



Numero 3 / 2025

Francesco ROTONDI

L'impatto dell'I.A. sulle fonti regolative del diritto del lavoro

L'impatto dell'IA. sulle fonti regolative del diritto del lavoro

Francesco ROTONDI

Giuslavorista. Consigliere Esperto CNEL. Avvocato

1. L'incontro dell'IA con le organizzazioni coinvolge la persona umana nei suoi profili più delicati e sensibili: ciò accade sia in fase di gestione (cd. "algoritmica") che in fase di controllo del lavoro; e sia nella sfera normativa che in quella etica.

Qualche esempio può essere utile a cogliere i suddetti profili nei contesti applicativi dell'IA.

Si pensi, in primo luogo, all'impatto della gestione algoritmica dei rapporti di lavoro sui processi di selezione e *matching* tra domanda e offerta di lavoro.

Una seconda macroarea di impatto dell'IA riguarda i cd. *people analytics*, ossia l'analisi ed elaborazione dei dati dei lavoratori anche in funzione predittiva, a fini di verifica dell'adempimento, valutazione del rendimento, premiazione del merito, .

Un ulteriore ambito applicativo dei sistemi di IA riguarda la salute e sicurezza sul lavoro. Qui le applicazioni spaziano dai dispositivi di protezione individuale intelligenti (elmetti e dispositivi indossabili per controllare lo stato di attenzione e altre caratteristiche psico fisiche delle persone), a strumenti e dispositivi (occhiali, monitor, guanti, simulatori, ecc.) che grazie alla realtà aumentata e al metaverso permettono di svolgere in maggiore sicurezza e con maggiore puntualità compiti delicati o rischiosi. Allo stesso tempo si hanno utilizzi dei *data analytics* per la prevenzione della incidentalità.

La rassegna dei possibili utilizzi è appena agli inizi di fronte alla loro potenziale infinitezza; anche in considerazione del fatto che l'IA vive in simbiosi binaria con le organizzazioni del lavoro, fungendo ora da strumento dell'organizzazione, ora da sua infrastruttura.

E'agevole rilevare le potenziali criticità per i diritti dei lavoratori, in ambiti ampi ed eterogenei, quali la nozione di subordinazione e diligenza, i limiti ai poteri datoriali di controllo, la discrezionalità dei poteri valutativi e selettivi che governano gli atti di costituzione, gestione e risoluzione dei rapporti di lavoro.

In generale, appare chiaro come la dimensione dominante in cui muovono gli aspetti rilevati sia quella normativa. Così come, in particolare, appare chiaro che il limite che i suddetti poteri datoriali incontrano, ruota oltre che attorno al tema della protezione dei dati personali, attorno al divieto di discriminazione.

Ma esiste, come abbiamo anticipato, un nesso tra cd. "etica d'impresa" e intelligenza artificiale: è questo nesso che ha caratterizzato l'approccio normativo europeo al tema dell'intelligenza artificiale fin dall'inizio, segnando la differenza fondamentale tra Europa e USA.

In realtà, nell'ultimo ventennio del secolo scorso aveva iniziato a diffondersi, soprattutto nelle grandi imprese multinazionali, il fenomeno della responsabilità sociale d'impresa (R.S.I.) o *Corporate Social Responsibility*. Ma la CSR era una forma di protezione dei soggetti diversi dagli

azionisti (*stakeholders*), mentre i codici etici di nuova generazione si caratterizzano per la loro simbiosi con l'intelligenza artificiale, di cui sembrano assurgere a sottoprodotti.

Cogliere questa simbiosi significa comprendere le ragioni di un approccio equilibrato e aperto, non integralmente liberista, ma nemmeno oppositivo, all'intelligenza artificiale,

Anche qui, qualche esempio può essere utile: come il caso del sistema di intelligenza artificiale utilizzato per stimare il rischio sanitario di oltre 200 milioni di americani, i cui algoritmi tendevano ad assegnare un livello di rischio inferiore ai pazienti afroamericani, perché la spesa sanitaria media individuale risultava meno elevata per la popolazione afroamericana, con la conseguenza di negare a quest'ultima l'accesso a cure adeguate.

Ovvero il caso della municipalità nordamericana che discriminava inconsapevolmente i neri nel ridisegno della mobilità urbana, effettuato sulla base delle indicazioni fornite dagli smart-phone (posseduti in numero inferiore dai neri).

