

Intelligente e artificiale: una tecnologia organizzativa da regolare

Giorgio Gosetti

Università degli Studi di Verona

Marco Peruzzi

Università degli Studi di Verona

I processi di innovazione nel lavoro sono sempre più legati all'introduzione dell'Intelligenza artificiale. Così come la digitalizzazione, anche l'Intelligenza artificiale sta generando un forte impatto sul mondo del lavoro, in termini quantitativi, ma soprattutto qualitativi. Infatti, sono ridisegnati processi e contenuti del lavoro e di conseguenza anche le condizioni della qualità della vita lavorativa. In questo scenario diventa quanto mai necessario riflettere sulla progettazione organizzativa supportata dalle tecnologie e soprattutto sui sistemi di regolazione a diversi livelli.

Innovation processes in the workplace are increasingly driven by the introduction of Artificial intelligence. Like digitisation, artificial intelligence is having a significant impact on the world of work, not only in quantitative terms but, above all, in qualitative ones. Work processes and content are being reshaped, and so are the conditions affecting the quality of working life. In this context, it is more necessary than ever to reflect on technology-supported organisational design and, above all, on regulatory frameworks at various levels.

DOI: 10.53223/Sinappsi_2025-02-10

Citazione

Gosetti G., Peruzzi M. (2025), Intelligente e artificiale: una tecnologia organizzativa da regolare, *Sinappsi*, XV, n.2, pp.128-155

Parole chiave

Intelligenza artificiale
Organizzazione del lavoro
Tutele del lavoro

Keywords

*Artificial intelligence
Work organisation
Labour protection*

Introduzione

I processi di innovazione nel lavoro sono sempre più legati all'introduzione dell'Intelligenza artificiale (IA). Così come la digitalizzazione, anche l'IA sta generando un forte impatto sul mondo del lavoro, in termini quantitativi, ma soprattutto qualitativi. Infatti, vengono ridisegnati processi e contenuti del lavoro e, di conseguenza, anche le condizioni della qualità del lavoro e della vita lavorativa. In questo scenario è quanto mai necessario riflettere sulla progettazione organizzativa supportata dalle

tecnologie e soprattutto sui sistemi di regolazione a diversi livelli.

A partire da queste premesse il saggio si articola in due parti. Nella prima, un approccio sociologico ai temi del lavoro cerca di porre in evidenza i tratti salienti dell'IA e le implicazioni sul lavoro, con particolare attenzione agli aspetti organizzativi e di qualità della vita lavorativa. Nella seconda, improntata a un approccio giuslavoristico, è sviluppata una riflessione sull'architettura regolativa chiamata a governare il fenomeno e sul tipo e grado

Il testo è frutto di un lavoro di riflessione congiunta fra i due Autori. Ai fini formali e per le rispettive competenze disciplinari vanno comunque attribuiti a Giorgio Gosetti il paragrafo 1 e a Marco Peruzzi il paragrafo 2.

di tutele che può fornire al lavoratore, sul piano individuale e collettivo. Le due sezioni di questo contributo sono attraversate da un aspetto che le accomuna, vale a dire la prospettiva che intende leggere l'IA specificatamente come una *variabile organizzativa*. La quinta rivoluzione tecnologico/industriale va considerata specificatamente per le implicazioni che può avere sulla vita lavorativa, cogliendo le connotazioni che la differenziano dalle precedenti. Siamo infatti di fronte a una tecnologia di quinta generazione che sta producendo effetti di profondo mutamento nelle forme organizzative, nei meccanismi operativi e nei contenuti del lavoro, sconfinando in un mutamento antropologico. Agisce sul modo nel quale intendiamo il lavoro. È del tutto evidente che uno scenario di mutamento di tale portata ci richiede di considerare forme e contenuti della regolazione.

1. Digitalizzazione e Intelligenza artificiale

Iniziamo questa prima parte da una frase che apre e chiude il cerchio delle nostre riflessioni: “il rischio al quale ci espone l'IA, come ogni tecnologia ad ampio spettro, è quello di trasformare gradualmente il mondo al punto da diventare sempre più indispensabile per il suo funzionamento” (Adler 2024, 375). Proviamo a riempire di contenuto il cerchio perimetrato da questa frase, avendo sullo sfondo quel “principio di moderazione” proposto da Adler, una sorta di “forma di sussidiarietà” che si traduce nella prospettiva di “fare appello all'Intelligenza artificiale solo quando il suo contributo netto è positivo” (*Ibidem*).

Elementi per definire l'IA

In un lavoro comparso recentemente sempre su questa rivista (Canal *et al.* 2024), proponendo un'analisi a partire dai dati della V Indagine Inapp sulla Qualità del lavoro in Italia (QdL), eravamo giunti alla conclusione che il “rapporto uomo-macchina, alla base di tanta letteratura che ha studiato le condizioni lavorative, va riproblematizzato alla luce dei recenti sviluppi delle innovazioni tecnologiche”. Le analisi dei dati ci hanno posto davanti a differenti profili di lavoratori legati al diverso tipo di qualificazione e specializzazione. Infatti “l'utilizzo della tecnologia pervasiva può (...) diventare un'occasione per generare *capabilities*, per declinare concretamente un saper apprendere e un saper agire dentro i luoghi di lavoro”, proprio perché interviene sul

tipo e sul livello di qualificazione e specializzazione del lavoratore. Il contenuto del lavoro legato alla tecnologia digitale è infatti direttamente riferibile alle possibilità per il lavoratore di un'azione autonoma e di un controllo sul proprio lavoro. Fuori da ogni determinismo tecnologico, il grado di qualificazione del lavoratore appare comunque segnato dal tipo di tecnologia utilizzata: “specializzazioni in tecnologie hardware si legano a profili simili a quelli dei lavoratori non digitali e, addirittura, in alcuni casi paiono più fragili (ad esempio per la maggiore presenza di contratti non standard); al contrario, la specializzazione in tecnologie software si associa a migliori profili occupazionali. (...) I lavoratori specializzati nell'utilizzo di strumenti produttivi di tipo software beneficiano di un ritorno diretto e complessivo in termini di qualità del lavoro, rispetto a coloro non interessati da strumenti digitali. Al contrario, l'utilizzo di tecnologie hardware non si associa a effetti positivi, anzi genera ricadute negative in termini ergonomici”. Una differenza che ritroviamo anche esaminando la percezione di essere esposti a un rischio di “obsolescenza tecnologica” o di “dissonanza” fra attività svolta e aspirazioni lavorative (Canal *et al.* 2024, 81-82). Tutti elementi che rinnovano l'invito a studiare e monitorare gli effetti quantitativi (sulle dinamiche occupazionali) e qualitativi (sulla qualità del lavoro e della vita lavorativa) generati dall'introduzione delle innovazioni tecnologiche.

Il percorso che intendiamo operare in questo saggio parte da queste premesse, ma fa proprio un ulteriore presupposto: rispetto alla recente ondata di innovazioni nel mondo del lavoro legate alla digitalizzazione, siamo di fronte a una nuova fase di mutamento, che ruota e ruoterà sempre più attorno all'IA. Le tecnologie che possiamo collocare sotto l'etichetta di IA si muovono in continuità con quelle della digitalizzazione ormai matura, ma ci pare evidenzino alcuni tratti distintivi rispetto ad essa. Tratti che ci portano a parlare ormai di una quinta rivoluzione tecnologica: dopo il vapore che muove i sistemi di produzione e trasporto, l'elettricità che alimenta le catene di montaggio, l'automazione e l'informatica che generano sistemi di produzione centrati sul trattamento di informazioni, la digitalizzazione che connette e pervade i sistemi di produzione e di vita, siamo arrivati a un *nuovo mondo* segnato dall'IA. Un mutamento da affrontare in termini interdisciplinari sotto il profilo analitico

e regolatorio (Gosetti *et al.* 2024), che riguarda un passaggio antropologico profondo. Nelle prossime pagine, inizialmente partiremo da alcuni aspetti definitori e di inquadramento sociologico del tema, per poi trovare elementi di discussione sotto il profilo giuridico e della regolazione a diversi livelli. Il terreno unificante rimane quello del lavoro, nelle sue diverse dimensioni.

Nel cercare elementi per una definizione di Intelligenza artificiale, ci viene incontro quanto recentemente stabilito dall'Unione europea (European Commission 2025): *"AI system means a machine-based system that is designed to operate with varying levels of autonomy and that may exhibit adaptiveness after deployment, and that, for explicit or implicit objectives, infers, from the input it receives, how to generate outputs such as predictions, content, recommendations, or decisions that can influence physical or virtual environments"*. Partendo dalla varietà di elementi che caratterizza l'IA e che è alla base della difficoltà a produrre una definizione esaustiva di IA, le linee tracciate dall'Unione europea sostengono che *"the notion of an 'AI system' should be clearly defined while providing the flexibility to accommodate the rapid technological developments in this field. The definition of an AI system should not be applied mechanically; each system must be assessed based on its specific characteristics"*. Nello specifico, a corollario della definizione, sono individuati sette elementi basilari: *"(1) a machine-based system; (2) that is designed to operate with varying levels of autonomy; (3) that may exhibit adaptiveness after deployment; (4) and that, for explicit or implicit objectives; (5) infers, from the input it receives, how to generate outputs (6) such as predictions, content, recommendations, or decisions (7) that can influence physical or virtual environments"*. Ci riferiamo quindi a un *sistema*, che richiama immediatamente una logica di parti in forte interazione, in grado di agire a differenti livelli di *autonomia* operativa, di *adattarsi*, modulando l'impiego di risorse a seconda degli *obiettivi* da perseguire, essendo in grado di *rilevare* e *trattare* informazioni, di formulare *previsioni* a partire da esse e, soprattutto, di prendere *decisioni* immediatamente *influenti* sull'ambiente fisico e virtuale. Un sistema che possiamo analizzare guardando al suo ciclo di vita, quindi distinguendo due fasi, quella di *'pre-deployment or building'* da quella di *'post-deployment or use'*. I sette

elementi ricordati poco sopra compaiono a diversa gradazione nelle due fasi e, come appare evidente, una definizione di IA così formulata risponde anche all'esigenza dell'Unione europea di includere un'ampia gamma di sistemi e di creare le premesse per azioni di governo.

Ma agli elementi di questa prima disamina, che potremmo qualificare di ordine politico-istituzionale, possiamo affiancarne altri, ricorrendo alla letteratura di carattere sociologico che sempre più spesso in questa fase sta discutendo il tema dell'IA. Sceglieremo una lettura che ci porta nel paragrafo successivo a tracciare alcune traiettorie interpretative provvisorie. Ci pare di poter dire che, per ora, prevalgono più le domande delle risposte.

Possiamo innanzitutto partire da una considerazione generale sostenendo che *"l'AI consente alle macchine di apprendere, pensare e agire in base a processi squisitamente umani, di ordine razionale, sollevando problemi tecnici, filosofici ed etici. (...) Mettere in grado una macchina di fissare i propri obiettivi e perseguirli, monitorare i risultati delle proprie azioni e modificarli in base al comportamento imprevedibile di altre entità attive nello stesso ambiente, apprendere automaticamente attraverso l'esperienza, in modo da migliorare il proprio rendimento, comprendere il linguaggio degli esseri umani, rispondere a domande, eseguire ordini, analizzare e tradurre testi rappresentano traguardi strepitosi se si pensa che sono stati raggiunti in poche decine di anni"* (De Masi 2018, 658-659). L'idea di una macchina che opera in autonomia è spesso prevalente: *"Non è facile precisare in cosa consista tale autonomia, ma ora requisito minimo sarebbe che l'ideatore umano non sia indotto a precisare ogni fase dell'esecuzione del programma. (...) L'IA (...) non ha un rapporto diretto con l'intelligenza umana; (...) assembla, pezzo per pezzo, una cassetta degli attrezzi che s'iscrive nella continuità dei processi di automatizzazione e di amplificazione delle capacità umane"* (Adler 2024, 26-34). Si tratta quindi di *"pratiche tecniche e sociali, istituzioni e infrastrutture, politica e cultura. La ragione computazionale e il lavoro incarnato sono profondamente interconnessi: i sistemi di IA riflettono e producono relazioni sociali e comprensioni del mondo"* (Crawford 2021, 15). Seppur autonoma, l'IA è radicata, esprime un disegno di società e di governo politico dei processi (sociali, economici ecc.). Come si vedrà più oltre potremmo dire che è anche incorporata.

L'IA è strettamente legata ai dati che riesce a reperire e trattare: “Questa è una differenza sostanziale col modello industriale, dove c’era una tensione tra quantità e qualità. L’intelligenza delle macchine sintetizza tale tensione, visto che raggiunge il suo massimo potenziale qualitativo solo approssimandosi alla totalità” (Zuboff 2019, 106). Cerca di trattare dati generando “*prodotti predittivi* pensati per intuire cosa sentiremo, penseremo e faremo” (Zuboff 2019, 106). Il salto di qualità negli ultimi anni è stato quello del passaggio all’IA generativa, a sistemi in grado di produrre immagini, testi e altro: “La prima ondata di GAI (*Generative Artificial Intelligence*) si concentra soprattutto sulla possibilità di conversare servendosi di una lingua naturale. Chiamati ‘*Large Language Models*’ (LLM), questi programmi hanno già dimostrato la loro sbalorditiva capacità di eseguire una serie di compiti con performance sovrumane. (...) L’AI generativa potrebbe innescare un nuovo scarto culturale, spostando il focus sulle macchine, permettendoci di utilizzare la potenza dell’intelligenza sintetica come strumento per accelerare il progresso (...). Le AI generative scardinano i nostri luoghi comuni sui computer. Dall’alba dell’era elettronica, i computer sono stati considerati una sorta di controparte misteriosa degli esseri umani: dotati di precisione infallibile e velocità smodata, freddi e calcolatori, privi di giudizi morali e di *savoir faire* sociale. Oggi, le GAI danno spesso sfoggio di impeccabili capacità interpersonali, di empatia e compassione, per quanto per ora siano prive di esperienze effettive nel mondo reale” (Kaplan 2024, 11-13).

L’IA è quindi un sistema in grado di (a) *rilevare* dati, leggendo il contesto nel quale opera, (b) *conservare* dati, immagazzinando anche grandi quantità di informazioni senza perdere memoria di quanto avviene nel contesto, (c) *elaborare* dati, consentendo di interpretare fenomeni che avvengono nel contesto, (d) *proiettare*, connettendo passato, presente e futuro a partire dai dati a disposizione, per prefigurare scenari possibili, (e) *decidere*, passando all’azione e generando conseguenze delle azioni intraprese, (f) *monitorare*, tenendo sotto controllo gli sviluppi del processo decisionale, (g) *valutare*, producendo analisi sugli effetti del processo decisionale e generando informazioni che vengono immagazzinate e consentono di innescare nuovi processi decisionali, (h) *apprendere*, ridefinendo i propri processi

operativi a partire da elementi raccolti nel corso dell’azione e dagli esiti dei processi.

Ci pare di poter dire che il vero e proprio salto di qualità è dato dal fatto che, a un certo punto del suo percorso, l’IA è stata in grado contemporaneamente di (a) fornire risposte, (b) ‘automatizzare task’, attività altrimenti eseguite da esseri umani, (c) generare prodotti di vario tipo. L’IA generativa, in particolare, “è costituita da modelli che non operano sulla base di programmi ma che vengono addestrati con miliardi di parametri” che in questo modo imparano e assolvono una quantità elevata di funzioni. Questo non significa che sia in grado di *comprendere* completamente il contesto, di fare esperienze sensoriali dirette, sviluppare una creatività intuitiva: “l’area della creatività, dell’invenzione sorprendente, dell’apertura di nuove strade è ancora appannaggio degli esseri umani”. Secondo molti autori, quindi, l’IA vede (per ora?) un limite significativo rispetto a quella umana nella “comprensione profonda dei processi umani, nella creatività, nell’empatia” (Butera e De Michelis 2024, 27-31). Si tratta di una tecnologia che è fondata sulla quantità e qualità dei dati a disposizione, ha continuamente bisogno di dati per agire e addestrarsi, secondo una logica di *machine learning*.

Uno dei capitoli più interessanti del dibattito attuale sull’evoluzione dell’IA è quello che ruota attorno al *deep learning*, l’apprendimento profondo, legato alla simulazione di reti di neuroni che per gradi sviluppano processi di apprendimento, riconoscendo immagini e linguaggi. Coloro che sviluppano reti neurali adottano “il cervello come modello di massima per progettare sistemi intelligenti”. Più strati di neuroni virtuali dentro la rete definiscono la sua “profondità” e quindi il livello di capacità di “apprendere qualcosa del mondo che ha intorno” (Bengio 2016, 24-26; Mitchell 2022). Più in specifico potremmo dire che “un ‘grande modello del linguaggio’ (*Large Language Model*, LLM) è, essenzialmente, una *rete neurale* di grandi dimensioni”, denominata in questo modo perché costituisce “un sistema artificiale che simula la struttura del cervello”, composto da “unità di elaborazione (‘nodi’) connesse tra loro, come i neuroni; le connessioni tra i nodi simulano le sinapsi che, nel cervello, consentono la comunicazione tra i neuroni. I nodi della rete sono organizzati in ‘strati’; ciascun nodo appartiene a uno e un solo strato. Ogni rete è fatta di uno strato di input, uno di output, e

uno o più strati intermedi (detti 'nascosti')" (Marconi 2025, 101).

