di liberarci dai compiti ripetitivi, dalla burocrazia fine a sé stessa, dalla fatica cognitiva di gestire informazioni frammentate. Ma perché questa promessa si realizzi davvero, dobbiamo mantenere il controllo del processo. Il volante – la capacità di decidere dove vogliamo andare, quali valori vogliamo difendere, quale società vogliamo costruire – deve rimanere saldamente nelle nostre mani. Nelle mani di esseri umani competenti, responsabili ed eticamente consapevoli.

Solo così l'agentificazione potrà essere ciò che promette di essere: non un'ennesima ondata di automazione, che sostituisce lavoro umano con macchine più efficienti, ma la più grande liberazione collettiva dal lavoro ripetitivo della storia recente, un'opportunità per riscrivere il contratto sociale tra tecnologia, lavoro e dignità umana.

BIBLIOGRAFIA

Da assistenti a costruttori cosa cambia

Alessandro Longo, "AI via l'economia degli AI Agent: l'automazione vale già miliardi" in "unipd.it".

<<https://thesis.unipd.it/retrieve/a1fc566d-b322-4ca0-8637-df66e7153978/Tesi%202072876%20Cavazzana%20Elena%202025.pdf>>

McKinsey & Company. (2024). The state of AI in early 2024: Gen AI adoption spikes and starts to generate value. McKinsey Global Survey. <<https://www.mckinsey.com/capabilities/quantumblack/ourinsights/the-state-of-ai>>

McKinsey Global Institute, "The Economic Potential of Generative AI: The Next Productivity Frontier" del 2023.

<<https://www.mckinsey.com/capabilities/tech-and-ai/our-insights/the-economic-potential-of-generative-ai-the-next-productivity-frontier>>

David G. Widder and Mar Hicks, "Watching the Generative AI Hype Bubble Deflate" in "arxiv.org" del 18 agosto 2024.

<<https://arxiv.org/abs/2408.08778>>

"L'AI Agentica vs l'AI Generativa: le differenze fondamentali" in "astera.com" del 5 giugno 2025.

<<https://www.astera.com/it/type/blog/agentic-ai-vs-generative-ai/>>

Z. Xi, W. Chen, X. Guo, W. He, Y. Ding, B. Hong, M. Zhang, J. Wang, S. Jin, E. Zhou, "The rise and potential of large language model based agents: A survey," in "arxiv.com" del 14 settembre 2023.

<<https://arxiv.org/abs/2309.07864>>

Lareina Yee, Michael Chui, and Roger Roberts con Stephen Xu, "Superagency in the workplace: Empowering people to unlock AI's full potential" in "italianelfuturo.com" del 26 novembre 2025.

<<https://italianelfuturo.com/reports/mckinsey-2025-ai-lavoro-superagenzia/>>

McKinsey, "Agents for growth: Turning AI promise into Impact" in "mckinsey.com" del 3 novembre 2025.

<<https://www.mckinsey.com/capabilities/growth-marketing-and-sales/our-insights/agents-for-growth-turning-ai-promise-into-impact>>

Kate Crawford, "L'impero del calcolo" in Wired n.113 del 24 giugno 2025.

Il ruolo della governance e del controllo umano

European Union Regulation (EU), Regolamento (EU) 2024/1689 in "europa.eu".

<<https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng>>

Jessica Kelly and others, "Navigating the EU AI Act: A Methodological Approach to Compliance for Safety-critical Products" in "publicarest.fraunhofer.de" del dicembre 2024.

<<https://publicarest.fraunhofer.de/server/api/core/bitstreams/2724ce80-442b-470e-84af-5069d66ea193/content>>

Chaffer, T.J., von Goins II C., Cotlage D., Okusanya B., Goldston, J., "Decentralized Governance of Autonomous AI Agents" in "arxiv.org" del 2024.

<<https://arxiv.org/abs/2412.17114>>

Maslej N., Fattorini L., Brynjolfsson E., Etchemendy J., Ligett K., Lyons T., Manyika J., Ngo H., Niebles J.C., Parli V., Shoham Y., Wald R., Clark J., Perrault R., "Artificial Intelligence Index Report 2025" in "arxiv.org" del 2025.