Gli approcci rimediali a siffatti “errori” valutativi sono sostanzialmente tre. Il primo, che si potrebbe definire scienziato, permette di capire con semplicità tutti i pregiudizi di cui le IA sono insieme vittime e diffusori, fra manager solo maschi, cameriere solo donne e operai solo afroamericani.

Si tratta di un rimedio adatto a contrastare le discriminazioni.

Il secondo è quello accolto dall'AI Act del 2024, la quale si basa su un sistema di valutazione del rischio, inserendo tra questi, per esempio, i sistemi di intelligenza artificiale utilizzati nel **settore sanità**, ma anche quelli utilizzati nella gestione delle risorse umane; sistemi per il cui utilizzo vengono imposte **rigide valutazioni di conformità**, al fine di assicurarne l'accuratezza, la robustezza e la sicurezza.

2. In questo ambito devono menzionarsi i codici etici, deputati a riempire di contenuti, in chiave sussidiaria, i principi formulati dell'A.I. Act, a cominciare dal principio di personalità o antropocentrico, con cui si chiude il cerchio, trattandosi da un lato della regola per cui nessuna decisione idonea a modificare la sfera giuridica di una persona può essere adottata integralmente da un sistema di intelligenza artificiale; dall'altro, il principio di accountability, che, per ogni violazione delle norme e dei codici etici, esige l'esistenza di un responsabile (*ca va sans dire*: umano),

Il Codice Etico IA deve stabilire criteri chiari per valutare il livello di trasparenza degli algoritmi utilizzati dall'azienda. Un modello IA è considerato trasparente quando:

1. Gli utenti e i soggetti coinvolti sono informati sull'uso dell'IA e sulle modalità con cui questa influisce sulle loro interazioni con l'azienda.
2. Le decisioni algoritmiche sono spiegabili e verificabili, con la possibilità di accedere a documentazione chiara sui criteri di funzionamento dell'IA.
3. I dataset utilizzati per addestrare gli algoritmi sono documentati e verificabili, per garantire l'assenza di bias o errori sistematici.

4. Esistono meccanismi di audit interni ed esterni, per garantire che gli algoritmi siano utilizzati in modo conforme alle normative vigenti.

Questi principi devono tradursi in obblighi precisi per sviluppatori, responsabili della compliance e decisori aziendali, affinché ogni sistema IA implementato sia documentato, tracciabile e accessibile agli stakeholder.

Uno dei problemi più critici nell'uso dell'IA è la spiegabilità delle decisioni, ossia la capacità di fornire una motivazione chiara e intelligibile per ogni scelta effettuata dall'algoritmo.

Il Codice Etico IA deve prevedere che:

- Ogni sistema IA sia dotato di un modulo di spiegabilità, che consenta di comprendere le logiche di funzionamento e i dati utilizzati.
- Gli utenti abbiano il diritto di richiedere una spiegazione dettagliata, soprattutto nei settori ad alto impatto (HR, finanza, sanità).
- Le decisioni IA siano validate attraverso controlli umani, affinché sia sempre possibile correggere errori o anomalie.

L'adozione di tecniche di Explainable AI (XAI) è una best practice fondamentale per garantire che anche gli algoritmi più complessi (machine learning, deep learning) siano comprensibili e verificabili.

Per garantire l'applicazione effettiva del principio di trasparenza, il Codice Etico IA deve prevedere strumenti operativi concreti, tra cui:

- Dichiarazioni di trasparenza per ogni sistema IA, con documentazione accessibile che descriva finalità, modalità di funzionamento e criteri decisionali.
- Registro delle decisioni algoritmiche, per tracciare le scelte IA e consentire verifiche e contestazioni in caso di anomalie.
- Audit periodici sulla trasparenza IA, per valutare e correggere eventuali criticità nei sistemi automatizzati.
- Coinvolgimento di esperti multidisciplinari, che supervisionino i processi di verifica della trasparenza degli algoritmi.

Un'azienda che investe nella trasparenza IA ottiene benefici tangibili, tra cui:

- Fiducia da parte degli utenti.

- Maggiore reputazione aziendale.
- Riduzione del rischio normativo e operativo.
- Migliore accesso a finanziamenti e investimenti.

3. L'Intelligenza Artificiale ha il potenziale per aumentare l'efficienza e l'oggettività nei processi decisionali, ma se non progettata e gestita correttamente può rafforzare pregiudizi e discriminazioni già presenti nei dati storici. Per questo motivo, il principio di equità, non discriminazione e rispetto della dignità deve essere un pilastro fondamentale del Codice Etico IA, imponendo all'azienda di adottare misure preventive e correttive per garantire che i sistemi IA siano imparziali, equi e privi di distorsioni discriminatorie.