Il *deep learning* è un "metodo tra i molti nel campo del *machine learning*", un sottocampo dell'IA che intende le macchine capaci di imparare e apprendere dalle "proprie esperienze" (Mitchell 2022, 9). L'infrastruttura alla base è fatta di 'soluzioni hardware e software' espressamente orientate a realizzare e ottimizzare le reti neurali profonde (Ford 2022). Quello degli algoritmi è comunque un comportamento che "non ha nulla a che vedere con quello del cervello umano". Le reti neurali profonde "per il momento, (...) non hanno alcuna rappresentazione del mondo né della sua organizzazione spaziale né delle relazioni causali. (...) Un modello generativo di linguaggio produce frasi (sequenze di parole più probabili) che non hanno senso per lui. Non c'è intenzionalità nel suo testo, anche se noi lettori tendiamo ad attribuirgliene, abituati come siamo a testi che fino a oggi erano sempre stati prodotti da esseri umani dotati di intenzione" (Mézard 2024, 43).

Gli algoritmi del *deep learning*, inteso appunto come apprendimento profondo, sono in grado di 'lavorare su dati non strutturati', 'imparano da sé'. Mentre l'apprendimento automatico del *machine learning* opera a partire da dati strutturati, organizzati ed elaborati, nel caso del *deep learning* "testi e immagini sono analizzati direttamente", non serve un intervento preordinatore dell'essere umano, e le loro peculiarità sono messe a confronto con altri dati "in un progressivo affinamento. Si chiamano processi di ottimizzazione, basati su reti neurali molto complesse e stratificate" che implicano un certo costo in termini di memoria e potenza di calcolo (Perilli 2025, 64). Andando in questa direzione si arriva a quella che viene definita 'intelligenza computazionale', vale a dire la capacità di un computer di imparare a eseguire compiti associando dati disponibili e 'osservazione sperimentale', pensando a un apprendimento esperienziale che vede la macchina capace di "modificare le sue proprie istruzioni per adattarsi alle diverse situazioni" (Perilli 2025, 98). Le macchine imparano e cambiano rispondendo agli stimoli che arrivano dall'esterno, si evolvono gradualmente. In sostanza, il *deep learning* parte dal presupposto che non sia "necessario" e "opportuno dire alla macchina cosa deve fare in dettaglio" e che i sistemi "possano apprendere ad apprendere, sviluppando

una propria capacità a partire da indicazioni generiche – algoritmi – riguardanti il modo in cui esaminare i dati". Le reti neurali addestrate a partire da esempi sono "capaci di apprendere probabilisticamente, senza supervisione, a partire da dati la cui qualità è ritenuta irrilevante" (Varanini 2024, 163).

Non va comunque dimenticato, per inquadrare ulteriormente questo passaggio, che "per capire come l'IA sia fundamentalmente politica, dobbiamo andare oltre le reti neurali e il riconoscimento di modelli (*pattern recognition*) per chiederci invece, che cosa viene ottimizzato, per chi e chi è che decide" (Crawford 2021, 17). Dobbiamo quindi declinare la chiave di lettura sociologica applicata alle dinamiche dell'economia e del lavoro includendo anche variabili politiche, in grado di aiutarci a capire le micro, meso e macro-strategie che si stanno giocando attorno all'IA. Una declinazione di strategie che avviene contemporaneamente al livello dei luoghi di lavoro (micro) che stiamo continuamente ridisegnando (fino a farci pensare se per alcune attività esista ancora un 'luogo' del lavoro), dei territori (meso), sempre più a perimetrazione flessibile a seconda dello sviluppo delle catene del valore e, necessariamente, delle relazioni globali (macro). L'IA entra direttamente in quella continua stratificazione spazio-temporale che riguarda i cosiddetti processi di glocalizzazione che connotano il mutamento frequente dei tre livelli e della loro fusione continua.

Linee interpretative dell'IA

Ripercorrendo quanto abbiamo sostenuto finora, pensiamo sia possibile individuare alcune linee interpretative dell'IA in quanto tecnologia che agisce direttamente su organizzazione e contenuto del lavoro, e più specificatamente sulla qualità delle condizioni lavorative. Proveremo nelle prossime pagine a proporre alcune traiettorie partendo da *parole chiave* che ci aiutano a sintetizzare.

(a) Lavoratori

Un modo per interpretare il ruolo dell'IA è quello di andare a comprendere quanto sia in grado di modificare il lavoro umano, magari potenziandolo, generando quindi 'lavoratori arricchiti' attraverso la tecnologia (associando, ad esempio, i tradizionali *job enlargement* e *job enrichment*), oppure tenda a sostituire del tutto o in parte il lavoro delle persone. Più che ampliare le competenze e le capacità di

azione del lavoratore o magari sostituirlo, mettendo a disposizione informazioni e riconfigurando i posti di lavoro sotto il profilo organizzativo e contenutistico, l'IA è in grado di ridefinire radicalmente il lavoro, reinventarlo, ridisegnare alla base i processi di creazione del valore. Andando in questa direzione finirebbe appunto per ridefinire quello che finora abbiamo chiamato *lavoro*, inteso come attività associata a una persona impegnata nello svolgere operazioni per produrre beni o servizi dentro meccanismi operativi organizzati e in siti identificabili, sostituendolo con attività non formalizzate che generano valore (economico), ma difficilmente assimilabili appunto alla nostra tradizionale definizione di lavoro.

Dentro questa traiettoria trasformativa, per ora, si collocano quei mutamenti di organizzazione e soprattutto di contenuto del lavoro che stanno innescando un superamento della tradizionale distinzione fra colletti bianchi e colletti blu. Si stanno imponendo i cosiddetti *new collar*, differentemente colorati, che, anche quando svolgono compiti una volta associabili, ad esempio, ai colletti blu, si trovano ora a trattare informazioni come i colletti bianchi, o a essere affiancati dai – e integrati nei – sistemi di IA che annullano *alla fonte* la precisa distinzione bianco-blu, omogeneizzando colletti e divise. La diffusione dell'IA nel lavoro costituisce sicuramente un ulteriore fattore di spinta verso quella ibridazione alla base della società dei lavori (Gosetti 2024), quindi verso una combinazione di competenze tecniche e *soft skills*, contenuti scientifici riferibili alle aree tecnologiche tradizionali e alle aree umanistiche (integrando discipline), competenze di specializzazione orizzontale e verticale. Competenze che mutano nel tempo, seguendo il corso di vita dell'organizzazione e della persona. Gli aspetti tecnologici di cui ci stiamo occupando, infatti, devono anche interessare una riflessione sui passaggi generazionali e sull'adattamento delle competenze ai cicli di vita lavorativa. I nuovi contesti lavorativi richiederanno forse sempre più una capacità sistemica di lettura del proprio contesto lavorativo, una leadership legata all'IA che sa usare anche la IA per generare collaborazione e condivisione nei luoghi di lavoro. L'IA sarà sempre più orientata a supportare la selezione del personale, la formulazione della strategia organizzativa e la valutazione delle performance. Basti pensare, solo per fare un esempio, ai dispositivi che vengono

impiegati nel mondo del calcio per i processi di scouting (individuando i 'talenti'), elaborare strategie e tattiche (prima della partita e durante il suo svolgimento), studiare i vari aspetti del giocatore nel corso della prestazione sportiva e proporre decisioni consequenziali, individuare percorsi personalizzati di sviluppo delle competenze.

Nei prossimi anni il lavoro sarà diverso da quello che abbiamo visto finora e si apriranno spazi per le diversità (potremmo forse dire la neurodiversità). Dovremo studiare pertanto quanto l'IA sarà in grado di creare 'ambienti' inclusivi, di rendere 'abitabile' il lavoro nelle sue diverse dimensioni. L'IA non sarà solamente sostitutiva o integrativa/complementare/supportiva, ma trasformerà il lavoro, diventerà un generatore di lavoro diverso, che ci chiederà nuovi lavoratori: *model & prompt engineers, data curators & trainers*, addestratori delle macchine, specialisti in etica applicata alle tecnologie e così via. Essendo una tecnologia non in grado di mettere in campo l'intelligenza emotiva, è del tutto probabile che vi sia bisogno di lavoratori, o forse ancor più di lavoratrici (dato che spesso queste competenze vengono attribuite alla componente femminile), che dispongano di questo tipo di competenze, così come di creatività e capacità analitiche. Ma, per le informazioni che abbiamo, siamo sicuramente in presenza di una tecnologia che richiede ancora molto lavoro umano, non di rado precario, sottopagato, ripetitivo, che facilita i processi dell'IA.

(b) Autonomia

Spesso quando si parla di IA si cita l'autonomia come elemento distintivo (e poco sopra vi abbiamo fatto ricorso cercando una definizione di IA). L'autonomia è, ad esempio, quella che nel lavoro vede l'incontro fra robotica e IA, quando siamo alle prese con un "utensile intelligente, capace di realizzare, a seconda del momento, intenzioni differenti", assumendo appunto una "certa autonomia, nella misura in cui il suo 'comportamento' non sarebbe dettato a ogni istante da un agente umano". I robot "intelligenti" inoltre sono in grado di "adattarsi da sé a una certa varietà di missioni" (Adler 2024, 177) L'autonomia è anche quella dell'apprendimento automatico, degli algoritmi "che possono imparare cose nuove interagendo col mondo. (...) Una IA deve avere la capacità di imparare cose nuove da sola, piuttosto che seguire le istruzioni dei suoi creatori originari"

(Harari 2024, 389-390). Questo appare quindi come un tratto distintivo rispetto alle tecnologie precedenti.

L'IA è quindi una "risorsa crescente di capacità di *agire* interattiva, autonoma e spesso autoapprendente, (...) una *riserva di capacità di agire a portata di mano*", che "*inscrive il mondo*" generando artefatti che interagiscono con il mondo e lo costruiscono (Floridi 2022, 53-54). Proprio in virtù della sua autonomia l'IA *scrive* il mondo, entra nel processo di descrizione del mondo perché contamina il linguaggio e le modalità di dialogo, fattori performativi. Se, come taluni sostengono, l'uomo diventa macchina proprio mentre la macchina tende a umanizzarsi, la riflessione sull'autonomia dei sistemi di IA diventa forse uno dei capitoli principali della riflessione sociologica applicata al lavoro. Un aspetto che, legato al cambiamento nel lavoro citato poco sopra, segna un passaggio *antropologico* che richiede anche una disamina sotto il profilo *etico*. L'IA genera decisioni in maniera autonoma rispetto alla nostra mente, è in grado di procedere piano piano a imporre sistemi di decisione che le sono propri e che riteniamo convenienti da adottare perché capaci di farci arrivare velocemente in fondo al processo operativo. L'IA *raccoglie* informazioni e le *finalizza*, ma – e questo è un altro oggetto di studio socio-lavorista per i prossimi anni – è l'intervento umano che le *interpreta*, così come interpreta i risultati ottenuti dai processi *gestiti* dall'IA, e quindi, in ultima istanza, pare riappropriarsi di quella chiusura del cerchio che normalmente si tende ad attribuire a una tecnologia che decide autonomamente. La capacità generativa dell'IA, in forma *autonoma*, entra direttamente nei cicli di valorizzazione del capitale, e ci richiede un dispiego di (nuovi e/o rinnovati) strumenti concettuali e operativi per interpretare i processi di generazione del valore e le nuove forme di alienazione della persona rispetto al lavoro. Entra direttamente nella produzione delle diverse forme di capitale, economico, culturale, sociale e simbolico; entra nella relazione fra produzione e riproduzione.

(c) Controllo

L'IA ripropone al centro dell'interesse dell'analisi socio-lavorista il tema del controllo. La pervasività dell'IA, che tende ad allontanarsi dal diretto controllo umano acquisendo sempre più autonomia, condizione perché possa essere ritenuta appunto *intelligente*, pone sul tavolo il problema del

controllo da due punti di vista, dalle due direzioni che caratterizzano il (nuovo) rapporto uomo-macchina: controllo del lavoratore verso la macchina e della macchina verso il lavoratore. La macchina per conquistarsi la qualifica di *intelligente* deve sempre più allontanarsi dalle istruzioni dirette e continue dell'operatore, per adottare autonomamente comportamenti. Secondo alcuni autori, i sistemi di Intelligenza artificiale necessitano dell'intelligenza umana per strutturarsi, finendo per generare sistemi di controllo dei comportamenti umani. Le persone, nel lavoro e nella vita in generale, spesso senza soluzione di continuità, subiscono il controllo e il governo dell'Intelligenza artificiale, ma anche le soluzioni tecnologiche altamente sofisticate richiedono l'intervento umano, in quanto sono gli esseri umani "che portano a termine compiti che le macchine non possono svolgere" (Crawford 2021, 250). In ogni caso la preoccupazione di perdere il controllo è sempre presente: "una delle nostre paure nei confronti dell'Intelligenza artificiale ha a che fare con la possibilità che la nostra individualità rimanga schiacciata da algoritmi che decidono per noi. Anche se le macchine non volessero ridurci al silenzio, potrebbero volerlo fare coloro che le controllano: puoi pensarla così oppure no, ma il computer dice no, quindi taci" (Runciman 2024, 7-9).

Potremmo però anche in questo caso leggere il tema del controllo secondo la prospettiva dell'analisi della qualità del lavoro e della vita lavorativa (Gosetti 2022), quindi come possibilità per il lavoratore di condividere i processi decisionali che definiscono le condizioni lavorative, e, in primo luogo, gli aspetti di organizzazione e contenuto del lavoro. Si torna quindi al tema della partecipazione, nelle diverse declinazioni (individuale e collettiva, diretta e indiretta, informale e formale ecc.) e a differenti livelli (informazione, consultazione, co-decisione) come strategia per *generare controllo* da parte del lavoratore. L'IA può agire direttamente sulla gestione dei "margini di incertezza" (Crozier e Friedberg 1978), che è alla base delle relazioni nelle organizzazioni e costitutiva del potere organizzativo. I sostenitori della superiorità umana sulle macchine ci ricordano che "gli esseri umani sono la principale fonte dell'incertezza" e che "quando sono coinvolte delle persone, la fiducia in algoritmi complessi può creare illusioni di certezza" (Gigerenzer 2023, 13). Proprio su questi aspetti però si vanno (ri)giocando le relazioni nelle organizzazioni altamente tecnologizzate,

sull'asimmetria di controllo dei margini di incertezza organizzativa e, in ultima analisi, sulle possibilità che uno strumento potente come l'IA ha di consentire a chi governa le organizzazioni di favorire/comprimere l'autodeterminazione del lavoratore (Gosetti 2022). L'IA può essere uno strumento potente per creare e rafforzare asimmetrie nei luoghi di lavoro, e quindi favorire in diversi modi logiche di dominio.

(d) Generatività

Una distinzione che va operata è quella fra l'IA che si limita a riprodurre quello che ha immagazzinato, a riproporlo tale e quale o semmai elaborato, rispondendo a richieste di informazioni, e l'IA che genera qualcosa di nuovo, innesca produzioni di elementi non immagazzinati (tesi, immagini, musica ecc.). Il salto di qualità, secondo molti autori, è stato proprio quello di proporsi come IA di carattere *generativo* (Kaplan 2024). La generatività dell'IA applicata al lavoro diventa anche generatività in termini socio-tecnici. Quindi non solo nelle modalità dette sopra, ma in una nuova compenetrazione e socio-tecnica. Utile diventa quindi rileggere i mutamenti in atto nel lavoro nei termini che ci ha insegnato il Tavistock Institute (Gosetti 2024), osservando i processi di implementazione ed evoluzione delle tecnologie radicati entro contesti sociali di utilizzo, e come questo processo di compenetrazione debba costituire anche una postura nella progettazione della tecnologia. Dobbiamo cominciare a riflettere inoltre sulle possibilità che le macchine generative siano in grado anche di produrre autonomamente istruzioni alle quali adeguare il loro comportamento. Si tratta di riflettere quindi sulla relazione fra prevedibilità e imprevedibilità, probabilità e improbabilità dei comportamenti. La prospettiva da taluni paventata è che l'IA crei universi paralleli al nostro (Domingos 2018), che, affiancandoci, definisca percorsi di vita altamente prevedibili che prendano il sopravvento sulla nostra imprevedibilità e capacità di farci sorprendere dalle scoperte inaspettate, di essere, in definitiva, soggetti capaci di indeterminatezza.

Il carattere di generatività contribuisce a traslare l'IA verso l'intelligenza collettiva, produce elementi nuovi che influenzano insiemi collettivi di riflessione attorno all'elemento generato. Può favorire lo sviluppo di comunità e diventare quindi generatrice di comunità. La stessa informazione prodotta dall'IA diventa conoscenza nella misura in cui viene

inserita in una struttura comunitaria, e quindi 'genera' comunità. Un tema di studio interessante è quindi quello del processo di contemporanea collettivizzazione e incorporazione dell'IA, della compenetrazione a livello micro (persona), meso (gruppi e organizzazioni) e macro (sistemi di produzione e riproduzione sociale) che sta alla base del carattere generativo dell'IA.