<<https://arxiv.org/abs/2504.07139>>

Santoni de Sio, F., Van Den Hoven J., "Meaningful Human Control over Autonomous Systems: A Philosophical Account" in "frontiersin.org" del 2018.

<<https://www.frontiersin.org/journals/robotics-and-ai/articles/10.3389/frobt.2018.00015/full>>

Crawford K., "L'impero del calcolo", "Wired Italia" n. 113 del 24 giugno 2025.

Dove vengono usati gli Agenti?

Alberto Bozzo, "Agenti di intelligenza artificiale in ambito sanitario: applicazioni, automazione e tendenze future" in "consultingpb.com" del 3 giugno 2025.

<<https://www.consultingpb.com/blog/diritto-rovescio/agent-ai-in-sanita/>>

Preeti Anchan, "I migliori strumenti di IA per gli operatori sanitari della comunità" in "clickup.com" del 15 novembre 2025.

<<https://clickup.com/it/blog/541916/strumenti-di-ia-per-operatori-sanitari-di-comunita>>

Kees Hertogh, "Agentic AI in action: Healthcare innovation at Microsoft Ignite 2025" in "microsoft.com" del 18 novembre 2025.

<<https://www.microsoft.com/en-us/industry/blog/healthcare/2025/11/18/agent-ai-in-action-healthcare-innovation-at-microsoft-ignite-2025/>>

Pedram Hosseini, "AI in Healthcare: Key Takeaways from Menlo Ventures 2025 Consumer AI Report" in "linkedin.com" del 3 luglio 2025.

<<https://www.linkedin.com/pulse/ai-healthcare-key-takeaways-from-menlo-ventures-2025-report-hosseini-ujinc>>

Agid, "Piano triennale per l'intelligenza artificiale nella PA" del 2025.

<https://www.agid.gov.it/sites/agid/files/2025-01/Piano_Triennale_per_l_informatica_nella_PA_2024-2026_Aggiornamento_2025.pdf>

Fabrizio Davide, Pietro Torre, "Agenti AI per la PA del futuro: ecco il valore aggiunto" in "agendadigitale.eu" del 12 novembre 2024.

<<https://www.agendadigitale.eu/cittadinanza-digitale/agenti-ai-per-la-pa-del-futuro-ecco-il-valore-aggiunto/>>

Matthew Finio, Amanda Downie, "Agenti di intelligenza artificiale nel marketing" in "ibm.com" del 2025.

<<https://www.ibm.com/it-it/think/topics/ai-agents-in-marketing>>

Quali sono i problemi irrisolti?

Maurizio Carmignani, "Le Allucinazioni dell'AI aumentano: ecco il perché" in "agendadigitale.eu" del 9 giugno 2025.

<<https://www.agendadigitale.eu/cultura-digitale/le-allucinazioni-dellia-aumentano-i-nuovi-modelli-sbagliano-piu-dei-vecchi/>>

Andrea Signorelli, "Le Allucinazioni dell'Intelligenza Artificiale" in "zanichelli.it" del 2025.

<<https://ia.zanichelli.it/scoprire-intelligenza-artificiale/le-allucinazioni-dell-intelligenza-artificiale/>>

Celm Dilmegani e Aleya Daldal, "AI Hallucination: Compare top LLMs like GPT-5.2" in "aiultiple.it" del 15 Dicembre 2025.

<<https://research.aimultiple.com/ai-hallucination/>>

Luca Barbanotti, "Oltre la scatola nera: come l'intelligenza artificiale agentiva sta ridefinendo la spiegabilità" in "sas.com" del 2025.

<https://www.sas.com/it_it/news/leading-art-innovation/innovation-sparks/tecnologie/agentice-ai-oltre-la-black-box-cambia-nostro-rapporto-con-spiegabilita.html>

"L'impatto dell'intelligenza artificiale sul processo decisionale" in "digital4pro.com".