Le aziende devono quindi impegnarsi a testare, validare e monitorare periodicamente i propri algoritmi, evitando che essi possano produrre risultati distorti o ingiusti. Questo principio non è solo un obbligo normativo, ma un fattore strategico, che permette di costruire modelli IA affidabili e sostenibili, rafforzando la fiducia di utenti, clienti e stakeholder.

La lotta alla discriminazione nell'uso dell'IA è sancita da diverse normative e regolamenti, che pongono obblighi chiari per le imprese:

1. AI Act Europeo

- Impone alle aziende di identificare, valutare e mitigare i bias algoritmici nei sistemi IA ad alto rischio.
- Prevede l'adozione di audit obbligatori e controlli periodici, per garantire che gli algoritmi non generino discriminazioni sistemiche.

2. GDPR (Regolamento UE 2016/679)

- L'articolo 22 vieta decisioni interamente automatizzate senza garanzie di equità e supervisione umana.
- Stabilisce che gli utenti devono poter contestare le decisioni IA e ottenere una revisione umana, se l'algoritmo influisce significativamente sui loro diritti.

3. Carta dei Diritti Fondamentali dell'Unione Europea

- Stabilisce che nessun individuo può essere discriminato in base a genere, etnia, religione, opinioni politiche, età o disabilità.
- Questo principio si estende anche alle decisioni automatizzate basate su IA.

4. Linee guida OCSE sull'IA

- Stabiliscono che gli algoritmi IA devono essere equilibrati e privi di distorsioni, e muniti di sistemi di validazione periodici.
- Raccomandano che le aziende adottino procedure di monitoraggio e testing continuo.

5. Direttiva sull'Equità Salariale UE (2023)

- Impone obblighi di trasparenza nei processi di selezione e valutazione del personale, regolando l'uso degli algoritmi HR per evitare discriminazioni salariali o professionali.

Alla luce di queste disposizioni, il Codice Etico IA deve stabilire standard chiari per prevenire e correggere i bias algoritmici, garantendo che ogni decisione IA sia equa, verificabile e priva di distorsioni.

Affinché il principio di non discriminazione sia realmente applicabile, l'azienda deve adottare strategie strutturate di monitoraggio e correzione, che coinvolgano l'intero ciclo di vita degli algoritmi IA, tra le quali il testing e la validazione preventiva dei modelli IA, il test di fairness sui dataset di addestramento, per individuare eventuali squilibri nei dati utilizzati dagli algoritmi. metriche di equità personalizzate, che garantiscano un trattamento equo tra categorie di persone diverse (es. genere, etnia, età, status socioeconomico).

4. Un dataset imparziale è fondamentale per garantire decisioni IA eque e non discriminatorie. L'azienda deve adottare politiche e tecniche di data governance rigorose, idonee a evitare che determinati gruppi sociali siano sottorappresentati, o a rimuovere pregiudizi storici.

Per garantire l'equità dell'IA, il Codice Etico IA deve prevedere obblighi di spiegabilità, affinché ogni soggetto coinvolto possa comprendere come e perché un algoritmo ha preso una determinata decisione. Questo significa:

- Predisporre interfacce utente accessibili, che forniscano informazioni chiare sulle decisioni IA.
- Garantire il diritto degli utenti a contestare decisioni algoritmiche, con la possibilità di richiedere una revisione umana.
- Pubblicare report periodici sulla fairness dei sistemi IA, per garantire trasparenza e accountability.

La mitigazione del bias algoritmico non può essere un'azione occasionale, ma deve essere garantita attraverso un sistema di monitoraggio e auditing continuo predisposto dal codice etico., che deve anche prevedere la revisione indipendente da parte di esperti esterni, e Report periodici al Board of Directors, che devono includere analisi dei rischi di bias e azioni correttive adottate. Queste misure garantiscono che l'azienda mantenga un elevato livello di controllo sull'equità degli algoritmi, riducendo il rischio di discriminazioni involontarie.

Se un'azienda utilizza sistemi IA discriminatori, i rischi possono essere gravi e concreti:

- Azioni legali per discriminazione, con potenziali richieste di risarcimento elevate.

- Sanzioni normative, poiché l'AI Act e il GDPR prevedono multe fino al 6% del fatturato annuo globale per l'uso scorretto dell'IA.
- Danni reputazionali e perdita di fiducia, che possono compromettere il rapporto con clienti, investitori e istituzioni.