(e) Scelta

A partire da quanto sostenuto finora, formuliamo anche un'ulteriore ipotesi interpretativa: l'IA è una tecnologia che *decide*, ma non *sceglie*, prende decisioni, ma non si relaziona con universi valoriali a partire dai quali adotta alcuni comportamenti piuttosto che altri, *scegliendo* appunto in relazione a *opzioni di valore*. La *decisione* è un processo computazionale trasferibile su sistemi che trattano dati per risolvere problemi; la *scelta*, sebbene possa centrarsi anch'essa su sistemi di calcolo che prevedono di processare informazioni, richiede un *giudizio* ancorato a valori (Perilli 2025). Harari ci dice che "mentre perseguono vari obiettivi, i computer in rete sviluppano un modello comune del mondo che li aiuta a comunicare e cooperare" (Harari 2024, 396). Le tecnologie coordinandosi producono ordine, quindi generano processi di decisione che ordinano attraverso azioni, ma non sono processi di scelta, se per scelta intendiamo l'ancoraggio dell'azione, riflessivo, preventivo e in itinere, a universi valoriali, condivisi e condivisibili. L'IA chiama direttamente in causa la dimensione etica, quindi l'aggancio delle scelte a universi di valori, anche quando si tratta di progettare e adottare l'IA nel lavoro. Più in generale vi è chi ha parlato di "design etico", "*values in design*" (Adler 2024), quando siamo in presenza di una progettazione della tecnologia che discute le implicazioni per il lavoratore, l'essere umano al lavoro, considerando che "il problema centrale dell'etica dell'IA non è di ordine teorico, ma pratico" (Adler 2024, 359). Si tratta, infatti, di verificare costantemente in che modo nella concretezza dei meccanismi operativi quotidiani l'IA agisca sui limiti e le potenzialità del lavoratore (quindi sull'autodeterminazione) e sugli universi valoriali che sono posti a riferimento della discussione su di essi/e (limiti e potenzialità).

Come spesso viene ricordato l'IA può muoversi entro un universo di pregiudizi. L'addestramento delle reti neurali, operato da esseri umani con i

loro universi socio-culturali di riferimento (valori e valutazioni), può portare con sé elementi di pregiudizio, trascinare dentro i sistemi decisionali, in maniera più o meno consapevole, stereotipi consolidati in grado di produrre discriminazione. Problema che, secondo alcuni potremmo risolvere se “dall’addestramento delle macchine si escludesse del tutto ogni decisione umana” (Spitzer 2024, 230). Argomentazione di indubbio carattere provocatorio, che ha il vantaggio però di farci di nuovo riflettere sulla necessità di trasparenza e di controllo delle responsabilità umane nella progettazione e nell’addestramento delle macchine, così come nel loro utilizzo. Anche in questo caso abbiamo sul tavolo alcuni binomi che sollecitano ulteriormente le nostre riflessioni, anche guardando al futuro: determinismo vs libertà, predeterminazione vs responsabilità.

Le macchine adottando una razionalità strumentale, potente, orientata a risolvere problemi, secondo una modalità di ragionamento deduttivo, lineare, mentre l’intelligenza umana è capace di approcciare alla realtà anche secondo una logica *abduktiva*, che rivede continuamente le proprie decisioni lungo un percorso di verifiche progressive di ipotesi. A fare la differenza è forse proprio il fatto che nel processo decisionale la componente umana inserisce un coefficiente di scelta legato all’attribuzione di valore, condizione non contemplata dalla decisionalità della macchina. La decisione dell’IA chiude il processo, anche velocemente, raggiungendo l’obiettivo prefissato, e anche quando raccoglie elementi per modificare la sua azione, quindi è allenata ad apprendere, tende a codificare un comportamento in vista di un obiettivo, mentre la scelta dell’intelligenza umana apre, asseconda l’imprevedibilità, la pone a confronto con universi valoriali di riferimento.

(f) *Habitus*

Mentre cerca di avvicinarsi il più possibile al linguaggio umano per interagire con l’essere umano, l’IA interviene sul linguaggio, modificandolo anche profondamente. La relazione lavoratore-macchina ridefinisce il linguaggio, creare nuovi linguaggi, e il linguaggio, come sappiamo da ormai tanto tempo, è uno dei principali fattori per generare valore (economico). L’IA agisce su alcuni riferimenti socio-culturali della nostra interazione umana, modifica prassi relazionali, costruisce mondo, trasforma

modalità espressive, e nel fare questo agisce sui riferimenti normativi della relazione interpersonale, entra nella relazione fra attore e sistema e nei processi a doppia strutturazione di cui spesso ci parla la sociologia. Noi strutturiamo il mondo con le nostre rel-azioni, ma agiamo dentro un mondo strutturato che (im)pone riferimenti (materiali e culturali) alle nostre rel-azioni, e così via. E l’IA entra potentemente in questo processo di continua *doppia strutturazione* sociale strutturata. Le nuove forme di generazione del valore economico che vedono l’IA attraversare immaterialità e corporeità, essere sempre più indossata, entrano a far parte dell’habitus. Il valore generato dall’IA non ha bisogno di un luogo di lavoro preciso. Si tratta di un’immaterialità deterritorializzata, che prevede però ‘luoghi’ di controllo, potenziali e reali situazioni di dominio legate alla capacità di interagire con l’IA, rendendola interlocutore. Allo stesso tempo l’IA agisce sul corpo, viene incorporata, entra nella sfera emotiva (*emotion AI*), capta sentimenti, motivazioni, atteggiamenti, usa dati per studiare le emozioni e i comportamenti. Legge le intonazioni della voce, le espressioni del volto, il linguaggio del corpo. Di fatto entra direttamente in relazione con quello che Bourdieu chiamava l’habitus, l’insieme di disposizioni durevoli e orientamenti maturati nel corso della vita, che sono alla base delle nostre azioni (intenzioni, aspirazioni ecc.). Mettendo assieme dati oggettivi, caratteristiche fisiche, dati comportamentali, informazioni sugli atteggiamenti, espressioni linguistiche e facciali, entra nel nostro *habitus*, e quindi nei nostri universi di costruzione del senso (McQuaid 2022). Dell’Intelligenza artificiale dobbiamo cogliere pertanto anche il suo aspetto materiale, *incarnato*, osservando come metta assieme risorse naturali, lavoro umano, storie personali e collettive, classificazioni (Crawford 2021). Potremmo dire che finisce per essere incorporata in un *habitus*, costruito di disposizione immateriali e materiali, di identità e vitalità alla base del nostro agire. Diventa *generativa* anche nei confronti del nostro *habitus*.

Un capitolo particolarmente interessante sul quale cimentare l’attenzione della ricerca sociologica dei prossimi anni è quello della diffusione dei cosiddetti ‘nanorobot’, componenti elettroniche di ridottissime dimensioni, che potremmo *incorporare* per diverse finalità e che di fatto agiranno sul nostro *habitus*, sulle nostre modalità di vita e di presa delle

decisioni. Componenti che tenderanno a generare un costruito in cui soluzioni tecnologiche e corpo umano fisico si fondono senza soluzione di continuità.

(g) Voice

Pur non avendo *coscienza* di quello che genera e del processo attraverso il quale genera, l'IA ha *conoscenza*, e per certi versi conserva memoria, di quello che genera, perché controlla le sue azioni, è attenta ai – e valuta i – risultati di esse. È una tecnologia che nel lavoro quotidiano, a partire dai dati che raccoglie sul comportamento adottato, è capace di riposizionarsi nei processi operativi e nel contesto più complessivo. Questo aspetto ci porta alla necessità di comprendere quali responsabilità vengono attribuite ai sistemi di IA, e come si generano responsabilità e si attribuiscono responsabilità ai vari ruoli implicati relativamente alla progettazione e al funzionamento dei sistemi di IA. Secondo alcuni l'IA di fatto ha una capacità di risolvere i problemi che non richiedono un 'libero arbitrio' e sostanzialmente rappresenterebbe un'estensione delle nostre capacità (Domingos 2018). Ma anche questi sono interrogativi sui quali dovremo riflettere nei prossimi anni. Infatti, l'IA ci riporta all'esigenza di creare contesti di confronto e di scelta condivisi. Pensare e progettare *luoghi* del lavoro, anche quando fisicamente le attività non richiedono una localizzazione precisa e visibile, che prevedano possibilità di *voice*, di agire criticamente verso le soluzioni che riguardano l'IA applicata a modelli organizzativi e contenuti del lavoro. Non devono quindi mancare campi, perimetri spazio-temporali, che in qualche modo possano rappresentare territori di senso, all'interno dei quali condividere le scelte relativamente a ragioni (perché) e modalità (come) di adottare le tecnologie dell'IA. Potremmo pensare ad *arene* dentro le quali si possa ritrovare quell'agire comunicativo, per certi versi habermasiano, che vede soggetti nella condizione di poter *discutere* l'adozione e le implicazioni dell'IA, controllando i potenziali effetti distorsivi della comunicazione. Affinché le parti in causa possano interpretare un diritto di *voice*, quindi, condividere la progettazione dell'organizzazione e dei contenuti del lavoro in presenza dell'IA, devono comunque essere preparate per farlo. Per agire una critica sono necessarie competenze e spazi nelle quali esercitarla. Già poco sopra, a questo proposito, abbiamo fatto cenno alla partecipazione

e alle diverse declinazioni nelle quali si concretizza nei luoghi di lavoro. L'esercizio del diritto di *voice*, richiede anche possibilità di verifica dell'esito delle istanze sottoposte a discussione, argomentate nelle *arene di discussione* dell'organizzazione (assemblee, riunioni, gruppi di lavoro ecc.). Quindi si tratta di una facoltà di agire nei contesti di giustificazione che discutono e legittimano le scelte organizzative.

Un aspetto che va studiato è se la stessa IA può costituire anche uno strumento di *voice*, sia per coinvolgere nella raccolta di suggerimenti, creando senso di appartenenza, sia per far emergere istanze, quindi per la voce alle persone dentro i luoghi di lavoro. Il pericolo di cadere in un 'algocrazia', un sistema di lavoro che vede gli algoritmi al governo, può essere scongiurato aumentando la conoscenza dei sistemi e la capacità di lettura critica, per non vedere la possibilità di governo, individuale e collettiva, espropriata dalla capacità tecnologica (tecnocrazia). Nelle agorà si dovrà discutere di etica e IA, del legame che intercorre fra innovazione, tecnologia e democrazia, mettendo a tema l'influenza che l'IA è in grado di esercitare sui processi sociali, sui comportamenti, ma ancor più sulle intenzioni e la maturazione delle convinzioni delle persone (sull'*habitus* citato poco sopra). Questo è un territorio nel quale ripensare anche le forme e modalità della rappresentanza collettiva, nel quale ridisegnare le modalità di relazione fra le parti e di contrattazione collettiva.

Pensiamo si possa chiudere questa prima parte mettendo al centro delle nostre riflessioni l'ipotesi che l'IA possa finire per rafforzare quella clessidra che nel (mercato del) lavoro abbiamo visto generarsi negli ultimi anni, vale a dire la polarizzazione (e quindi il distanziamento) fra una parte alta e una bassa del (mercato del) lavoro, soprattutto in termini di qualità del lavoro e della vita lavorativa. Un'ipotesi che già da prima della digitalizzazione interessa le analisi socio-lavoriste, e che anche l'avvento dell'IA pare vada confermando. A questo proposito, dobbiamo primariamente fare riferimento al "lato oscuro dell'IA", a quella sfera invisibile di lavoratori che generano dati, "lavoratori sottopagati necessari per contribuire a costruire, mantenere e testare i sistemi di Intelligenza artificiale", "microlavoratori" dediti a compiti digitali molto ripetitivi (etichettatura e caricamento di dati, revisione di contenuti ecc.) (Crawford

2021, 75). Una serie di attività a basso costo, non eseguibili dalle macchine o più convenientemente realizzabili da esseri umani, che genera un'area di sfruttamento lavorativo e alienazione. Ma l'IA, come abbiamo detto poco sopra, è anche uno strumento che potenzia significativamente le capacità del lavoratore, crea un lavoratore potenziato, in grado di sviluppare azioni 'abilitate' proprio dalla tecnologia di ultima generazione. Una tecnologia che crea nuovi lavori non necessariamente di bassa qualità, anzi. Ma dovremmo produrre ricerca proprio per individuare ipotesi analitiche da sottoporre a verifica, per cogliere le possibilità e modalità di partecipazione ai processi di progettazione delle tecnologie e quindi della qualità della vita lavorativa. Non siamo in grado di dire come si strutturi questa nuova fase di polarizzazione, come si vadano generando o consolidando asimmetrie fra datore di lavoro e lavoratore, come si possano ridurre le disuguaglianze, di genere e generazionali, in primo luogo legate alla richiesta di nuove competenze. Di fatto ci aspettano tante tappe di lavoro analitico sul campo e di confronto interdisciplinare.

Due su tutte ci paiono comunque le piste di lavoro prioritarie sulle quali avviare un confronto: (a) il tema dei valori e (b) i riflessi dell'introduzione dell'IA sulle dinamiche collettive nei luoghi di lavoro.

Già nelle pagine precedenti ci siamo soffermati sul tema dei valori. Lo abbiamo fatto nel sostenere che la macchina intelligente decide e non sceglie, quindi non ancora la propria azione a un universo valoriale di riferimento, che orienta appunto la scelta. Produce una decisione rispettando una prescrizione, un format procedurale al quale si attiene. Ma qualcuno potrebbe obiettare, giustamente, che la macchina, quando viene *addestrata a decidere*, assume i *valori dell'addestratore*. E il caso diventa ancora più intrigante sotto il profilo analitico quando l'addestramento lascia aperta alla macchina la *possibilità di apprendere in corso d'opera*. Se ipotizziamo per un momento che, come ci suggerisce Dejours, lavorare significhi "*colmare lo scarto* tra il prescritto e l'effettivo", e che "ciò che bisogna mettere in atto per colmare questo scarto non [possa] essere previsto in anticipo", è chiaro che il lavoratore deve mettere in campo competenze (e intelligenza) per attraversare il "cammino (...) tra il prescritto e l'effettivo" (Dejours 2020, 8). Questo può implicare anche mettere sotto una lente critica le prassi operative, porre in

discussione l'addestramento e l'addestratore della macchina. In ultima analisi, assumere la *critica* come valore di riferimento. Potremmo anche supporre che l'incertezza, l'instabilità dei problemi, apra spazi che non possono essere occupati dall'apprendimento e agire automatico dell'IA (Gigerenzer 2023). O, da un'altra prospettiva, che partendo proprio dalla capacità dell'IA di generare decisioni certe, traiamo spunto per pensare a scelte instabili, a produrre discontinuità di percorso e di risultato. E qui di nuovo si apre il confronto sui valori, come orientatori e mediatori sociali alla base dei processi di condivisione delle scelte di discontinuità. Un confronto che deve prendere le mosse primariamente da una riflessione sugli effetti delle decisioni che l'IA genera nel lavoro (per stare al contesto che qui stiamo privilegiando). L'attenzione della ricerca può quindi essere posta sulle diverse forme che assume il lavoro, dentro *luoghi* tradizionali o in nuove *prassi* di generazione dei risultati. Quindi, se anche le macchine agiscono a partire da valori, potremmo dire, incorporati, si tratta di capire quanto i lavoratori siano in grado di (vi siano condizioni e competenze per) agire il valore della discontinuità, dell'indeterminatezza (a livello individuale e/o collettivo). Dato che le analisi ci dicono che l'IA sta generando e rafforzando asimmetrie nei luoghi di lavoro, e quindi le decisioni prese dalle macchine intelligenti vanno a vantaggio soprattutto di alcune parti in gioco, la prima discontinuità da mettere sul tavolo della discussione potrebbe essere legata ai valori dell'uguaglianza e della libertà nel lavoro. Si tratta quindi di analizzare i processi di regolazione dell'IA e di regolazione delle applicazioni dell'IA, per comprendere come il lavoro sia un contesto di libertà e generatore di libertà.

Altro grande tema, che per certi versi si lega a quello appena argomentato, riguarda la facoltà intrinseca delle recenti innovazioni tecnologiche di agire sulla connessione, come lo sono state da sempre quelle digitali, che nella versione dell'IA paiono però in grado di farlo con un elevato livello di autonomia organizzativa. Una connessione non solo fra tecnologie e fattori della produzione, ma anche fra persone, fra individui che cercano senso attraverso (e nell') *l'agire individuale e collettivo*. *L'agire collettivo* dentro l'organizzazione si lega direttamente alla dimensione del *riconoscimento* come concetto chiave di una lettura sociologica del lavoro. Il riconoscimento è un concetto relazionale,

come possiamo apprendere da tanta letteratura sociologica a partire dagli studi di Axel Honneth. E recentemente proprio Honneth ci invita a riflettere sul ruolo del lavoro come risorsa per “contribuire in libertà, senza coazione e senza vergogna, alla prassi della formazione democratica della volontà” (Honneth 2025, 8). Un secondo filone di studio potrebbe quindi mettere a tema quanto l’IA agisca direttamente sulla formazione dell’agire collettivo dentro l’organizzazione, quanto sia in grado di generare modelli di comportamento, lasciando aperti o chiudendo spazi per il lavoro di squadra e l’azione collettiva più in generale. Significativi processi di individualizzazione in atto ormai da anni nel lavoro, associati a nuove frammentazioni delle attività, hanno indebolito le dimensioni collettive, producendo spesso isolamento del lavoratore. Un isolamento favorito anche da precise scelte organizzative, prese anche senza il diretto coinvolgimento del lavoratore o delle sue rappresentanze. L’IA è una tecnologia di organizzazione perché è una tecnologia di comunicazione, che usa un linguaggio e genera un linguaggio. E proprio la generazione di un linguaggio (non solo nelle organizzazioni lavorative, ovviamente) crea comunità di appartenenza e condizioni di esclusione. Se mettiamo sul tavolo varie ipotesi che stanno circolando relativamente alle nuove tecnologie; (i) che agiscono direttamente sul ‘senso del noi’, definiscono specificità delle comunità digitali (che condividono obiettivi e interessi) rispetto a quelle naturali (che condividono spazi e tempi) (Riva 2025), ridisegnando la tradizionale contrapposizione fra comunità e società ereditata dalla sociologia di Tönnies (1887); (ii) che generano strutture comunicative (pattern) (Esposito 2022), manipolando dati, generando percezioni della realtà, centralizzando le risorse necessarie a produrre influenza (potere) e nuove forme di controllo; che propongono nuove modalità operative in grado di agire sui concetti di intelligenza e di coscienza così come finora li abbiamo intesi (Perilli 2025); e così via.