<<https://www.digital4pro.com/2025/03/18/limpatto-dellintelligenza-artificiale-sul-processo-decisionale/>>

Thseen Zia, "La trappola degli agenti di intelligenza artificiale: le modalità di errore nascoste dei sistemi autonomi per cui nessuno si sta preparando." in "unite.ai" del 13 dicembre 2025.

<<https://www.unite.ai/it/when-ai-turns-rogue-exploring-the-phenomenon-of-agentice-misalignment/>>

Chi serve per governare questi problemi?

"Artificial Intelligence Index Report" 2025 in "stanford.eu".

<https://hai.stanford.edu/assets/files/hai_ai_index_report_2025.pdf>

Santoni De Sio, F., Van Den Hoven J., "Meaningful Human Control over Autonomous Systems" in "frontiersin.org" del 2018.

<<https://www.frontiersin.org/journals/robotics-and-ai/articles/10.3389/frobt.2018.00015/full>>

Philip Drammeh, "Multi-Agent LLM Orchestration Achieves Deterministic Response Quality" in "arxiv.org" del 19 novembre 2025.

<<https://arxiv.org/abs/2511.15755>>

Natalia Campana, "What Does An AI Auditor Do?" in "freelancemap.com" dell'11 settembre 2025.

<<https://www.freelancemap.com/blog/what-does-big-data-specialist-do/>>

"Who's Responsible When AI Gets It Wrong?" in "bronson.ca" del 7 ottobre 2025.

<<https://bronson.ca/algorithmic-accountability/>>

"The Role of the AI Officer: Guide to Responsible AI Leadership" in "aiguardianapp.com" del 2025.

<<https://www.aiguardianapp.com/ai-officer-responsibilities#:~:text=Similar%20to%20a%20Chief%20Data,initiatives%20are%20developed%20and%20deployed>>

Bradley Hook, "What is an AI Ethicist and How Do I Become One?" in "bradleyhook.com"

<<https://bradleyhook.com/what-will-a-world-teeming-with-ai-agents-be-like/>>

Ripensare le competenze: dalla logica del fare alla logica dell'essere

ILO (International Labour Organization) (2023-2024). "Generative AI and Jobs: A global analysis of potential effects on job quantity and quality" in "ilo.org".

<<https://www.ilo.org/publications/generative-ai-and-jobs-global-analysis-potential-effects-job-quantity-and>>

Il ritorno dei filosofi: perché le competenze umanistiche sono la nuova ingegneria

Olojede H.T, "Arts and Humanities in the Age of Artificial Intelligence", in "philarchive.org" del 2025.

<<https://philarchive.org/archive/POLIPO>>

Yana Samuel, "Cultivation of human centered artificial intelligence: culturally adaptive thinking in education (CATE) for AI", in "frontiersin.org" del 1° dicembre 2023.

<<https://www.frontiersin.org/journals/artificial-intelligence/articles/10.3389/frai.2023.1198180/full>>

Lina Khair Rahme, Rasha Halat, "Fostering Ethical and Critical Minds in the AI Era: A Practical Approach" in "harvard.edu" del 6 novembre 2023.

<<https://mepli.gse.harvard.edu/our-fellows-at-work/fostering-ethical-and-critical-minds-in-the-ai-era-a-practical-approach/>>

Il bivio del 2030: Tre scenari per un futuro "agentificato"

UK Government Office for Science, "AI 2030 Scenarios Report" in "gov.uk" del gennaio 2024. Questo rapporto utilizza cinque incertezze critiche per costruire scenari futuri sullo sviluppo dell'AI, offrendo un framework metodologico robusto per l'analisi degli scenari possibili.

<https://assets.publishing.service.gov.uk/media/6808fc002a86d6dfb2b52772/AI_2030_Scenarios_Report.pdf>

Deloitte, "Autonomous Workforce Planning: How Agentic AI Transforms Workforce Strategy" in "deloitte.com". Il report esplora come l'AI agentiva trasforma la pianificazione della forza lavoro da modelli statici a sistemi dinamici e adattivi.