5. L'integrazione di sistemi di Intelligenza Artificiale nei processi aziendali impone un cambiamento radicale nella gestione delle responsabilità. A differenza di altri strumenti tecnologici, l'IA non è un mero supporto operativo, ma un agente decisionale che incide direttamente sulla vita delle persone, sulle strategie aziendali e sulla conformità normativa dell'impresa. Per questo motivo, il principio di accountability, ovvero la capacità di attribuire responsabilità chiare e verificabili per ogni decisione automatizzata, è un pilastro essenziale del Codice Etico IA.

Il concetto di responsabilità aziendale per l'uso dell'IA si declina su più livelli, coinvolgendo l'intero sistema di governance dell'impresa. Se un algoritmo assume decisioni scorrette, discriminatorie o non conformi alla normativa, chi ne risponde? Se un sistema IA genera un danno a un dipendente, un cliente o un partner commerciale, quali misure vengono adottate per gestire il rischio? E ancora, chi ha l'autorità di approvare o bloccare l'implementazione di un nuovo modello IA?

Queste domande non possono rimanere senza risposta. L'accountability dell'IA impone che ogni processo basato su algoritmi sia supervisionato da una struttura di governance chiara, con responsabilità definite e procedure di controllo rigorose.

6. Il principio di accountability è sancito da diverse fonti normative europee e internazionali, che pongono l'accento sulla necessità di identificare e attribuire ruoli e responsabilità specifiche nell'uso degli algoritmi.

1. AI Act Europeo: impone obblighi stringenti di tracciabilità e supervisione per i sistemi IA ad alto rischio, prevedendo che i responsabili aziendali siano identificabili e responsabili delle decisioni algoritmiche.
2. GDPR (Regolamento UE 2016/679): stabilisce il principio di responsabilità del titolare del trattamento, imponendo obblighi di risk assessment e gestione delle decisioni automatizzate.
3. D.Lgs. 231/2001: introduce il concetto di responsabilità amministrativa dell'ente, che può essere esteso anche alle violazioni derivanti dall'uso scorretto dell'IA, se questa incide su ambiti critici come la gestione del personale, le operazioni finanziarie o la protezione dei dati.
4. Direttiva NIS2 (UE 2022/2555): impone alle aziende di dotarsi di piani di gestione del rischio cibernetico e misure di sicurezza per proteggere i sistemi IA da attacchi informatici.

Alla luce di questo quadro normativo, il Codice Etico IA deve garantire che l'azienda abbia strutture di controllo solide, capaci di prevenire, rilevare e correggere eventuali violazioni o errori legati all'uso dell'IA.

Affinché il principio di accountability sia effettivo, è necessario che il Codice Etico IA stabilisca un sistema di governance strutturato, in cui siano definiti:

- Ruoli e responsabilità per l'uso dell'IA, attribuiti a specifiche figure aziendali.

- Processi di validazione e approvazione degli algoritmi, che impediscano l'uso di IA non conforme agli standard etici e normativi.
- Meccanismi di monitoraggio e auditing, per garantire un controllo continuo sulle performance degli algoritmi.

L'attribuzione delle responsabilità deve seguire un modello chiaro e gerarchico, che preveda:

1. Il Board of Directors: ha la responsabilità ultima della governance algoritmica, approvando le strategie IA aziendali e verificando il rispetto delle normative.
2. L'AI Ethicist e il Comitato Etico IA: figure chiave per la supervisione delle implicazioni etiche e sociali dell'IA, con potere di bloccare modelli IA non conformi.
3. La funzione Compliance: responsabile della verifica normativa, della gestione del rischio e dell'attuazione delle policy di accountability.
4. I Data Scientists e gli sviluppatori IA: responsabili della trasparenza e della documentazione tecnica degli algoritmi, affinché siano verificabili e spiegabili.

Questa struttura consente di evitare che la responsabilità sia dispersa o non attribuibile, creando un sistema di controllo rigoroso e in grado di rispondere a eventuali contestazioni legali o etiche.

7. Per garantire la responsabilità aziendale nell'uso dell'IA, il Codice Etico deve prevedere meccanismi di verifica preventiva e di monitoraggio continuo, che impediscano l'adozione di algoritmi senza adeguati controlli.

- Ogni nuovo modello IA deve essere sottoposto a una valutazione del rischio, che includa analisi di conformità normativa, etica e tecnica.

- Le decisioni automatizzate devono essere tracciabili e documentate, in modo che possano essere riesaminate in caso di errori o reclami.

- Devono essere previsti controlli periodici, con revisioni annuali sugli algoritmi in uso per verificare che non abbiano generato bias, discriminazioni o malfunzionamenti.