Se tutte queste sono traiettorie di riflessione imprescindibili per la ricerca sociologica, siamo allora chiamati a studiare e comprendere come le tecnologie dell’IA ri-definiscano modelli di comportamento collettivi, che diventano confini alle scelte individuali. Le tecnologie di ultima generazione, compresa l’IA alla quale ci riferiamo in

questo saggio, attraversano dimensioni individuali e collettive. La sociologia del lavoro ha dedicato negli anni molti studi al tema della partecipazione e alle condizioni che la rendono effettiva. Come abbiamo evidenziato, l’IA, agendo direttamente sui processi decisionali, investe le condizioni che producono una partecipazione realmente incisiva, a livello individuale e collettivo. È una tecnologia che dobbiamo *problematizzare*, nel senso di tradurre in grandi argomenti sfidanti per la ricerca sul campo, proprio mettendo a tema principalmente il coinvolgimento del lavoratore, da solo e in forma associata. Se da anni ormai le analisi, divenute letteratura consolidata e acquisita per la disciplina, ci parlano di una individualizzazione dei processi lavorativi, con relativo disancoramento del lavoratore da universi di solidarietà più o meno istituzionalizzati che garantivano tutele e diritti (Castel 2004 e 2019; Sennett 2000), dobbiamo a questo punto chiederci quali effetti provocherà l’IA proprio sul coinvolgimento della persona al lavoro e nel corso della sua vita lavorativa. A tema di ricerca finiscono pertanto sia i modelli organizzativi e quindi le modalità di collaborazione a livello verticale e orizzontale che si stanno (ri)disegnando nel lavoro, sia le forme di rappresentanza che devono necessariamente essere ripensate per intercettare la forte eterogeneità che anche le innovazioni tecnologiche contribuiscono a generare nel lavoro e nella vita lavorativa.

2. Intelligenza artificiale e tutele sul lavoro: l’impatto dell’AI Act nel sistema integrato delle fonti

La (ri)costruzione delle regole di tutela del lavoratore di fronte all’uso datoriale di sistemi di IA richiede di verificare le fonti normative applicabili nell’ambito di processi organizzativi in cui possono essere veicolati, nascosti e/o amplificati poteri datoriali direttivi e di controllo, trattati dati personali, introdotte e/o esacerbate fonti di pericolo per la salute e sicurezza, la dignità, la libertà, nonché per il diritto del lavoratore a non essere discriminato (*ex multis*, Paliokaité *et. al.* 2025; Lai 2025; Treu 2024; Novella 2022; Tullini 2021; Tebano 2020). La sfida regolatoria non si muove, tuttavia, solo sul piano della mappatura dei raccordi e delle implicazioni normative: da un lato, si pone il confronto con fonti che non presentano una matrice lavoristica, come ad esempio il Regolamento sulla protezione dei dati

personali (GDPR), dall'altro, vi è la necessità di non perdere l'orizzonte della proiezione collettiva delle tutele, espressione chiave dell'impronta propria del diritto del lavoro, collegata all'applicazione del principio di eguaglianza sostanziale. Su questa sfida si inserisce ora l'impatto dell'AI Act, normativa di prodotto chiamata a tradurre il complesso equilibrio tra protezione dei diritti fondamentali e sviluppo della competitività di mercato, tra la scelta di un modello regolativo *hard* e di dettaglio e la rapida evoluzione che connota la categoria tecnologica in questione. L'indagine che segue cerca di analizzare la risposta regolativa UE sull'IA e la sua rilevanza per l'ambito del lavoro, partendo da alcuni punti di osservazione: il suo inquadramento e la sua tenuta nel contesto internazionale, i termini con cui traccia le fattispecie di riferimento, anzitutto quella di 'sistema di IA', e articola l'approccio basato sul rischio, il carattere minimo e complementare del sistema di regole predisposto e il conseguente ruolo delle Parti sociali.

Confini di regole...

Nella relazione di accompagnamento alla proposta di Regolamento sull'IA, la Commissione europea affermava che l'interesse dell'Unione è "preservare la leadership tecnologica dell'UE", nonché "tutelare la sovranità digitale dell'Unione e sfruttare gli strumenti e i poteri di regolamentazione di quest'ultima per plasmare regole e norme di portata globale" (par. 1.1. e par. 2.2)².

Come è stato sottolineato dalla dottrina, tali dichiarazioni di intenti si confrontavano con un dato di fatto: l'Europa non è un produttore di tecnologia, né possiede importanti piattaforme o sistemi di comunicazione digitale, a differenza di Cina e Stati Uniti. Il piano su cui poteva o può eventualmente competere è quello della regolamentazione,

attraverso l'innesco del c.d. 'effetto Bruxelles', già ottenuto con il GDPR (Bradford 2020)³, con conseguente affermazione del modello normativo UE di risposta all'IA come punto di riferimento a livello globale (Finocchiaro 2024). In tale ottica, si ponevano e si pongono le norme sull'applicazione extraterritoriale contenute nell'AI Act, volte a vincolare alle garanzie UE anche fornitori e deployer stabiliti in Paesi Terzi quando il sistema o modello sia immesso nel mercato europeo, ovvero quando il sistema produca un output utilizzato nell'Unione (v. art. 2, par. 1, Reg. 2024/1689; Bassini 2024).

A pochi giorni dall'entrata in applicazione delle prime disposizioni dell'AI Act, l'impressione data da alcune scelte della Commissione europea è che a livello globale si stia registrando non tanto un 'effetto Bruxelles' quanto un 'effetto *su* Bruxelles'.

L'inquadramento delle fonti UE nel più ampio contesto internazionale aveva permesso fino ad ora di rilevare come i diversi livelli di azione condividessero alcune scelte di fondo, metodologiche e di contenuto, in particolare un approccio basato sul rischio e l'individuazione, quali garanzie imprescindibili di una prospettiva antropocentrica, della trasparenza e supervisione umana (Peruzzi 2023; Iannuzzi 2025).

In tale prospettiva si sono collocati, in particolare, la Convenzione del Consiglio d'Europa (di cui fanno parte tutti i 27 Stati membri dell'UE), completata a marzo 2024 e firmata il 5 settembre 2024 (tra gli altri) dall'UE, gli Stati Uniti, il Regno Unito e Israele⁴, e il Processo di Hiroshima. Quest'ultimo, condotto dai leader del G7, ha portato, prima, il 30 ottobre 2023, all'adozione di principi-guida internazionali e di un Codice di condotta internazionale per le organizzazioni che sviluppano AI avanzata, poi, il 13 settembre 2024, alla sottoscrizione di un Piano d'azione, in cui si valorizzano, tra i vari profili, da un

2 Si veda *Proposta di Regolamento che stabilisce regole armonizzate sull'intelligenza artificiale*, COM(2021) 206 final.

3 Rispetto all'applicazione extraterritoriale del Regolamento (UE) 2016/679 sulla protezione dei dati personali, si segnala la recente vicenda che ha interessato l'applicazione cinese *DeepSeek* e il blocco della stessa da parte di molti garanti privacy europei a fronte dell'assenza di garanzie fornite circa la conformità del trattamento dei dati personali degli utenti dell'UE alla normativa europea in materia. Come rileva il Provvedimento n. 33 del 30 gennaio 2025 dell'autorità italiana, il gestore del servizio era vincolato al rispetto del GDPR in forza di quanto stabilito dall'art. 3, par. 2, secondo cui il "regolamento si applica al trattamento dei dati personali di interessati che si trovano nell'Unione, effettuato da un titolare del trattamento o da un responsabile del trattamento che non è stabilito nell'Unione, quando le attività di trattamento riguardano: a) l'offerta di beni o la prestazione di servizi ai suddetti interessati nell'Unione".

4 Convenzione quadro sull'Intelligenza artificiale, i diritti umani, la democrazia e lo stato di diritto. Le precondizioni poste per la firma da alcuni Paesi, in particolare dagli Stati Uniti, spiegano il ristretto ambito di applicazione della Convenzione, costituito dalle attività del ciclo di vita dei sistemi di Intelligenza artificiale intraprese dalle autorità pubbliche o da attori privati che agiscono per loro conto (art. 3, par. 1, lett. a).

lato, il principio di accountability, in base al quale “la responsabilità per le decisioni concernenti i rapporti di lavoro, anche qualora prese da sistemi di IA, dovrebbe rimanere in capo ai datori di lavoro”, dall’altro, l’importanza del dialogo sociale e della contrattazione collettiva a tutti i livelli⁵.

La prospettiva di una IA affidabile, sicura e priva di discriminazioni, e di uno sviluppo responsabile, basato su tutela della trasparenza e garanzia del controllo umano, erano principi che, inoltre, si ponevano alla base dell’Executive Order adottato il 30 ottobre 2023 dalla Presidenza Biden⁶.

Il quadro globale è, tuttavia, radicalmente mutato, con la revoca di quest’ultimo ordine esecutivo da parte di Trump il giorno stesso dell’insediamento alla Casa Bianca e la mancata firma di Stati Uniti e Regno Unito alla dichiarazione finale dell’AI Action Summit di Parigi dell’11 febbraio 2025 (firmata invece dalla Cina)⁷. Sempre al Summit di Parigi J.D. Vance, Vicepresidente degli Stati Uniti, ha criticato aspramente l’approccio a suo dire iper-regolativo dell’Unione: poche ore dopo, la Commissione europea ha ritirato la proposta di Direttiva in materia di responsabilità civile per danni da IA⁸, che pur era rimasta ferma da due anni, non senza critiche da parte di molti giuristi. A smentire il collegamento tra le critiche e il ritiro è stata la Commissaria europea per la sovranità tecnologica Henna Virkkunen in una intervista al *Financial Times* del 14 febbraio 2025.

Sempre l’11 febbraio, dando seguito al documento strategico *Bussola per la competitività dell’UE* adottato a gennaio, la Commissione europea ha pubblicato il Programma di lavoro 2025, fortemente incentrato sulla semplificazione regolativa, come sottolineato sul sito istituzionale. In tale sede, oltre al primo intervento c.d. *omnibus* in materia di sostenibilità, incidente, tra le altre, sulla neo-adottata Direttiva *Due diligence* 2024/1760, è annunciata una valutazione d’impatto sul

pacchetto digitale, comprensivo di GDPR e AI Act. La strategia si pone in linea con il Report di Mario Draghi *The future of European competitiveness. A competitiveness strategy for Europe* del settembre 2024 in cui, con specifico riferimento all’IA, si segnalano le lodevoli ambizioni del GDPR e del nuovo Regolamento, ma anche “la loro complessità e il rischio di sovrapposizioni e discrepanze”, tali da poter “compromettere gli sviluppi nel campo dell’IA da parte degli attori industriali dell’UE. (...) Poiché nella competizione globale sull’IA stanno già prevalendo le dinamiche ‘chi vince prende di più’, l’UE si trova ora ad affrontare un inevitabile compromesso tra tutele normative ex ante più forti per i diritti fondamentali e la sicurezza dei prodotti, e regole più leggere per promuovere gli investimenti e l’innovazione dell’UE, ad esempio attraverso il *sandboxing*, senza abbassare gli standard dei consumatori” (Report Part B, p. 79)⁹.

Se la posizione della Commissione, articolata nel Programma di lavoro, è quella di riverificare o rivedere il punto di equilibrio tra la tutela dei diritti fondamentali e le esigenze di sviluppo tecnologico e relativa competitività di mercato, il tema di un possibile intervento normativo dedicato al lavoro continua, cionondimeno, a riproporsi nel dibattito istituzionale. La prospettiva, già promossa dall’European Trade Union Confederation (ETUC) in una risoluzione del 6 dicembre 2022, era stata condivisa sia da Gualmini e Benifei all’interno del Parlamento europeo, sia da Schmit, Commissario europeo per il lavoro e i diritti sociali nella Commissione von der Leyen I. Roxana Minzatu, attuale Vicepresidente esecutiva della Commissione europea per i diritti sociali, nelle osservazioni di apertura alla riunione della Commissione interparlamentare sull’IA e il mercato del lavoro del 17 febbraio 2025 ha confermato la necessità di una riflessione su un possibile ampliamento delle tutele, a partire dal modello fornito dal capitolo sulla

5 Ministerial Declaration, G7 Labour and Employment Ministers’ Meeting in Cagliari 12-13 September 2024, *Towards, an inclusive human-centered approach for new challenges in the world of work*; Annex 1, *G7 Action Plan for a human-centered development and use of safe, secure and trustworthy AI in the World of Work*, spec. pp.16-18.

6 Executive Order 14110 of October 30, 2023, *Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence*, Federal Register/Vol. 88, No. 210/Wednesday, November 1, 2023, p. 75191.

7 *Statement on Inclusive and Sustainable Artificial Intelligence for People and the Planet*, Palais de l’Élysée, Tuesday February 11th 2025.

8 Commissione europea, *Proposta di Direttiva del Parlamento europeo e del Consiglio relativa all’adeguamento delle norme in materia di responsabilità civile extracontrattuale all’Intelligenza Artificiale (direttiva sulla responsabilità da Intelligenza Artificiale)*, Bruxelles, 28 settembre 2022, COM(2022) 496 final.

9 Si veda https://commission.europa.eu/topics/eu-competitiveness/draghi-report_en#paragraph_47059.

gestione algoritmica della Direttiva (UE) 2024/2831 in materia di lavoro mediante piattaforme digitali. Questa prospettiva ha trovato peraltro una prima declinazione nella proposta presentata dalla Commissione per l'occupazione e gli affari sociali del Parlamento europeo il 12 giugno 2025¹⁰. A ogni modo, il recepimento della direttiva piattaforme da parte degli stati membri, da effettuare entro il 2 dicembre 2026, costituirà un importante banco di prova per una verifica dei modelli normativi che meglio possono declinare nella prospettiva lavoristica le garanzie di trasparenza e sorveglianza umana rispetto all'utilizzo dei sistemi automatizzati, anche di IA.

... confini di fattispecie

Ragionare di tutele del lavoratore di fronte all'uso datoriale di sistemi di IA significa ragionare di sistema integrato di fonti regolative. È un'operazione a cui invita lo stesso AI Act nel momento in cui si qualifica come quadro normativo volto a stabilire uno standard di tutele minimo e complementare: minimo perché non osta all'introduzione di regole più favorevoli per i lavoratori a livello nazionale anche mediante contratto collettivo (art. 2, par. 11); complementare perché non deve pregiudicare le altre fonti UE o interne applicabili, essendo anzi chiamato a svolgere un ruolo funzionale, facilitante e abilitante nei confronti dei diritti e dei mezzi di ricorso esistenti (Cons. n. 9). Ciò trova esplicita conferma e declinazione nella disciplina del *deployer*, ossia di chi utilizza il sistema di IA sotto la propria autorità (art. 3, punto 4), soggetto centrale in una prospettiva d'analisi lavoristica, considerato che tale sarà, in questo quadro di regole, la (prima, ma non sempre unica, v. *infra*) qualificazione assunta da un datore di lavoro che faccia uso di sistemi di IA. Rispetto a tale figura, si chiarisce, in particolare, che l'obbligo di utilizzare il sistema seguendo le istruzioni consegnate dal fornitore non deve compromettere l'adempimento degli obblighi stabiliti da altre fonti (art. 26, par. 3). Nella (ri)costruzione dell'architettura regolativa applicabile, il raccordo è anzitutto con la disciplina del GDPR, se il sistema automatizzato implica il trattamento di dati personali, nonché, a

livello di normativa nazionale, con l'art. 4 St. Lav. se lo strumento consente di raccogliere informazioni sull'attività lavorativa dei dipendenti, nonché con l'art. 1-bis D.Lgs. 152/1997 e gli obblighi informativi ivi previsti laddove si tratti di sistema decisionale o di monitoraggio integralmente automatizzato. Altrettanto fondamentale è il raccordo con la normativa in materia di salute e sicurezza e con quella antidiscriminatoria, entrambe di matrice UE. L'orizzonte della dimensione collettiva propria dell'ottica lavoristica invita, infine, a declinare questi profili anche nella prospettiva del ruolo delle Parti sociali, nei varchi d'accesso a loro forniti dalle diverse fonti coinvolte, ma anche negli interventi che a livello di contrattazione collettiva possono sviluppare, soprattutto a sostegno di un coinvolgimento delle rappresentanze dei lavoratori nei processi di valutazione e gestione del rischio e dell'impatto dell'IA.

L'operazione richiede, d'altra parte, a monte, di individuare i confini delle fattispecie regolate dall'AI Act, per comprenderne l'ambito di applicazione materiale e l'incidenza sul quadro normativo delle tutele.

Si è già potuto evidenziare nei paragrafi precedenti gli elementi costitutivi della definizione di sistema di IA prevista dall'art. 3, n. 1 (v. *supra*).