<<https://www.deloitte.com/us/en/insights/topics/talent/future-of-workforce-planning/autonomous-workforce-planning.html>>

Felten, E.W., Raj M., Seamans R., "Forecasting the Impact of Artificial Intelligence on the Labor Market" in "ssrn.com". Propone metodologie innovative per misurare e prevedere l'impatto dell'AI sul mercato del lavoro, utilizzando dati dinamici da job posting online in tempo reale.

<https://papers.ssrn.com/sol3/papers.cfm?abstract_id=4710299>



Massimo Tamiatti

Mi occupo da sempre di approfondimenti e divulgazione di tematiche riguardanti le nuove tecnologie, i nuovi lavori, le nuove competenze. Lo faccio in una maniera completamente nuova e diversa rispetto agli approcci tradizionali che si hanno sulla ricerca, perché oltre a consultare le fonti tradizionali consulto le fonti che fanno riferimento alle grandi società di consulenza strategica, che non hanno una visione statica della realtà ma dinamica, poiché si concentrano sul futuro. Utilizzo sistematicamente anche il web, alla ricerca dei contenuti più aggiornati possibili, magari attraverso webinar con testimoni privilegiati sull'argomento, oggetto dell'indagine o attraverso qualche importante influencer specializzato sul tema trattato. Evidentemente oggi, utilizzo gli strumenti dell'Intelligenza Artificiale Generativa. Non trascuro il presente, cerco di interpretare i segnali che abbiamo sotto i nostri occhi, che spesso siamo abituati a guardare e non a vedere, e che ci suggeriscono quello che potrebbe accadere. Non trascuro il passato, quindi la storia, ma in realtà sono molte le discipline che mescolo: l'economia, la statistica, la sociologia, la psicologia, l'antropologia e soprattutto uso l'immaginazione ancorandola ai fatti che accadono ogni giorno nel mondo. Considero i numeri molto importanti, ma non sono tutto, al centro ci deve sempre essere una particolare attenzione ai testimoni privilegiati delle cose che accadono e alle loro storie personali. Frontiere non ha il dono della preveggenza, ma il nostro team può fare delle ipotesi e tra queste ipotesi mi sono reso conto che è molto probabile incontrare il futuro. Diffondere la cultura dell'anticipazione aiuta le persone a comprendere cosa stia capitando nel mondo, in particolare in quello del lavoro, li aiuta a prepararsi, elaborando strategie e organizzandosi per arrivare a scelte giuste e consapevoli molto importanti per l'effetto che avranno sulle loro vite.

Profilo istituzionale

Sono stato Chief Research Officer dell'Agenzia Piemonte Lavoro, dove ho coordinato gruppi di ricercatori nell'ambito delle politiche del lavoro e della formazione professionale fino al 30 giugno 2023. Dal primo luglio sono entrato a far parte dello staff del "Laboratorio Riccardo Revelli" come socio onorario. Sono un esperto di servizi pubblici per l'impiego e attualmente sono docente del Master sulle Risorse Umane all'Università di Torino. Faccio parte del Comitato Scientifico dell'ISMEL e sono componente del Comitato Scientifico di Polis policy, Accademia di Alta Formazione. Sono dal 2022 il fondatore e direttore del sito on line di divulgazione scientifica www.frontieretecnologialavoro.com e membro dell'Associazione dei Futuristi italiani di Roberto Poli. Sono esperto nell'utilizzo di tecniche di previsione applicate ai mercati del lavoro locali, in collaborazione con la start up Skopia, specializzata negli studi di futuro e nell'anticipazione, finalizzata ad aiutare le organizzazioni nello sviluppo di strategie anticipanti a supporto di decisioni complesse in ambito di forte incertezza. Ho una laurea in lettere e filosofia, con una tesi in storia contemporanea e una in scienze politiche, entrambe conseguite presso l'Università di Torino e sono autore di numerose pubblicazioni legate al mercato del lavoro e in particolare ai nuovi scenari del mercato del lavoro, alle nuove professioni e alle nuove competenze.

Hanno collaborato con me a questo progetto: Gabriele Lamberti, Benjamin Dafku, Ivan e Federico Giacobino e Fabio Prettico.