Il monitoraggio non deve limitarsi agli aspetti tecnici, ma deve prevedere un coinvolgimento attivo di tutti gli stakeholder, garantendo che lavoratori, clienti e utenti abbiano canali di comunicazione aperti per segnalare problemi legati all'IA.

L'accountability non può esistere senza un sistema di enforcement chiaro. Il Codice Etico IA deve prevedere conseguenze concrete per chi utilizza l'IA in modo non conforme, con un sistema di sanzioni proporzionato alla gravità della violazione.

- Errori minori o violazioni procedurali: possono essere sanzionati con richiami formali e obbligo di formazione aggiuntiva per i responsabili.

- Violazioni gravi della normativa: possono portare alla sospensione di un algoritmo, alla revoca delle credenziali di accesso o a sanzioni disciplinari nei confronti dei responsabili.

- Danni a terzi o violazioni legali: in caso di uso illecito dell'IA, l'azienda deve prevedere piani di risarcimento e la collaborazione con le autorità di vigilanza per mitigare le conseguenze.

L'azienda deve, inoltre, adottare un sistema di whistleblowing dedicato all'IA, che permetta a dipendenti e stakeholder di segnalare in modo sicuro eventuali usi scorretti degli algoritmi.

Le normative europee e internazionali riconoscono il ruolo critico della formazione e della competenza nella gestione dell'IA. Tra i principali riferimenti legislativi:

1. AI Act Europeo

- Stabilisce che le aziende che utilizzano IA ad alto rischio devono formare adeguatamente il personale responsabile della gestione e supervisione degli algoritmi.

2. Direttiva UE sulla Due Diligence Aziendale (CSDDD)

- Impone che le aziende adottino programmi di formazione obbligatoria sulla governance dell'IA, con particolare attenzione all'etica algoritmica.

Queste normative rendono evidente che la formazione sull'IA non è un elemento opzionale, ma un obbligo normativo e strategico, necessario per garantire che i sistemi IA siano gestiti in modo trasparente, equo e sicuro.

8. La necessità di garantire la presenza di un controllo umano adeguato sulle decisioni IA è sancita da diverse normative e regolamenti, che pongono obblighi chiari per le imprese:

1. AI Act Europeo: stabilisce che tutti i sistemi IA ad alto rischio devono essere sottoposti a supervisione umana, prevedendo che le decisioni automatizzate siano monitorate da personale qualificato, con la possibilità di intervenire o correggere le scelte dell'algoritmo.

2. GDPR (Regolamento UE 2016/679), Art. 22: vieta decisioni automatizzate che abbiano un impatto significativo sulle persone senza l'intervento umano adeguato e la possibilità di revisione manuale delle decisioni.

3. Raccomandazioni OCSE sull'IA: stabiliscono che ogni sistema IA deve avere un meccanismo di "human-in-the-loop" o "human-on-the-loop", per garantire che l'uomo rimanga il principale responsabile delle decisioni.

4. Direttiva sull'Equità Salariale UE (2023): impone trasparenza nell'uso dell'IA per la gestione delle risorse umane, con la necessità di supervisione nelle decisioni di assunzione, valutazione e licenziamento dei lavoratori.

5. Regolamenti settoriali in ambito sanitario e finanziario: stabiliscono che gli algoritmi utilizzati per la diagnosi medica o la concessione di credito non possono operare senza una verifica umana effettiva.

Alla luce di queste disposizioni, il Codice Etico IA deve garantire che ogni decisione IA che impatta persone o processi aziendali critici sia sempre supervisionata da un essere umano competente e responsabile.

L'implementazione della supervisione umana nei sistemi IA può avvenire attraverso tre diversi modelli, che devono essere adottati in base al livello di rischio e alla sensibilità del processo decisionale:

1. Human-in-the-loop (HITL)

- L'essere umano partecipa attivamente al processo decisionale dell'IA, intervenendo prima che l'algoritmo prenda una decisione finale.
- Esempio: un sistema di selezione del personale basato su IA può suggerire candidati idonei, ma la decisione finale spetta a un recruiter umano.
- Utilizzato in ambiti ad alto rischio, come sanità, finanza e HR.

2. Human-on-the-loop (HOTL)

- L'IA prende decisioni in autonomia, ma un operatore umano monitora e può intervenire in caso di anomalie.
- Esempio: un sistema IA per il riconoscimento facciale negli aeroporti identifica automaticamente i passeggeri, ma gli agenti di sicurezza possono verificare e correggere eventuali errori.
- Usato per processi ad alto volume, dove un controllo umano continuo sarebbe inefficiente.