Alla luce degli orientamenti elaborati dalla Commissione europea a norma dell'art. 96, preme qui segnalare, anzitutto, come l'elemento dell'autonomia, qualificato dalla presenza di un certo grado di indipendenza dall'intervento umano, escluda dalla definizione i sistemi progettati per eseguire tutti i compiti attraverso funzioni azionate manualmente. Nel chiarire che l'autonomia non deve essere necessariamente completa, la Commissione spiega come a rilevare come dato-chiave sia la capacità del sistema di generare da solo l'output con cui interagisce con l'ambiente esterno, potendo l'input, in ipotesi, essere anche inserito manualmente. Al riguardo, è importante sottolineare – per non cadere in un'erronea sovrapposizione tra autonomia e apprendimento automatico, da un lato, autonomia e processo decisionale integralmente automatizzato, d'altro – che la suddetta capacità può

¹⁰ Progetto di Relazione recante raccomandazioni alla Commissione sulla digitalizzazione, l'Intelligenza artificiale e la gestione algoritmica sul luogo di lavoro – plasmare il futuro del lavoro (2025/2080(INL)). Relatore: Andrzej Buła. Le raccomandazioni alla Commissione indicano, in particolare, il contenuto di una proposta di direttiva sulla gestione algoritmica sul luogo di lavoro, con ambito di applicazione esteso a tutti i lavoratori e datori di lavoro nell'Unione, nonché ai lavoratori autonomi individuali e ai relativi committenti di servizi.

basarsi sull'applicazione di approcci sia di *machine learning*, sia *logic-based* e l'output può consistere non solo in una decisione (e in tal caso avremo in effetti un sistema decisionale integralmente automatizzato ai sensi dell'art. 22 GDPR e dell'art. 1-bis D.Lgs. 192/1997¹¹), ma anche in una previsione (stima di valori sconosciuti a partire da valori noti), in contenuti (nel caso di sistemi di IA generativi) o in mere raccomandazioni (non applicate automaticamente). A confermare che un sistema di IA, anche ad alto rischio, possa solo sostenere – e non prendere – le decisioni sono le disposizioni dell'AI Act che obbligano il *deployer* a determinate garanzie di trasparenza laddove il sistema adotti o *assista* nell'adozione di decisioni (art. 26, par. 11), obblighi di informazione che devono includere l'indicazione del diritto della persona interessata, oggetto di una decisione adottata dal *deployer* sulla base dell'output della macchina, alla spiegazione dei

singoli processi decisionali (art. 86; v. Cons. n. 93)¹². Si deve, d'altra parte, sottolineare che l'assenza di un impatto sostanziale sull'esito del processo decisionale, sia esso umano o automatizzato, permette di escludere a norma dell'art. 6, par. 3, la qualificazione ad alto rischio dei sistemi di IA elencati nell'allegato III, tra cui quelli utilizzati per la selezione e gestione del personale, salvo non effettuino profilazione di persone fisiche. Alla luce di quanto esposto, si comprende perché la fattispecie “sistema di IA” non possa essere completamente sussunta nella categoria normativa di “sistema decisionale o di monitoraggio integralmente automatizzato”¹³, prevista ad esempio nell'art. 1-bis del D.Lgs. 152/97. In tal senso si rileva una criticità nella formulazione dell'art. 11 del D.D.L. governativo in materia di IA, approvato in via definitiva dal Senato il 17 settembre 2025, laddove sovrapponendo le due fattispecie (e dimenticando, peraltro, le rappresentanze sindacali)

11 Come spiegato da Trib. Palermo, 20 giugno 2023, sono da considerarsi sistemi di monitoraggio/decisionali integralmente automatizzati, ai sensi dell'art. 1-bis, quelli “che non prevedono l'intervento umano contemporaneo o successivo rispetto alla decisione del sistema o al monitoraggio dei lavoratori, che, poi, di norma, è finalizzato proprio all'assunzione da parte del sistema di una decisione che li riguarda. Sono, quindi, in quest'ambito, integralmente automatizzati i sistemi che non prevedono l'intervento umano nella fase finale della decisione o del monitoraggio, a prescindere da un eventuale intervento dell'uomo nelle fasi antecedenti, quale quella di mero inserimento di dati, comunque elaborati”. In tal senso anche Trib. Torino, 5 agosto 2023: “un sistema decisionale integralmente automatizzato è quello che adotta decisioni senza richiedere l'intervento umano; (...) la circostanza che un intervento umano si possa rendere necessario per ovviare a problematiche che non attengono ad un percorso fisiologico, denota come l'intervento umano sia meramente eventuale, accessorio e, dunque, irrilevante”. Rispetto all'art. 22 GDPR, si può richiamare la pronuncia della Corte di giustizia dell'Unione europea (CGUE), 7 dicembre 2023, C634/21, *Schufa*, ECLI:EU:C:2023:957, in cui si spiega che la nozione di ‘decisione’ a cui rinvia l'articolo non è definita dal Regolamento; a fronte della formulazione della disposizione, la stessa deve riferirsi “non solo ad atti che producono effetti giuridici riguardanti il soggetto di cui trattasi, ma anche ad atti che incidono significativamente su di esso in modo analogo”. In tal senso, un determinato outcome, quale ad esempio un punteggio di solvibilità, rientra nel concetto di “decisione completamente automatizzata”, qualora incida in modo decisivo sulle scelte assunte, in ipotesi anche da un terzo, in relazione a un rapporto contrattuale. Nella più recente pronuncia del 27 febbraio 2025 sul caso *CK* (C-203/22), la Corte ha altresì chiarito che il diritto di informazione relativo ai processi decisionali automatizzati previsto dal GDPR copre “qualsiasi informazione pertinente relativa alla procedura e ai principi di utilizzo, con mezzi automatizzati, di dati personali al fine di ottenerne un determinato risultato”, a cui si devono aggiungere le informazioni relative all'importanza e alle conseguenze previste del trattamento controverso per l'interessato. Il diritto di ottenere “informazioni significative sulla logica utilizzata” in un processo decisionale automatizzato deve essere inteso come un diritto alla spiegazione della procedura e dei principi concretamente applicati per utilizzare, con mezzi automatizzati, i dati personali dell'interessato al fine di ottenerne un risultato specifico. “Non può soddisfare tali requisiti né la semplice comunicazione di una formula matematica complessa, come un algoritmo, né la descrizione dettagliata di tutte le fasi di un processo decisionale automatizzato, in quanto nessuno di questi metodi costituirebbe una spiegazione sufficientemente concisa e comprensibile” (par. 59). La ponderazione tra diritto di accesso e tutela del segreto commerciale deve essere effettuata caso per caso dall'autorità di controllo o dal giudice competente. Uno Stato membro, quindi, non può prevedere una disposizione che escluda, di regola, il diritto di accesso dell'interessato, previsto all'articolo 15 del GDPR, qualora tale accesso comprometta un segreto commerciale o un segreto aziendale del titolare del trattamento o di un terzo, stabilendo cioè in modo definitivo, a priori, l'esito della ponderazione dei diritti e degli interessi in gioco.

12 Rispetto all'obbligo informativo nei confronti anche dei rappresentanti dei lavoratori, da adempiere prima della messa in servizio o utilizzo del sistema di IA nell'ambiente di lavoro, cfr. art. 26, par. 7.

13 Per una definizione di sistemi di monitoraggio automatizzati, cfr. la Direttiva (UE) 2024/2831 sul lavoro mediante piattaforme digitali, in cui tali sistemi sono qualificati come “sistemi utilizzati per effettuare, o che sostengono, il monitoraggio, la supervisione o la valutazione, tramite strumenti elettronici, dell'esecuzione del lavoro delle persone che svolgono un lavoro mediante piattaforme digitali o delle attività svolte nell'ambiente di lavoro, anche raccogliendo dati personali” (art. 1, par. 1, lett. h).

stabilisce un obbligo di informazione del lavoratore dell'utilizzo dell'IA "nei casi e con le modalità di cui all'art. 1-bis" (art. 11, comma 2).

Gli orientamenti della Commissione confermano altresì che un sistema di IA, per essere definito tale ai fini del Regolamento, non debba necessariamente presentare adattabilità dopo la diffusione, intesa come capacità di autoapprendimento dopo l'installazione e durante l'uso, con possibile modifica dei pesi con cui le diverse variabili influenzano l'output e conseguente comportamento auto-evolutivo della macchina. Tale caratteristica può, invece, rilevare ai fini della categorizzazione del rischio: in forza dell'art. 6, par. 1, dell'AI Act, infatti, è considerato ad alto rischio un sistema destinato ad essere utilizzato come componente di sicurezza di un prodotto o che è esso stesso prodotto e che ai sensi della normativa di armonizzazione UE elencata nell'allegato I, tra cui rientra il nuovo Regolamento Macchine, debba essere sottoposto a una valutazione di conformità da parte di terzi. Tale è, in particolare, stando al citato Regolamento Macchine, un sistema con comportamento integralmente o parzialmente auto-evolutivo, quindi dotato di adattabilità durante l'uso, volto a garantire una funzione di sicurezza in un macchinario. Si pensi a un sistema di IA che rileva la presenza di operatori od ostacoli vicino a una pressa o a un veicolo industriale/robot che blocca/adatta il movimento o i parametri di funzionamento in caso di pericolo o per evitare collisioni; a un sistema di IA che adatta il funzionamento di un escavatore al terreno per ridurre o eliminare alcuni rischi (ad esempio di vibrazione); a un sistema di manutenzione predittiva che identifica guasti imminenti; a un sistema che modifica forza, velocità, ritmo di movimento di un braccio robotico in un impianto di produzione in base ai dati che raccoglie in tempo reale. Solo in presenza di adattabilità durante l'uso, tali sistemi, pur rientrando nella definizione di sistemi di IA, saranno sottoposti al principale quadro di garanzie previste dall'AI Act, dedicato ai sistemi ad alto rischio.

Ancora, gli orientamenti della Commissione ribadiscono come ai fini della definizione di sistema IA non rilevi il tipo di approccio che consente l'inferenza, quanto la capacità di inferenza stessa. La tecnica che permette di dedurre l'output può essere, pertanto, sia di ragionamento automatico (*logic-based*), sia di apprendimento automatico. Lo stesso *machine learning* si può articolare in varie tipologie, dal supervisionato al non supervisionato, all'auto-supervisionato, come sottocategoria del secondo, fino all'auto-apprendimento per rinforzo.

Ciò che conta, per la Commissione, è che il sistema sia in grado, nel processo inferenziale, di gestire relazioni e modelli complessi nei dati.

A rimanere esclusi dalla definizione sono i sistemi "che non trascendono l'elaborazione dati di base", come quelli che consentono di migliorare la performance computazionale, ma non di adeguare il modello decisionale; le applicazioni software standard per fogli di calcolo (privi di applicazioni di IA) o i sistemi di gestione di database che permettono di ordinare o filtrare i dati in base a criteri specifici; ancora, i sistemi di previsione semplici, di analisi descrittiva o di stima statica.

L'elenco esemplificativo degli esclusi non consente, d'altra parte, di segnare il confine della fattispecie. Del resto, il confine non può che rimanere poroso, se a guidarne la demarcazione è il parametro della complessità. È la stessa Commissione a segnalare come l'ampia ed elastica definizione in esame non permetta di determinare automaticamente o di elencare in modo esaustivo i sistemi che vi rientrano: la qualificazione richiede una verifica caso per caso¹⁴.

Certo: se il criterio della complessità non ci consente di predeterminare in modo automatico l'applicabilità della definizione, è proprio la complessità del fenomeno da regolare che spiega la necessità dell'approccio basato sul rischio, come tecnica normativa volta a garantire l'effettività dei diritti fondamentali (Loi 2001).

Sul punto, tuttavia, è necessaria una precisazione.

¹⁴ Le problematiche derivanti da una definizione molto ampia e dai confini porosi sono state segnalate con forza dai partecipanti alla consultazione pubblica sulle linee guida. Nel report di sintesi pubblicato dalla Commissione, si sottolinea, in particolare che "Le risposte indicano un ampio consenso sulla necessità di approfondire ulteriormente le definizioni di 'adattività', 'inferenza' e 'autonomia'. Molti portatori di interesse hanno espresso preoccupazione per il fatto che le definizioni attuali potrebbero includere, inavvertitamente, anche sistemi software tradizionali, come quelli basati su regole o modelli statistici di base, che non presentano le caratteristiche proprie dell'Intelligenza artificiale", Commissione europea (2025, 21).

Come evidenzia la Commissione nei propri orientamenti sulla definizione di sistema di IA, l'approccio *risk-based* trova implementazione, all'interno del Regolamento sull'IA, a partire da una distinzione per categorie di rischio (inaccettabile, alto e non-alto), a cui corrispondono previsioni di divieto ovvero l'obbligo di rispettare il principale *corpus* di prescrizioni stabilite dal Capo III¹⁵. La tecnica di giuridificazione del rischio utilizzata è, d'altra parte, costituita dall'individuazione di insiemi di pratiche d'uso e di ipotesi di deroga tipizzate, costruite a livello normativo con enunciati troppo generici o di dubbia interpretazione, tali da rendere ancora una volta le fattispecie di non chiara demarcazione (Mantelero e Peruzzi 2024). Gli orientamenti sulle pratiche vietate hanno certamente sciolto alcuni interrogativi, come quello sui sistemi di inferenza delle emozioni, sorto a fronte di un contrasto di lettura tra quanto indicato nel Cons. n. 18 e quanto riportato sulla pagina istituzionale della Commissione (ora è chiaro che i sistemi che rilevano il dolore o la stanchezza non inferiscono emozioni). I 434 paragrafi in cui si articolano le linee guida dedicate (solo) all'art. 5 sono, d'altra parte, indicativi della problematicità della lettura e del tracciamento degli ambiti di applicazione delle regole. A questi orientamenti si aggiungeranno quelli – attesi, a norma dell'art. 6, par. 5, entro il 2 febbraio 2026 – relativi all'applicazione delle deroghe alla classificazione ad alto rischio dei sistemi elencati nell'allegato III. Tali indicazioni saranno tanto più rilevanti se si considera che il fornitore, in molti casi, come quello dei sistemi utilizzati per la selezione e la gestione del personale,

effettua la classificazione sulla base di un processo interno di self-assessment, dovendo sì documentare la valutazione, ma non anche depositare tale documentazione al momento della registrazione del sistema nella banca dati UE, posto che a essere inserita, tra le informazioni da allegare, è solo l'indicazione delle condizioni che, in base all'art. 6, par. 3, hanno consentito di non ritenere detto sistema ad alto rischio. La documentazione è conservata internamente dal fornitore e messa a disposizione delle autorità competenti, su richiesta. L'autorità di vigilanza del mercato (per l'Italia, l'Autorità per la cybersicurezza nazionale, secondo quanto previsto dal D.D.L. governativo¹⁶), qualora abbia motivi sufficienti per ritenere che la classificazione non sia corretta, effettua una valutazione del sistema ed eventualmente richiede al fornitore di adottare le misure necessarie per garantire la sua conformità (art. 80).

Come già segnalato, a norma dell'art. 6, par. 3, il rischio significativo di danno – cui segue la classificazione ad alto rischio – è ritenuto escluso laddove il sistema non influenzi materialmente il risultato del processo decisionale, ovvero sempre presente qualora il sistema effettui profilazione di persone fisiche.

La tipizzazione delle condizioni che escludono il rischio, effettuata dalla disposizione, solleva non pochi dubbi interpretativi, soprattutto in rapporto ad alcune aree critiche d'uso elencate dall'allegato III. Il rischio significativo di danno sarebbe, ad esempio, da escludere, qualora il sistema fosse solo volto a rilevare lo scostamento di una valutazione umana

15 Per i sistemi non classificabili come ad alto rischio trovano applicazione, se del caso, le garanzie di trasparenza di cui al Capo IV ovvero può essere prevista l'applicazione volontaria delle garanzie del Capo III tramite codici di condotta (elaborabili anche da parte di organizzazioni di *deployer*, quindi anche di rappresentanza datoriale, insieme con rappresentanti dei portatori di interessi, quindi anche sindacati, a norma dell'art. 95). Come chiarisce il Cons. n. 166, "è importante che i sistemi di IA collegati a prodotti che non sono ad alto rischio in conformità del presente regolamento e che pertanto non sono tenuti a rispettare i requisiti stabiliti per i sistemi di IA ad alto rischio siano comunque sicuri al momento dell'immissione sul mercato o della messa in servizio. Per contribuire a tale obiettivo, sarebbe opportuno applicare come rete di sicurezza il Regolamento (UE) 2023/988" relativo alla sicurezza generale dei prodotti.

16 La scelta di affidare il ruolo a una agenzia soggetta alle funzioni di indirizzo e coordinamento della Presidenza del Consiglio, ovvero, come avvenuto in Spagna, a una nuova agenzia di vigilanza ad hoc, comunque presieduta dal titolare della Segreteria di Stato per la digitalizzazione, è consentita dall'AI Act che non richiede il carattere indipendente dell'*authority*, ma solo che eserciti i poteri "in modo indipendente, imparziale e senza pregiudizi, in modo da salvaguardare i principi di obiettività delle loro attività e dei loro compiti" (art. 70). Sul punto, cfr. le considerazioni di Falletta e Marsano (2024). Come spiegano Schmidl e Rohner (2025, 1025), la differenza tra la "completa indipendenza" richiesta dal GDPR alle relative *authority* (e che anche il PE aveva invero inserito pur solo nel considerando n. 77 nel suo testo di compromesso dell'AI Act) e l'indipendenza prevista dal Regolamento sull'IA è sostanziale: "mentre la completa indipendenza delle autorità per la protezione dei dati richiede che esse possano 'scegliere e disporre di personale proprio', soggetto solo alla 'direzione esclusiva del membro o dei membri' dell'autorità per la protezione dei dati (cfr. articolo 52, paragrafo 5, del GDPR), e che debbano disporre di un 'bilancio annuale separato e pubblico' (articolo 52, paragrafo 6, del GDPR), requisiti analoghi non si trovano nella legge sull'AI".

già completata rispetto a uno schema prestabilito, potendo procedere a una sua sostituzione o determinando un condizionamento della stessa solo con adeguata revisione umana. L'indicazione interroga sui termini con cui la presenza della revisione umana in un processo decisionale automatizzato sia in grado di escludere l'alto rischio (quando, cioè, rilevi e si possa ritenere di misura adeguata).

Parimenti il rischio è escluso se il sistema di IA esegue un compito procedurale limitato (ad esempio classificazione di documenti per categorie), migliora il risultato di un'attività umana precedentemente completata (ad esempio allineamento del testo di un documento già redatto a un determinato stile), ovvero esegue un compito solo preparatorio per una valutazione pertinente ai casi d'uso di cui all'allegato III, quindi ad esempio ai fini della selezione o gestione del personale. Il Cons. n. 53 segnala, tra i possibili compiti preparatori, "soluzioni intelligenti per la gestione dei fascicoli, che comprendono varie funzioni quali (...) il collegamento dei dati ad altre fonti di dati". La delimitazione della fattispecie andrebbe però meglio chiarita, perché il generico enunciato normativo non consente di capire se o quando, ad esempio, uno *screening* di candidature effettuato da un sistema di IA, indicato tra le finalità d'uso critiche dell'allegato III (sistemi utilizzati "per filtrare le candidature e valutare i candidati", ma anche sistemi "destinati a essere utilizzati per adottare decisioni riguardanti le condizioni dei rapporti di lavoro, la promozione"), possa integrare la condizione di deroga alla classificazione ad alto rischio, in quanto compito prodromico a restringere la rosa di candidati da sottoporre a una valutazione umana. Il tema è tanto più rilevante considerati i rischi di discriminazione che tali processi di preselezione possono implicare. Si pensi a un sistema addestrato con tecniche di apprendimento supervisionato sulla base di un set di dati di training non sufficientemente rappresentativi: imparerà a penalizzare i candidati (o le candidate...) che non presentano i pattern ricorrenti nel gruppo che risulta sovrarappresentato e, quindi, agli 'occhi' della macchina, statisticamente più adatto all'obiettivo (c.d. *proxy discrimination*); parimenti, una discriminazione (indiretta) potrebbe annidarsi nell'applicazione di tecniche di apprendimento supervisionato (per addestrare un modello funzionale, ad esempio, al ranking di candidati a una

promozione o a un bonus) in cui la selezione delle variabili indipendenti collegate alla variabile target (la c.d. etichetta) includa fattori apparentemente neutri, ma ad impatto differenziato, come la tipologia del contratto (full-time/part-time) o il numero di giorni di assenza (Gosetti *et al.* 2024). Ad ogni modo, quantomeno nell'esempio summenzionato, pare difficile escludere che il processo, riducendo l'ambito della scelta, non influenzi materialmente la decisione umana finale. Soprattutto, a risolvere a monte molte questioni interpretative poste dalla costruzione delle ipotesi di deroga dovrebbe essere l'applicabilità dell'ipotesi di conferma del rischio sempre stabilita dal par. 3, ossia la profilazione di persone fisiche. Per la definizione, l'AI Act rinvia a quanto stabilito nel GDPR, che qualifica la fattispecie come un trattamento automatizzato finalizzato a valutare, ai fini di analisi e/o previsioni, aspetti e caratteristiche personali individuali, come il rendimento, l'ubicazione, l'affidabilità, il comportamento. Essendo molto probabile che tale condizione ricorra in caso di utilizzo di processi automatizzati di monitoraggio o decisionali sul lavoro, può ritenersi tendenzialmente da escludere una classificazione come 'a basso rischio' dei sistemi destinati a essere utilizzati nell'esercizio di prerogative datoriali nei confronti dei lavoratori.

Necessità e declinazione dell'approccio risk-based. Gli obblighi di valutazione e gestione del rischio del datore di lavoro nell'AI Act, come employer e fornitore

Come messo in luce dalla dottrina, la tecnica normativa dell'approccio basato sul rischio, nel tradurre "una prospettiva procedurale del diritto", rappresenta una "modalità di reazione del diritto ad un elevato grado di differenziazione" e di complessità del reale, tale da impedire una predeterminazione rigida e standard dei contenuti delle misure di tutela (Loi 2001, 47-48). Si tratta, pertanto, di un modello regolatorio di risposta all'inadeguatezza che può presentare la regolazione diretta per misure di dettaglio 'nominate' di fronte alla dinamicità evolutiva dei fattori di pericolo e alla modularità degli elementi che li caratterizzano.

Se, come insegna il diritto della sicurezza sul lavoro, la tecnologia è sempre stata una determinante chiave di questa complessità, nel caso dell'IA, soprattutto del *machine learning*, la problematica si intensifica e acuisce, perché le peculiarità adattive e auto-evolutive della

fonte di potenziale pericolo, nonché la variabilità della sua possibile autonomia d'azione fanno dipendere l'effettività delle tutele da un necessario processo iterativo di valutazione e gestione dei rischi. Sempre le caratteristiche dell'IA – ossia la dipendenza dai dati per l'applicazione delle tecniche di inferenza, l'opacità del processo di ottenimento dell'output e la capacità di quest'ultimo di influenzare gli ambienti fisici e virtuali – impongono che il modello di controllo del rischio si muova sul duplice piano della governance dei dati e della verifica periodica di impatto dei processi automatizzati. Ciò è tanto più vero all'aumentare della profondità delle reti neurali e dell'apprendimento automatico (profondità che sta alla base, ad esempio, dello sviluppo dei modelli di GPAI¹⁷), considerato che in tal caso l'investigazione può realizzarsi solo attraverso un'osservazione dall'esterno del modello addestrato, nel suo funzionamento e nei risultati che consegna (Italiano *et al.* 2022).

La scelta di metodo, nella costruzione del sistema integrato delle fonti, si confronta, d'altra parte, nella prospettiva lavoristica con la necessità di "delimitazione di zone di indisponibilità, costituite dai diritti sociali fondamentali" (Loi 2001, 48). In altre parole, se il paradigma procedurale dell'approccio *risk-based* si presenta funzionale a garantire l'effettività dei diritti nel descritto scenario di complessità, la definizione di tali diritti deve essere presidiata a livello di fonti legislative e costituzionali, in modo che solo queste "rappresentino l'adeguato spazio giuridico in cui effettuare il bilanciamento degli interessi coinvolti" (Loi 2001, 68-69). È in tale spazio che deve essere, cioè, individuato il grado di tutela dei diritti, venendo affidata all'attore privato, titolare dell'obbligo di valutazione e gestione, solo la scelta delle misure tecniche e organizzative, ossia dei metodi e degli strumenti, che meglio lo garantiscono nello specifico contesto d'incidenza della fonte di pericolo.

Il modello previsto dalla Direttiva 89/391/CEE, concernente la salute e sicurezza durante il

lavoro, rappresenta l'esempio di matrice lavoristica della convivenza di questi due paradigmi (Loi 2001). È un esempio che salda il grado di tutela dei diritti a un principio di massima sicurezza tecnologica¹⁸, sganciandolo "da considerazioni di carattere puramente economico" (Cons. n. 13), e rispetta la necessaria proiezione collettiva delle tutele richiesta dalla materia, prevedendo il coinvolgimento di rappresentanze specializzate nel processo di valutazione e gestione del rischio (v. *infra*). Significativamente, all'interno della direttiva 89/391 la combinazione dei due paradigmi incide sulla stessa declinazione dell'approccio basato sul rischio, essendo questo informato dalla prospettiva della c.d. prevenzione primaria, ossia di un modello di tutela dei diritti fondamentali che incorpora "un bilanciamento già intervenuto per opera del Legislatore", prevedendo "la tensione alla radicale eliminazione del rischio" e "la copertura surrogatoria in senso limitativo o riduttivo della dose di rischio ineliminabile" (Buoso 2020, 48-49). In tal senso, il paradigma procedurale è segnato da una precisa sequenza di obiettivi: prima l'eliminazione del rischio; in subordine, la riduzione al minimo dei rischi ineliminabili (Lazzari e Pascucci 2024).

Si è già potuto in altra sede riflettere su quanto l'importanza della grammatica dei diritti fondamentali trovi conferma anche nel contesto dei due principali plessi regolativi coinvolti nell'integrazione tra fonti (AI Act e GDPR) e quanto la declinazione del paradigma procedurale *risk-based* all'interno di tali sistemi possa trovare allineamento con lo standard tracciato dal quadro prevenzionistico (Peruzzi 2024; cfr. al riguardo anche le diverse riflessioni di Gaudio 2024 e Treu 2025).

Pare, d'altra parte, utile tornare sull'incidenza dell'approccio basato sul rischio sulla figura del *deployer*-datore nel contesto normativo dell'AI Act. In forza dell'art. 26, tale soggetto è tenuto a dare applicazione alle misure di sorveglianza umana indicate dal fornitore, finalizzate espressamente, ai

17 Cfr. Orofino (2025). Come evidenzia la dottrina, "il successo delle reti neurali profonde nel riconoscimento delle immagini ha condotto a estenderne l'uso anche ai compiti linguistici" ed "è proprio nell'ambito linguistico che le reti neurali profonde hanno consentito di ottenere i risultati più significativi". Si fa riferimento a "i grandi modelli linguistici (*Large Language Model* o LLM)" ossia "enormi reti neurali pre-addestrate, la cui realizzazione è stata resa possibile dalla combinazione di tre fattori: enormi raccolte di testi da utilizzare per l'addestramento; altissima potenza di calcolo; nuove tecnologie per la realizzazione delle reti". "In generale, il funzionamento degli LLM resta ancora assai misterioso nel senso che non esiste una precisa teoria che spieghi perché tali sistemi forniscano le sorprendenti prestazioni di cui sono capaci" (Sartor *et al.* 2024, 252-256).

18 Rispetto alla possibilità di collocare, nell'ambito di operatività di tale principio, l'introduzione della tecnologia avanzata, nella forma di Intelligenza artificiale e robotica, cfr. di recente Faioli (2025).

sensi dell'art. 14, in sequenza, a prevenire o ridurre al minimo i rischi per la salute e sicurezza o i diritti fondamentali; ed è altresì chiamato a monitorare debitamente il funzionamento del sistema (anche attraverso la registrazione e conservazione dei log di tracciamento), dovendo sospenderne l'uso, qualora abbia ragione di ritenere, a fronte del monitoraggio, che nonostante il rispetto delle istruzioni ricevute il sistema presenti un rischio per la salute e sicurezza e i diritti fondamentali. Se un obbligo di valutazione di impatto sui diritti fondamentali (c.d. FRIA) è espressamente prescritto, dall'art. 27, solo per gli organismi di diritto pubblico o gli enti privati che forniscono servizi pubblici (come ospedali, scuole, banche e compagnie di assicurazione), la riflessione sull'importanza del paradigma *risk-based* in un contesto di utilizzo dell'IA invita, infine, a considerare una sua introduzione (volontaria) trasversale nelle posizioni d'obbligo datoriali ai fini di un loro corretto adempimento (Mantelero 2024; Peruzzi 2024)¹⁹.

Al di là di questi profili, si deve d'altra parte tener conto della possibilità che in capo al datore di lavoro gravi la posizione non solo di *deployer*, ma anche di fornitore, con tutti i relativi obblighi in termini di valutazione e gestione dei rischi del sistema di IA in questione, se classificabile come ad alto rischio.

Anzitutto, a norma dell'art. 3, n. 3, è fornitore non solo chi sviluppa un sistema di IA o un modello di IA per finalità generali (GPAI), ma anche chi li fa sviluppare, per poi commercializzarli o anche solo metterli in servizio (ossia consegnarli al *deployer* per uso proprio), purché apponga il proprio nome (ad esempio OpenAI, L.P.) o marchio (ad esempio ChatGPT) (Bassini 2024; il marchio può essere anche non registrato, cfr. Feiler e König 2025). In tal senso,

il datore può figurare al contempo come fornitore e *deployer* del sistema, laddove utilizzi software sviluppati internamente o fatti sviluppare e in ipotesi anche commercializzati, a cui abbia apposto il proprio nome o marchio. Il cumulo di posizioni si configura anche se il datore sviluppi o faccia sviluppare un sistema che integra un modello di IA (c.d. fornitore a valle), indipendentemente dal fatto che il modello sia fornito dallo stesso o da terzi²⁰. Il datore potrebbe quindi figurare, al contempo, fornitore a valle e *deployer* del sistema, nel momento in cui dapprima sviluppi o faccia sviluppare (con suo nome o marchio) un sistema che integra un modello di GPAI, quale ad esempio un modello di LLM per l'analisi del linguaggio naturale (ad esempio GPT-4, Bert), adattandolo o 'accordandolo' (c.d. *fine-tuning*) per eseguire come compito l'analisi del contenuto dei curricula ricevuti con selezione dei migliori candidati, quindi lo applichi direttamente per tale finalità d'uso²¹. Il *fine-tuning* può avvenire, ad esempio, con ri-addestramento supervisionato del modello, con nuovi dati specifici per il compito e l'ambito/settore interessato (ad esempio con i dati storici dei CV ricevuti dall'azienda in precedenza). Resta ferma, anche in questo, ai fini della qualificazione del soggetto come fornitore, la messa in servizio (se non anche commercializzazione) di un sistema di IA autonomo (quindi con una propria struttura applicativa, ad esempio una interfaccia utente, integrazione con basi di dati, flussi decisionali automatizzati) sotto il proprio nome o marchio.

È importante considerare, tuttavia, che il *deployer* è soggetto agli obblighi del fornitore, a norma dell'art. 25, anche laddove apporti a un sistema di IA ad alto rischio – già immesso sul

19 La prospettiva di una applicazione trasversale dell'approccio *risk-based* trova conferma nelle linee guida elaborate a livello interno dall'Agid per la Pubblica amministrazione e dal Ministero del Lavoro per il mondo del lavoro. In particolare, nella *Bozza di linee guida per l'adozione di IA nella pubblica amministrazione*, adottata e sottoposta a consultazione pubblica dall'AGID con la determinazione n. 17 del 17 febbraio 2025, si propone, accanto alla FRIA e DPIA, anche l'AIIA (*Artificial Intelligence Impact Assessment*). Nelle *Linee guida per l'implementazione dell'IA nel mondo del lavoro*, presentate dal Ministero del Lavoro il 14 aprile 2025 con apertura di una consultazione pubblica, si evidenzia la necessità di una valutazione di impatto – con verifica dei rischi e implementazione delle misure di mitigazione – sull'occupazione e transizione lavorativa, privacy e protezione dei dati, sicurezza e condizioni di lavoro, discriminazione e *bias* algoritmici (p. 40).

20 Laddove il modello sia sviluppato dallo stesso fornitore (è il caso di OpenAI e ChatGpt), il soggetto sarà gravato anche di tutti gli obblighi di cui al Capo V, se si tratta di modello di GPAI. Cfr. il codice di buone pratiche sull'IA per finalità generali e gli Orientamenti in materia (C(2025) 5045 final) pubblicati dalla Commissione a luglio 2025.

21 Cfr. il sistema proposto da Sareddy e Farhan (2025), basato sull'integrazione di un modello di LLM (BERT), modificato nella sua architettura (Nish-BERT) e quindi adattato per lo screening dei CV attraverso tecniche di clusterizzazione e di apprendimento supervisionato.

mercato o fornitogli – una modifica sostanziale²² che non incida su tale classificazione; ovvero anche laddove modifichi la finalità prevista di un sistema di IA, sia anche un sistema per finalità generali (GPAI)²³, in modo tale da richiederne una classificazione ad alto rischio. Il tema è di estrema rilevanza, posto che in questo caso l'assunzione della qualifica di fornitore non avviene a fronte dell'apposizione di nome e marchio al prodotto (per quanto poi richiesta, in forza della conseguente applicazione dell'art. 16, par. 1, lett. b)²⁴. Se ancora una volta la fattispecie della 'modifica sostanziale' dovrà essere chiarita dagli orientamenti della Commissione ai sensi dell'art. 96, pare utile richiamare sul punto l'art. 43, par. 4, secondo cui non costituiscono modifiche sostanziali quelle apportate al sistema e alle sue prestazioni dal processo di apprendimento continuo sviluppato dalla macchina dopo l'immissione sul mercato o la messa in servizio, purché predeterminate dal fornitore al momento della valutazione iniziale di conformità (v. Cons. 128).

Il profilo da ultimo esaminato richiede di porre particolare attenzione all'ipotesi in cui il *deployer* di un sistema di GPAI venga gravato anche degli obblighi propri del fornitore per averne modificato la finalità prevista.