3. Human-in-command (HIC)

- L'IA opera con ampia autonomia, ma l'essere umano mantiene il pieno controllo strategico e può interrompere il funzionamento del sistema in qualsiasi momento.
- Esempio: sistemi di IA in ambito militare o industriale, che possono essere disattivati o modificati da un operatore umano in caso di malfunzionamento.
- Applicato in ambiti critici con rischio di danni gravi.

Il Codice Etico IA deve definire chiaramente quale modello di supervisione umana si applica a ciascun sistema IA aziendale, in base a criteri di rischio, impatto sociale e conformità normativa.

Affinché la supervisione umana sia efficace, il Codice Etico IA deve prevedere procedure operative dettagliate, che consentano ai supervisori di monitorare e correggere le decisioni IA in modo sistematico.

1. Deve essere identificato un responsabile umano per ogni sistema IA in uso in azienda.

- Il supervisore deve avere accesso agli strumenti di monitoraggio e revisione delle decisioni algoritmiche.

2. Se un algoritmo prende una decisione errata, il supervisore deve avere il potere di annullarla o modificarla.

- Tutti gli interventi umani devono essere documentati e tracciati, per garantire la trasparenza del processo.

3. Test di affidabilità e audit continuo

- Devono essere condotti test periodici per verificare che i supervisori siano in grado di individuare e correggere errori IA.

- Devono essere previsti audit esterni e interni, con report periodici al Comitato Etico IA e al Board of Directors.

Se un'azienda utilizza sistemi IA senza un controllo umano adeguato, i rischi possono essere elevati:

- Errori gravi nelle decisioni aziendali, con conseguenti danni economici o reputazionali.

- Violazioni normative e sanzioni, in caso di mancato rispetto degli obblighi previsti dall'AI Act e dal GDPR.

- Perdita di fiducia da parte di dipendenti, clienti e stakeholder, con impatti negativi sulla reputazione aziendale.

9. L'Intelligenza Artificiale sta modificando il modo in cui il lavoro viene svolto, monitorato e valutato, incidendo direttamente sul rapporto tra datore di lavoro e lavoratore. Se da un lato essa offre maggiore efficienza operativa e riduzione delle attività ripetitive, dall'altro pone sfide etiche e giuridiche fondamentali per la tutela della dignità, dell'autenticità e dell'integrità della prestazione lavorativa.

Il principio di integrità e genuinità della prestazione lavorativa stabilisce due concetti chiave:

1. L'IA non può sostituire completamente la prestazione umana, né ridurre il lavoratore a mero esecutore di istruzioni algoritmiche, privandolo della sua autonomia, competenza e responsabilità.

2. Il lavoratore non può delegare la propria prestazione all'IA, simulando un contributo personale quando, in realtà, l'attività viene svolta interamente da un sistema algoritmico.

Il Codice Etico IA deve quindi includere regole chiare e vincolanti per l'uso dell'IA nel mondo del lavoro, affinché essa sia uno strumento di supporto e non un surrogato della competenza umana.

3.10.1. Il quadro normativo e regolamentare

L'adozione dell'IA nei rapporti di lavoro pone questioni giuridiche complesse, che coinvolgono diritti fondamentali del lavoratore, obblighi del datore di lavoro, regolamentazione della trasparenza e responsabilità aziendale.

L'uso dell'IA per la gestione del personale, il monitoraggio della produttività e la valutazione delle performance è soggetto a una fitta rete normativa, che include normative europee, nazionali e internazionali, con particolare attenzione alle seguenti aree:

1. Diritti fondamentali del lavoratore e protezione della dignità professionale.
2. Obblighi di trasparenza e informazione del datore di lavoro.
3. Supervisione umana e diritto alla contestazione delle decisioni IA.
4. Tutela del lavoratore contro l'automazione illecita e la sostituzione impropria della prestazione lavorativa.

Le normative più rilevanti in materia includono:

1. AI Act Europeo e il trattamento dell'IA nel mondo del lavoro

L'AI Act Europeo, attualmente in fase di approvazione finale, introduce per la prima volta un quadro normativo vincolante sull'uso dell'IA nei rapporti di lavoro.

Classificazione delle IA ad alto rischio

L'AI Act definisce ad alto rischio i sistemi IA impiegati per:

- Selezione e assunzione del personale.
- Valutazione delle performance e promozioni.
- Decisioni relative alla retribuzione, incentivi o licenziamenti.