La riflessione impone a monte due considerazioni. Anzitutto, a poter essere messo in servizio da un fornitore e utilizzato poi da un *deployer* è solo un sistema di GPAI (ad esempio ChatGPT) e non anche un modello di GPAI (ad esempio GPT-5). In altre parole, "un modello GPAI non può essere messo in servizio in modo autonomo. Solo una volta che il modello GPAI è stato integrato in un sistema di Intelligenza artificiale, esso può essere messo in servizio" (Voigt e Hullen 2024, 26); parimenti, l'AI

Act non si applica o riguarda gli utilizzatori di modelli GPAI, ma solo i *deployer* di sistemi di IA: del resto 'il funzionamento' di un modello di IA avviene sempre tramite un sistema di IA – ed è per quest'ultimo che possono essere previsti obblighi specifici a carico degli utilizzatori, distributori o importatori (in particolare nel caso di sistemi di IA ad alto rischio)" (Voigt e Hullen 2024, 148). La seconda considerazione attiene, invece, ai termini in cui si può individuare una modifica di finalità prevista in caso di sistema di GPAI, tale da integrare una modifica sostanziale del sistema. Al riguardo, la dottrina ha messo in luce che "non tutte le finalità potenziali di un sistema GPAI devono essere intese come finalità previste, ma solo quelle che il fornitore ha già specificato come campi di applicazione"; "la nozione di 'finalità generale' si riferisce alle capacità oggettive, indipendentemente dalle intenzioni del fornitore. La finalità prevista, invece, si riferisce esclusivamente alle intenzioni soggettive del fornitore. Questa distinzione è supportata anche dal fatto che l'uso improprio ragionevolmente prevedibile non rientra nella finalità prevista. Di conseguenza, la finalità prevista non dovrebbe dipendere dalla conoscenza del fornitore di come il sistema di IA potrebbe essere utilizzato, ma solo dalle finalità che il fornitore ha volontariamente assegnato al sistema" (Voigt e Hullen 2024, 26).

Si pensi, pertanto, all'ipotesi di un datore di lavoro che, come *deployer*, utilizzi ChatGPT e che, accedendo alla funzione che consente di 'customizzare' una versione della applicazione per una specifica finalità, decida di impostarne vincoli di contenuto tramite istruzioni strutturate (c.d. *prompt engineering*, a sua volta realizzabile con

22 La modifica è definita sostanziale, a norma dell'art. 3, n. 23, quando avviene a seguito della immissione sul mercato o messa in servizio del sistema, non è prevista o programmata nella valutazione iniziale di conformità effettuata dal fornitore e ha l'effetto di incidere sulla conformità del sistema di IA ai requisiti di cui al capo III, sezione 2, ovvero comporta una modifica della finalità prevista per la quale il sistema di IA è stato valutato. Come evidenzia Bassini (2024, 133), "il mutamento di qualificazione, che comporta l'assunzione dei relativi obblighi, non è però unilaterale, ma risulta 'biunivoco': in questi casi, infatti, l'originario fornitore che ha immesso sul mercato o messo in servizio il sistema di Intelligenza artificiale non sarà più considerato fornitore di quel determinato sistema, ferma restando l'esigenza di cooperazione con il 'nuovo' fornitore e la necessità di fornire tutte le informazioni necessarie".

23 A norma dell'art. 3, si tratta di "un sistema di IA basato su un modello di IA per finalità generali e che ha la capacità di perseguire varie finalità, sia per uso diretto che per integrazione in altri sistemi di IA".

24 Come evidenziano Riede e Talhoff (2025, 529), "se un tale cambio di ruolo avviene involontariamente, si crea un rischio significativo di responsabilità per il 'nuovo fornitore', poiché è altamente improbabile che adempia inconsapevolmente ai propri obblighi in quanto fornitore di un modello di IA ad alto rischio. Tutti gli operatori nella catena del valore dell'IA che apportano qualsiasi tipo di modifica ai sistemi di IA dovrebbero monitorare attentamente la conformità all'art. 25 per evitare di incorrere in responsabilità come 'deemed provider'".

il supporto dell'IA²⁵), con limitazione del dominio d'uso e integrazione della base di conoscenza con allegazione di fonti informative dedicate, ovvero anche collegamenti a banche dati aziendali tramite API. Ci si può chiedere, alla luce di quanto sopra esposto, se la personalizzazione, volta ad esempio a richiedere all'applicazione di valutare dei CV e predisporre un ranking dei candidati, implichi una modifica della finalità d'uso prevista del sistema – ad esempio da assistente generale per interazione conversazionale a selezionatore del personale – e, quindi, una modifica sostanziale tale da determinare l'applicazione al *deployer* anche degli obblighi stabiliti per il fornitore dall'art. 16²⁶. Il passaggio richiede, come detto, che la modifica determini o confermi una classificazione del sistema come ad alto rischio a norma dell'art. 6, una classificazione che richiede, d'altra parte, una verifica non solo dell'inclusione della pratica d'uso tra quelle critiche elencate dall'allegato III (ad esempio, analizzare o filtrare le candidature e valutare i candidati), ma anche della sussistenza di un rischio effettivo per i diritti fondamentali, esclusa, come detto, laddove l'output non incida sostanzialmente sull'esito del processo decisionale, ma sempre presente laddove il sistema effettui una profilazione di persone fisiche (come avviene nel caso in esame).

IA e contrattazione collettiva: le esperienze in Italia

Si è più volte evidenziata la necessità di dedicare specifica attenzione al ruolo delle Parti sociali, nella costruzione della dimensione collettiva delle tutele

che lo stesso AI Act richiama laddove prevede che lo standard minimo di garanzie predisposto possa essere incrementato a favore dei lavoratori tramite contratto collettivo (v. *supra*). Sempre nel contesto del Regolamento, a parte il (mero) obbligo informativo alle rappresentanze imposto al datore-*deployer* prima della messa in servizio/ utilizzo di un sistema ad alto rischio (art. 26, par. 7), si prevede il possibile ruolo delle Parti sociali nell'elaborazione dei codici di condotta e dei relativi sistemi di governance (Ciucciavino 2024), intesi a declinare i processi di alfabetizzazione sull'IA del personale²⁷, nonché l'applicazione volontaria delle garanzie ai sistemi a basso rischio ovvero di requisiti supplementari, a promozione di processi inclusivi e partecipati, ai sistemi ad alto rischio (art. 95; Cons. n. 165). Varchi d'accesso per il coinvolgimento delle rappresentanze si aprono nel raccordo con le fonti lavoristiche (si pensi, al ruolo del Rappresentante dei lavoratori per la sicurezza (RLS) nella valutazione e gestione dei rischi su salute e sicurezza ovvero delle rappresentanze ai fini dell'installazione di strumenti che consentono il controllo tecnologico a distanza) e non (il riferimento è anzitutto alla consultazione delle rappresentanze che può essere richiesta, se del caso, nella valutazione di impatto, c.d. DPIA, dalla normativa sulla protezione dei dati personali). Una specifica declinazione collettiva delle tutele rispetto alla gestione algoritmica è, infine, predisposta dalla Direttiva Piattaforme 2024/2831 (Aimo 2024), per il proprio ambito di applicazione, non a caso guardata come modello

25 Come spiega la dottrina “Negli LLM più potenti (come GPT-4) è possibile fornire istruzioni di input (prompt) di lunghezza considerevole (circa 50 pagine di testo). Pertanto, tali istruzioni possono comprendere descrizioni dettagliate del risultato da ottenere e del procedimento da seguire, esempi di casi in cui il compito sia stato eseguito correttamente ecc.” (Sartor et al. 2024, 255).

26 Come segnalano Riede e Talhoff (2025, 533), “per questi motivi, gli utilizzatori di sistemi AI saranno generalmente invitati a non utilizzarli in contesti ad alto rischio, a meno che il sistema AI ad alto rischio in questione non sia stato specificamente progettato per tale finalità dal fornitore”.

27 Il profilo è tutt'altro che secondario, soprattutto se si legge il Cons. n. 20, ove si afferma che il processo include le conoscenze necessarie per comprendere in che modo le decisioni adottate con l'assistenza dell'IA possano incidere sulle persone interessate. L'obbligo di alfabetizzazione, previsto dall'art. 4 AI Act, entrato in applicazione il 2 febbraio 2025, riguarda tutti i *deployer* di sistemi di IA, siano essi ad alto o basso rischio. Come chiarito dalla Commissione europea sul proprio sito, nella sezione *AI Literacy - Questions & Answers*, l'alfabetizzazione, a norma dell'art. 3 AI Act, deve includere “le competenze, le conoscenze e la comprensione che consentono (...) di procedere a una diffusione informata dei sistemi di IA, nonché di acquisire consapevolezza in merito alle opportunità e ai rischi dell'IA e ai possibili danni che essa può causare”. Il contenuto dell'obbligo è d'altra parte modulare e dipende dalle caratteristiche dell'organizzazione, dello strumento introdotto e dei relativi rischi, dal tipo di lavoratori coinvolti e dalle loro competenze: “in molti casi” – segnala la Commissione – “affidarsi semplicemente alle istruzioni per l'uso dei sistemi di Intelligenza artificiale o chiedere al personale di leggerle potrebbe risultare inefficace e insufficiente”, “non esiste un approccio valido per tutti (...). I requisiti che deve presentare la formazione dipendono dal contesto concreto”. Il tema dell'alfabetizzazione del personale è oggetto di particolare attenzione sia nelle *Linee guida per l'utilizzo dell'AI nella Pubblica amministrazione*, sia in quelle per il mondo del lavoro presentate dal Ministero del Lavoro (v. *supra*).

per un possibile rafforzamento delle garanzie per i lavoratori in materia, a livello UE (v. *supra*).

Rinviando ad altra sede una riflessione più complessiva sul punto, pare utile qui richiamare alcune esperienze nazionali, indicative del 'fermento' che si sta registrando a livello di contrattazione collettiva in Italia.

In principio, si può ricordare il noto Accordo Tim del 29 aprile 2020 – sottoscritto ai sensi dell'art. 4 St. Lav. con SLC-CGIL, FISTel-CISL, UILCom-UIL, UGL Telecomunicazioni, unitamente al Coordinamento nazionale RSU – volto a disciplinare l'uso della tecnologia 'Routing Evoluto' Afiniti, in grado di creare abbinamenti tra i clienti chiamanti e gli operatori telefonici, utilizzando sistemi di IA (Imberti 2023). Ancora, nel contesto dell'Accordo territoriale sottoscritto il 14 giugno 2022 dalle sigle di Confcooperative e Cgil, Cisl e Uil competenti per la provincia di Cuneo – finalizzato a disciplinare l'allineamento al CCNL per i lavoratori dipendenti da aziende cooperative di trasformazione di prodotti agricoli e zootecnici e lavorazione prodotti alimentari – sono condivisi alcuni 'assunti fondamentali' a guida dell'introduzione di robots/cobots e/o IA nelle singole realtà aziendali. Nello specifico, si evidenzia la necessità, sul piano dei contenuti, di garantire il controllo umano, "con riduzione della possibilità di adozione di misure o di decisioni in modo automatizzato", nonché l'"affidabilità e verificabilità dei dati e dei percorsi valutativi"; sul piano procedurale, si pone in luce l'importanza dell'informazione in punto di salute e sicurezza, possibilità di valutazione e ricadute occupazionali, richiedendosi all'azienda di indicare preventivamente alle Rsa (o, in loro mancanza, alle organizzazioni sindacali territoriali stipulanti l'accordo territoriale): "tipologia di innovazione, nelle sue componenti tecniche essenziali e nel suo previsto impatto sul ciclo produttivo e sull'assetto organizzativo preesistente; in caso di intelligenze artificiali: i) il tipo di dati che le medesime siano chiamate a raccogliere; ii) i tempi e i modi di conservazione dei dati; iii) le finalità della raccolta dei medesimi".

Rispetto al tema della gestione algoritmica, particolare risonanza mediatica ha riscontrato il Protocollo di intesa del 28 novembre 2023 sottoscritto da Assoespressi e le aziende appaltatrici dell'Emilia-Romagna che si occupano delle consegne 'ultimo miglio' nella filiera Amazon, a cura dei driver:

al suo interno, Assoespressi si è resa disponibile a effettuare incontri specifici con le aziende ed eventualmente con le organizzazioni sindacali, su richiesta delle stesse, per verificare situazioni di carichi di lavoro eccessivi con particolare riferimento a singole rotte o eventualmente su zone particolari. Sempre in tale sede, si è prevista la promozione a livello aziendale di campagne di sensibilizzazione dei dipendenti per un utilizzo più corretto del device e si è evidenziata, sul tema, la piena rilevanza del ruolo del RLS.

Numerosi gli interventi che si sono registrati lo scorso anno.

Il Protocollo Saipem del febbraio 2024, sottoscritto ai sensi dell'art. 4 St. Lav., ha riguardato l'introduzione e disciplina di *smart cameras* nel cantiere di Argo/Cassiopea e, quindi, la sperimentazione di un sistema di IA nel campo della sicurezza sui posti di lavoro. A tal fine, è stato coinvolto anche l'Organismo paritetico nazionale di HSE (Salute e sicurezza), previsto dal Contratto collettivo nazionale di lavoro.

Spostando lo sguardo sull'ambito delle piattaforme, può essere significativo ricordare l'Accordo stipulato il 19 febbraio 2024 da Assogrocery con Cgil, Cisl e Uil per la disciplina del lavoro su piattaforma digitale degli shopper. All'art. 5, non solo si dispone un obbligo informativo alle rappresentanze sui processi anche parzialmente automatizzati, ma si prevede anche che a tali informazioni segua la possibilità di un confronto.

L'Accordo di rinnovo per le Banche di Credito cooperativo – Casse Rurali e artigiane, del 9 luglio 2024 – oltre a richiamare esplicitamente la *Joint declaration* sulla IA delle Parti sociali europee di settore del maggio 2024 – ha stabilito che l'individuazione di nuove mansioni e figure professionali sia oggetto di soluzioni condivise nel contesto di un Organismo nazionale bilaterale e paritetico sull'impatto delle nuove tecnologie/digitalizzazione.

L'Accordo integrativo aziendale di GlaxoSmithKline SpA sottoscritto in data 26 luglio 2024 con la RSU Femca, Uiltec e Filctem per il quadriennio 2024-2027 prevede, rispetto alla adozione di nuove tecnologie, tra cui quelle di IA, la "condivisione di iniziative efficaci (...) anche attraverso momenti informativi e formativi dedicati alla RSU e alla popolazione aziendale"; in particolare si stabilisce che la parti condividano "in

forma preventiva e strutturata (...) percorsi finanziati attraverso lo strumento di Fondimpresa estesi al maggior numero possibile di lavoratori. A tal fine, anche con riferimento all'evoluzione tecnologica e digitale degli strumenti di lavoro, le Parti intendono investire in percorsi formativi dedicati”.

L'Accordo del 23 settembre 2024 tra la Sanofi S.p.a. e le strutture territoriali di Femca-Cisl, Filctem-Cgil e Uiltec-Uil, in assistenza della RSU, ha previsto un *Patto per il digitale e l'Intelligenza artificiale*, con impegno a verificarne gli effetti. Il Patto è stato redatto dall'Osservatorio digitale bilaterale, costituito dalle parti nel 2022 ai sensi di quanto previsto dal CCNL Industria chimica. In tale sede, sono enunciati alcuni principi generali guida, a garanzia di un approccio alle tecnologie basato sulla centralità del ruolo umano, il principio di accountability, la protezione dei dati, il rispetto per l'integrità umana, l'equità e inclusione, il benessere comune e la sostenibilità ambientale, nonché la salvaguardia della trasparenza, la promozione del coinvolgimento, lo sviluppo della formazione.

L'Accordo di percorso sulla trasformazione del Gruppo Intesa San Paolo, del 23 ottobre 2024, impegna azienda e organizzazioni sindacali a un confronto costante per monitorare gli effetti per tutto il Gruppo, comprese le ricadute dirette sulla rete filiali, del passaggio alla nuova piattaforma digitale Isytech e dello sviluppo dell'IA, anche attraverso l'individuazione di nuovi mestieri, per i quali le parti si confronteranno, nell'ambito della contrattazione di secondo livello, per definire possibili percorsi di sviluppo professionale e delle competenze. A tal fine si prevede l'attivazione del Comitato trasformazione digitale, che si riunirà indicativamente ogni due mesi e avrà il presidio per tutte le tematiche della Banca digitale e della trasformazione digitale del Gruppo.

Più di recente, il CCNL del settore elettrico dell'11 febbraio 2025 dispone, da un lato che l'Osservatorio bilaterale di settore costituisca la sede anche per esaminare la tematica dell'applicazione dell'AI Act e dell'impatto dell'IA sul lavoro, dall'altro che a livello aziendale siano previsti incontri preventivi di confronto e consultazione tra le aziende e le Organizzazioni sindacali – previa formazione delle parti – per un coinvolgimento nell'analisi e verifica degli esiti dell'uso dell'IA.

Ancora, l'Accordo ponte' del 7 marzo 2025 tra le Segreterie nazionali e territoriali Slc Cgil, Fisl Cisl,

Uilcom Uil, unitamente alle RSU/RSA, e l'azienda Konecta include un patto di sistema sulla gestione etica degli strumenti di IA. In particolare, in merito agli strumenti di IA atti a supportare i lavoratori durante lo svolgimento della prestazione lavorativa, sono previste garanzie specifiche per l'utilizzo dello strumento Agent Assist, tra cui l'istituzione di un Comitato paritetico, 8 componenti aziendali e 8 componenti sindacali, che monitori il funzionamento dello strumento e verifichi il rispetto di quanto sancito nell'accordo.