Per queste applicazioni, il regolamento impone obblighi stringenti, tra cui:

- Supervisione umana obbligatoria sulle decisioni IA.
- Obbligo di documentazione e tracciabilità delle decisioni.

- Trasparenza nei criteri di valutazione algoritmica.
- Audit periodici per verificare equità e affidabilità dei modelli.

Criticità giuridiche aperte

- L'AI Act non fornisce una regolamentazione dettagliata sull'uso dell'IA da parte dei lavoratori, lasciando spazio a interpretazioni nazionali.
- Manca una definizione chiara sulla responsabilità del lavoratore che utilizza l'IA per svolgere la propria mansione, lasciando un vuoto normativo sul tema della delega impropria della prestazione lavorativa.

2. Il GDPR e la protezione dei dati personali nei rapporti di lavoro

Il GDPR (Regolamento UE 2016/679) gioca un ruolo fondamentale nella regolamentazione dell'IA in ambito lavorativo, imponendo limiti e obblighi sulla gestione dei dati personali nei processi automatizzati.

Articolo 22: Divieto di decisioni esclusivamente automatizzate

- Nessun lavoratore può essere soggetto a decisioni interamente automatizzate senza una base giuridica chiara e senza possibilità di contestazione.
- Il datore di lavoro è obbligato a informare il dipendente su come i dati personali vengono trattati dagli algoritmi IA.
- Il lavoratore ha il diritto di ottenere un intervento umano, richiedendo una revisione manuale delle decisioni IA.

Protezione della privacy e limiti al monitoraggio del lavoratore

- Il GDPR vieta forme di sorveglianza occulta, imponendo che ogni strumento IA che monitori i dipendenti sia proporzionato e trasparente.
- I lavoratori devono essere informati preventivamente su qualsiasi attività di monitoraggio IA che riguardi la loro prestazione.

In Italia, l'uso dell'IA nel monitoraggio dei lavoratori è regolato dallo Statuto dei Lavoratori (L. 300/1970), che impone limiti stringenti alla sorveglianza tecnologica e ai sistemi di controllo a distanza. Questa norma:

- Vieta l'uso di strumenti tecnologici per il controllo a distanza dei lavoratori senza il previo accordo con i sindacati o l'autorizzazione dell'Ispettorato del Lavoro.

- Questo divieto si applica anche agli strumenti IA utilizzati per monitorare la produttività, le performance o il comportamento dei dipendenti.

- Sentenze recenti confermano che l'uso di IA nei processi di selezione del personale o nella valutazione delle prestazioni deve rispettare i principi di trasparenza e giustificabilità.

- È stato ribadito che il datore di lavoro è sempre responsabile per decisioni IA che influiscono sulla carriera dei dipendenti, anche se il processo decisionale è automatizzato.

L'Organizzazione Internazionale del Lavoro (OIL) ha stabilito una serie di principi per garantire che l'automazione e l'IA non erodano i diritti fondamentali dei lavoratori.

Le decisioni di licenziamento basate non possono basarsi sul monitoraggio dei lavoratori e su analisi algoritmiche, devono essere sempre processate e trattate da un essere umano secondo il principio già discusso di verticalità, per evitare discriminazioni e decisioni arbitrarie.

10. L'integrazione dell'Intelligenza Artificiale nelle attività lavorative può generare una serie di rischi per l'integrità della prestazione lavorativa, minando il rapporto di fiducia tra lavoratore e datore di lavoro e modificando il concetto stesso di lavoro personale e autentico.

Se da un lato l'IA può ottimizzare i processi, aumentare la produttività e ridurre il carico di lavoro, dall'altro vi è il pericolo che:

1. I lavoratori abusino dell'IA per delegare occultamente la propria prestazione.
2. I datori di lavoro utilizzino l'IA per esercitare un controllo eccessivo e invadente sulla prestazione lavorativa.
3. L'IA possa sostituire gradualmente il lavoro umano, con un impatto negativo sulla qualità dell'occupazione.

In questo contesto, è fondamentale identificare e regolamentare i rischi per garantire che l'IA non alteri la natura genuina del rapporto di lavoro.

1. Abuso dell'IA da parte del lavoratore: la delega occulta della prestazione lavorativa

Uno dei rischi principali riguarda l'uso improprio dell'IA da parte del lavoratore, che potrebbe sfruttare strumenti di automazione per delegare interamente la propria mansione all'algoritmo, senza che il datore di lavoro ne sia consapevole, come: per es., fenomeni di elusione della prestazione lavorativa,

Se da un lato vi è il rischio di abuso dell'IA da parte del lavoratore, dall'altro vi è il pericolo opposto, ovvero l'uso dell'IA da parte del datore di lavoro per monitorare eccessivamente la prestazione lavorativa, con conseguente perdita di autonomia decisionale.

Alcuni strumenti IA possono essere impiegati per monitorare in tempo reale l'attività lavorativa, tra cui:

- Software di productivity tracking, che analizzano il numero di e-mail inviate, documenti redatti e interazioni digitali del dipendente.
- Analisi comportamentale basata su IA, che valuta pause, velocità di digitazione e tempo impiegato per ogni task.
- Sistemi di riconoscimento facciale e registrazione audio, utilizzati per garantire la “presenza attiva” del lavoratore nei processi digitali.

I problemi giuridici legati alla sorveglianza digitale sono ovviamente riconducibili alla violazione dell'articolo 4 dello Statuto dei Lavoratori, che vieta controlli a distanza sui dipendenti., come pure alla violazione del GDPR, in quanto il monitoraggio continuo deve essere proporzionato e giustificato.

L'Intelligenza Artificiale è una tecnologia con un impatto trasformativo sulle strutture aziendali, capace di ridefinire processi decisionali, produttività e gestione delle risorse umane. Tuttavia, la sua adozione deve avvenire in un quadro regolamentato e strutturato, affinché i principi di legalità, equità, trasparenza e genuinità della prestazione lavorativa non restino mere enunciazioni, ma si traducano in pratiche operative concrete.

Per garantire che l'IA non comprometta la qualità della prestazione lavorativa, la dignità del lavoratore o l'affidabilità delle decisioni aziendali, è necessario che il Codice Etico IA sia supportato da un solido sistema di governance, capace di monitorare, valutare e adattare continuamente l'uso dell'IA nell'organizzazione.

11. Si sono già menzionati i principi fondamentali della governance algoritmica. Si accenna di seguito, sul piano degli attori della governance dell'IA., ad una figura specializzata che, di fronte all'incalzare delle normative (AI Act, legislazioni statali in USA, linee guida in Cina, normative emergenti in Spagna, Singapore e altrove) e all'evoluzione tecnologica, molte imprese avvertono l'esigenza di creare. Si tratta dell'AI Ethicist o Ethics Officer per l'IA: un professionista che conosca sia la parte normativa (GDPR, AI Act, PIPL cinese, PDPA di Singapore, LOPDGGD spagnola, etc.) sia i fondamenti di machine learning, fairness, explainability e governance societaria; caratterizzato da una funzione ibrida: si colloca a metà tra il dipartimento Legal/Compliance e l'area HR/IT, poiché deve valutare l'impatto etico sulla forza lavoro e la conformità legale.

12. Europa è il primo mercato al mondo con una legislazione dedicata alla tutela, alla mitigazione e all'annullamento dei rischi derivanti dall'uso improprio dell'AI ed interviene anche su temi di avanguardia tecnologica, tra i quali il riconoscimento visivo, emotivo, la manipolazione dei comportamenti, la vigilanza predittiva, il punteggio sociale e individuale. L'implementazione di un sistema di governance algoritmica solido è una condizione imprescindibile per garantire un uso etico e sostenibile dell'IA.

L'IA deve essere regolata da un Codice Etico IA chiaro e vincolante, integrato nei processi aziendali.

La governance algoritmica deve essere gestita da figure specializzate, come l'AI Ethicist e il Comitato Etico IA.

La genuinità della prestazione lavorativa deve essere tutelata, prevenendo l'uso improprio dell'IA da parte di lavoratori e datori di lavoro.

L'IA deve essere soggetta a revisione, auditing e aggiornamento continuo, per garantirne l'affidabilità e l'allineamento alle normative vigenti.

Una governance algoritmica responsabile non è solo un obbligo normativo, ma una scelta strategica, che permette alle aziende di innovare in modo etico e sostenibile, rafforzando la fiducia di clienti, investitori e lavoratori.

La governance algoritmica in tal modo diviene uno strumento reale di intervento dalla autentica portata operativa ed anticipa nella realizzazione quanto delineato dall'ONU nel Report *"Governare l'AI per l'umanità"* curato da un Advisory Board multispecialistico di competenze aziendali e istituzionali. Le conclusioni dell'Advisory Board sono per alcune indicazioni e per una sola direttiva: rendere la gestione della tecnologia dell'AI sia nell'impiego, sia nello sviluppo responsabile, equa e giusta.