Il 19 marzo 2025 è stato sottoscritto un accordo tra la società di assicurazioni CNP Vita Assicura e le RSA Fisac-Cgil, Snfia e Uilca-Uil, per regolare l'adozione dell'IA nei processi aziendali, con la tutela dei diritti dei dipendenti e promozione al contempo dell'innovazione responsabile. In particolare, l'accordo prevede che “per tutti coloro che dovessero essere impattati nelle loro attività lavorative dall'introduzione delle suddette nuove tecnologie, l'azienda garantisce corsi di formazione ad hoc che consentano una riqualificazione e, ove necessario e/o compatibile con le esigenze aziendali, una ricollocazione in altra mansione, coerente con la professionalità, l'inquadramento e l'esperienza acquisita”. Al contempo, si vieta l'utilizzo dell'IA “nei processi industriali aziendali se l'impiego di questa nuova tecnologia genera ricadute occupazionali, intendendosi per tali la perdita del posto di lavoro, sul personale assunto a tempo indeterminato alla data del presente accordo”.

Se il 31 marzo 2025 è stato indetto uno sciopero da parte delle sigle sindacali del settore TLC per lo stallo nella rinegoziazione del CCNL, pare utile ricordare come proprio all'interno della piattaforma sindacale presentata per tale rinnovo si trovi uno dei profili più interessanti e innovativi nella prospettiva d'analisi che ci occupa, ossia l'istituzione della figura del delegato per la gestione dei dati (a calco, quindi, della figura del RLS prevista per l'ambito della salute e sicurezza sul lavoro²⁸).

Il 15 aprile 2025 è stato sottoscritto l'accordo di rinnovo del CCNL per i settori chimico e farmaceutico, valido per il triennio 2025-2028, all'interno del quale trovano sede delle linee guida dedicate all'utilizzo dell'IA nell'Organizzazione aziendale e un Patto sociale per le competenze attento ai profili della transizione digitale ed ecologica. Sul primo versante, viene data particolare rilevanza alla mappatura e analisi di tutti i rischi, da effettuare prima dell'introduzione dei sistemi, nonché

28 La prospettiva di un rappresentante specializzato in materia era stata avanzata in dottrina da Ingrao (2023).

alla trasparenza e al coinvolgimento dei lavoratori e dei loro rappresentanti: adeguata informazione e formazione ai lavoratori coinvolti; informazione preventiva alla RSU sull'introduzione di progetti di IA; coinvolgimento di lavoratori e RSU per la comprensione di scopi, procedure, rischi e benefici.

Considerazioni conclusive

“Riusciamo a vedere solo per un breve tratto davanti a noi, ma già lì vediamo moltissime cose da fare”. Con queste parole Alan Turing chiudeva il saggio *Computing Machinery and Intelligence* comparso nel 1950 sulla rivista *Mind*. L'IA è una tecnologia da organizzare e regolare. In queste pagine, da due angolature, abbiamo provato a mettere in fila alcuni contenuti per capire di che cosa stiamo parlando: un primo passaggio che riteniamo necessario. In termini analitici, di progettazione organizzativa e di regolazione davanti

a noi “vediamo moltissime cose da fare”. È vero che cogliamo solo un breve tratto del percorso, ma osando, sul filo dell'ossimoro, dovremmo sviluppare una lettura sul breve periodo con una prospettiva che guarda lontano. Il lavoro sta mutando rapidamente, il mondo del lavoro sta raggiungendo un livello di eterogeneità interna forse mai visto in epoche precedenti. Ai lavori segnati da processi organizzativi e contenuti operativi piuttosto *tradizionali*, vediamo affiancarsi *nuovi scenari*, ormai del tutto maturi, dentro i quali l'IA sta giocando un ruolo di estrema rilevanza. All'interno di questo perimetro dobbiamo generare dialogo fra le discipline che studiano il lavoro, così come fra gli attori che regolano il lavoro. Nei prossimi anni sarà centrale creare competenze e spazi per discutere quanto sta avvenendo attorno al lavoro, per comprendere e governare un processo di *mutamento antropologico* che passa *anche* attraverso il lavoro.

Bibliografia

- Adler D. (2024), *Il duplice enigma. Intelligenza Artificiale e intelligenza umana*, Torino, Einaudi
- Aimo M. (2024), Trasparenza algoritmica nel lavoro su piattaforma: quali spazi per i diritti collettivi nella proposta di Direttiva in discussione?, *Lavoro Diritti Europa*, n.2, 6 maggio
- Bassini M. (2024), Oggetto, campo di applicazione e ambito territoriale, in Pollicino O., Donati F., Finocchiaro G., Paolucci F. (a cura di), *La disciplina dell'Intelligenza Artificiale*, Milano, Giuffrè, pp.113 ss.
- Bengio Y. (2016), Macchine che imparano, *Le Scienze, Speciale I Quaderni: Intelligenza Artificiale*, n.576
- Bradford A. (2020), *The Brussels Effect: How the European Union Rules the World*, Oxford, Oxford University Press
- Buoso S. (2020), *Principio di prevenzione e sicurezza sul lavoro*, Torino, Giappichelli
- Butera F., De Michelis G. (2024), *Intelligenza Artificiale e lavoro, una rivoluzione governabile*, Venezia, Marsilio
- Canal T., Gosetti G., Luppi M. (2024), Qualità del lavoro e digitalizzazione. Riflessioni aperte sul caso Italiano, *Sinapsi*, XIV, n.2, pp.66-92
- Castel R. (2019), *La metamorfosi della questione sociale. Una cronaca del salariato*, Sesto San Giovanni (MI), Mimesis
- Castel R. (2004), *L'insicurezza sociale. Cosa vuol dire essere protetti*, Torino, Einaudi
- Ciucciiovino S. (2024), Risorse umane e Intelligenza Artificiale alla luce del regolamento (UE) 2024/1689, tra norme legali, etica e codici di condotta, *Diritto delle Relazioni Industriali*, n.3, pp.573-614
- Commissione europea (2025), *Analysis of EU AI Office stakeholder consultations: defining AI systems and prohibited applications. Final study report*, Luxembourg, Publications Office of the European Union
- Crawford K. (2021), *Né intelligente né artificiale. Il lato oscuro dell'IA*, Bologna, il Mulino
- Crozier M., Friedberg H. (1978), *Attore sociale e sistema. Sociologia dell'azione organizzata*, Milano, Etas
- De Masi D. (2018), *Il lavoro nel XXI secolo*, Torino, Einaudi
- Dejours C. (2020), *Lavoro vivo*, Sesto San Giovanni (MI), Mimesis
- Domingos P. (2018), Le nostre frontiere digitali, *Le Scienze*, n.603, pp.84-89
- Esposito E. (2022), *Comunicazione artificiale. Come gli algoritmi producono intelligenza sociale*, Milano, Bocconi University Press
- European Commission (2025), *Annex to the Communication to the Commission - Approval of the content of the draft Communication from the Commission - Commission Guidelines on the definition of an artificial intelligence system established by Regulation (EU) 2024/1689 (AI Act)*, Brussels, 6.2.2025, C(2025) 924 final) ANNEX

- Faioli M. (2025), Adeguatezza ex art. 2086 c.c. e obbligo di introdurre tecnologia avanzata per mitigare i rischi da lavoro, *federalismi.it*, n.6, pp.143-157
- Falletta P., Marsano A. (2024), Intelligenza Artificiale e protezione dei dati personali: il rapporto tra Regolamento europeo sull'Intelligenza Artificiale e GDPR, *Rivista italiana di informatica e diritto*, n.1, pp.119-137
- Feiler L., König B. (2025), Article 3. Definitions, in Pehlivan C.N., Forgó N., Valcke P. (eds.), *The EU Artificial Intelligence (AI) Act. A Commentary*, Alphen aan den Rijn, Wolters Kluwer, p.57
- Finocchiaro G. (2024), *Intelligenza Artificiale. Quali regole?*, Bologna, il Mulino
- Floridi L. (2022), *Etica dell'Intelligenza Artificiale. Sviluppi, opportunità, sfide*, Milano, Raffaello Cortina Editore
- Ford M. (2022), *Il dominio dei robot. Come l'Intelligenza Artificiale rivoluzionerà l'economia, la politica e la nostra vita*, Milano, il Saggiatore
- Gaudio G. (2024), Valutazione di impatto e management algoritmico, *Rivista giuridica del lavoro e della previdenza sociale*, n.4, p.538 ss.
- Gigerenzer G. (2023), *Perché l'intelligenza umana batte ancora gli algoritmi*, Milano, Raffaello Cortina Editore
- Gosetti G. (2024), *Percorsi di sociologia del lavoro. Da Taylor alla società dei lavori*, Milano, Franco Angeli
- Gosetti G. (2022), *La qualità della vita lavorativa. Lineamenti per uno studio sociologico*, Milano, Franco Angeli
- Gosetti G., Peruzzi M., Calabrese B. (2024), *Digitalizzazione e lavoro. Elementi per un'analisi giuridica e sociologica*, Torino, Giappichelli
- Harari Y.N. (2024), *Nexus. Breve storia delle reti di informazione*, Milano, Bompiani
- Honneth A. (2025), *Il lavoratore sovrano. Lavoro e cittadinanza democratica*, Bologna, il Mulino
- Iannuzzi A. (2025), La definizione di Intelligenza Artificiale fra Regolamento europeo e norma tecnica internazionale, in Pizzetti F. (a cura di), *La regolazione europea dell'Intelligenza Artificiale nella società digitale*, Torino, Giappichelli, pp.11 ss.
- Imberti L. (2023), Intelligenza Artificiale e sindacato. Chi controlla i controllori artificiali?, *federalismi.it*, n.29, pp.192-201
- Ingrao A. (2023), Controllo a distanza e privacy del lavoratore alla luce dei principi di finalità e proporzionalità della sorveglianza, *Labour & Law Issues*, 9, n.1, pp.101-121
- Italiano G.F., Civitarese Matteucci S., Perrucci A. (2022), L'Intelligenza Artificiale: dalla ricerca scientifica alle sue applicazioni. Una introduzione di contesto, in Pajno A., Donati F., Perrucci A. (a cura di), *Intelligenza Artificiale e diritto: una rivoluzione? Diritti fondamentali, dati personali e regolazione*, vol. I, Bologna, il Mulino, pp.43 ss.
- Kaplan J. (2024), *Generative A.I. Conoscere capire e usare l'Intelligenza Artificiale generativa*, Roma, Luiss University Press
- Lai M. (2025), Brevi note in tema di Intelligenza Artificiale e salute e sicurezza del lavoro, *Lavoro Diritti Europa*, n.1, 9 marzo
- Lazzari C., Pascucci P. (2024), Sistemi di IA, salute e sicurezza sul lavoro: una sfida al modello di prevenzione aziendale, fra responsabilità e opportunità, *Rivista giuridica del lavoro e della previdenza sociale*, n.4, Parte I, pp.587 ss.
- Loi P. (2001), *La sicurezza. Diritto e fondamento dei diritti nel rapporto di lavoro*, Torino, Giappichelli
- Mantelero A. (2024), The Fundamental Rights Impact Assessment (FRIA) in the AI Act: Roots, legal obligations and key elements for a model template, *Computer Law and Security Review*, n.54, pp.9 ss.
- Mantelero A., Peruzzi M. (2024), L'AI Act e la gestione del rischio nel sistema integrato delle fonti, *Rivista giuridica del lavoro e della previdenza sociale*, n.4, Parte I, pp.517 ss.
- Marconi M. (2025), L'intelligenza di Turing, in Turing A.M., *Macchine calcolatrici e intelligenza*, Torino, Einaudi
- McQuaid J. (2022), Spiare le emozioni, *Le Scienze*, n.642, pp.42-49
- Mézard M. (2024), L'intelligenza artificiale a supporto della scienza, *eco*, n.5, pp.40-43
- Mitchell M. (2022), *L'Intelligenza Artificiale. Una guida per esseri umani pensanti*, Torino, Einaudi
- Novella M. (2022), Impresa, in Novella M., Tullini P. (a cura di), *Lavoro digitale*, Torino, Giappichelli, pp.5 ss.
- Orofino M. (2025), La regolazione dei modelli di IA per finalità generali, in Pizzetti F. (a cura di), *La regolazione europea dell'Intelligenza Artificiale nella società digitale*, Torino, Giappichelli, pp.95 ss.
- Paliokaitė A., Christenko A., Aloisi A., Riedel L., Ratti M., De Stefano V. (2025), *Study exploring the context, challenges, opportunities, and trends in algorithmic management*, Luxembourg, Publications Office of the European Union
- Perilli L. (2025), *Coscienza artificiale. Come le macchine pensano e trasformano l'esperienza umana*, Milano, il Saggiatore

- Peruzzi M. (2024), IA e obblighi datoriali di tutela del lavoratore: necessità e declinazioni dell'approccio risk-based, *federalismi.it*, n.30, pp.225-249
- Peruzzi M. (2023), *Intelligenza Artificiale e lavoro. Uno studio su poteri datoriali e tecniche di tutela*, Torino, Giappichelli
- Riede L., Talhoff O. (2025), Article 25. Responsibilities Along the AI Value Chain, in Pehlivan C. N., Forgó N., Valcke P. (eds.), *The EU Artificial Intelligence (AI) Act. A Commentary*, Alphen aan den Rijn, Wolters Kluwer, pp.516 ss.
- Riva G. (2025), *Io, noi, loro. Le relazioni nell'era dei social e dell'IA*, Bologna, il Mulino
- Runciman D. (2024), *Affidarsi. Come abbiamo ceduto il controllo della nostra vita a imprese, Stati e intelligenze artificiali*, Torino, Einaudi
- Sareddy M.R., Farhan M. (2025), Automatic Resume Screening Approach Based on AI and Ethereum Blockchain for Human Resource Management Using Gaaru, *SN Computer Science*, 6, n.4 <DOI:10.1007/s42979-025-03787-8>
- Sartor G., De Gregorio G., Muto G. (2024), Large Language Models, Intelligenza Artificiale generativa e Artificial Intelligence Act, in Pollicino O., Donati F., Finocchiaro G., Paolucci F. (a cura di), *La disciplina dell'Intelligenza Artificiale*, Milano, Giuffré, pp.259 ss.
- Schmidl M., Rohner A. (2025), Article 70. Designation of National Competent Authorities and Single Point of Contact, in Pehlivan C.N., Forgó N., Valcke P. (a cura di), *The EU Artificial Intelligence (AI) Act. A Commentary*, Alphen aan den Rijn, Wolters Kluwer, pp.1019 ss.
- Sennett R. (2000), *L'uomo flessibile. Le conseguenze del nuovo capitalismo sulla vita personale*, Milano, Feltrinelli
- Spitzer M. (2024), *Intelligenza Artificiale. Opportunità e rischi di una rivoluzione tecnologica che sta cambiando il mondo*, Milano, Corbaccio
- Tebano L. (2020), *Lavoro, potere direttivo e trasformazioni organizzative*, Napoli, Editoriale scientifica
- Tönnies F. (1887), *Gemeinschaft und Gesellschaft. Abhandlung des Communismus und des Socialismus als empirischer Culturformen*, Leipzig, Fues
- Treu T. (2025), *Il controllo umano delle tecnologie: regole e procedure*, WP C.S.D.L.E. "Massimo D'Antona".IT n.492, Catania, Centre for the Study of European Labour Law "MASSIMO D'ANTONA", University of Catania
- Treu T. (2024), *Intelligenza Artificiale (IA): integrazione o sostituzione del lavoro umano?*, WP C.S.D.L.E. "Massimo D'Antona". IT n.487, Catania, Centre for the Study of European Labour Law "MASSIMO D'ANTONA", University of Catania
- Tullini P. (2021), La questione del potere nell'impresa. Una retrospettiva lunga mezzo secolo, *Lavoro e diritto*, n.3-4, pp.442 ss.
- Turing A. (1950), Computing Machinery and Intelligence, *Mind*, LIX, n.236, pp.433-460
- Varanini F. (2024), *Splendori e miserie delle intelligenze artificiali. Alla luce dell'esperienza umana*, Milano, Guerini e associati
- Voigt P., Hullen N. (2024), *The EU AI Act. Answers to Frequently Asked Questions*, Berlin, Heidelberg, Springer
- Zuboff S. (2019), *Il capitalismo della sorveglianza. Il futuro dell'umanità nell'era dei nuovi poteri*, Roma, Luiss University Presse

Giorgio Gosetti

giorgio.gosetti@univr.it

Insegna Sociologia del lavoro presso il Dipartimento di Scienze umane dell'Università degli Studi di Verona ed è Responsabile scientifico del Centro di ricerca RE-WOrk (REsearching for REmaking Work and Organizing). Fra le sue pubblicazioni recenti si segnalano: (con Peruzzi M., Calabrese B.), *Digitalizzazione e lavoro. Elementi per un'analisi giuridica e sociologica*, Giappichelli, 2024; *Percorsi di sociologia del lavoro. Da Taylor alla società dei lavori*, Franco Angeli, 2024; *La qualità della vita lavorativa. Lineamenti per uno studio sociologico*, Franco Angeli, 2022.

Marco Peruzzi

marco.peruzzi@univr.it

Insegna Diritto del lavoro presso il Dipartimento di Scienze giuridiche dell'Università degli Studi di Verona. Fra le sue pubblicazioni recenti si segnalano: IA e obblighi datoriali di tutela del lavoratore: necessità e declinazioni dell'approccio risk-based, *Federalismi.it*, 2024; Gestione algoritmica del lavoro, protezione dei dati personali e tutela collettiva, *Lavoro e diritto*, n.2, 2024; *Intelligenza artificiale e lavoro*, Giappichelli, 2023.