

VALERIA BARONE

Il problema della discriminazione “di precisione”

Abstract

The article provides a critical examination of the concept of “algorithmic decision-making” and its implications for the legal field, with particular attention to its impact on the foundational principles of the legal system. It foregrounds the uniquely human character of law as an artifact of linguistic and intellectual creation. Central to the legal system is the act of “articulating” the law – a practice inherently reliant on human judgment, interpretive reasoning, and the contextualization of facts and norms. While advanced technologies, such as natural language processing (NLP) and computational systems, enable machines to generate and process language, these systems lack the inherent capacity for genuine comprehension or autonomous decision-making. Their outputs are mere simulations produced through statistical models and pre-defined algorithms, devoid of intentionality and deeper understanding.

The article situates this discussion within a historical framework, tracing the origins of computational law to Leibniz’s vision of a *mathesis universalis*, thereby engaging with long-standing debates regarding the computability of legal reasoning. It further explores contemporary applications, such as smart contracts and autonomous vehicles, while critically interrogating whether the translation of legal norms into technical commands undermines the normative and ethical essence of the law. In conclusion, the article argues that the delegation of judicial decision-making to algorithms risks eroding the human essence of the legal system. Genuine legal reasoning requires core human attributes, such as empathy, contextual discernment, and a commitment to justice – qualities inherently absent in algorithmic processes. The reduction of law to a mechanistic system threatens its integrity as a human construct, designed not merely to regulate behavior but to navigate the complexities of human society and strive toward the perpetually elusive ideal of justice.

Keywords

- Algorithmic decision-making
- Computational law
- Legal interpretation
- Artificial intelligence and law
- Human-centered justice.

Sommario

1. Premessa
2. Verso una discriminazione “di precisione”
3. Sistema giudiziario. L’esemplarità del caso Loomis
4. Tutela della salute
5. Settore finanziario, bancario e assicurativo
6. Lavoro e selezione del personale
7. Conclusioni

1. Premessa

Il presente articolo si propone di denunciare le nuove forme di discriminazione rese possibili dall'intelligenza artificiale e dalle tecnologie digitali, mettendo in luce le implicazioni etiche, giuridiche e sociali di questi fenomeni. L'obiettivo è mostrare come l'uso massivo di dati e di algoritmi possa rafforzare e persino perfezionare le disuguaglianze esistenti, rendendole più difficili da individuare e contrastare.

Il saggio inizia proponendo l'introduzione dell'espressione "discriminazione di precisione" al fine di individuare le peculiarità delle più recenti forme di discriminazione rispetto a quelle "tradizionali", basate su categorie come razza, genere o età. Le forme di discriminazione finora manifestate consistevano (e consistono) in trattamenti differenziati applicati a interi gruppi sociali. Quello che oggi si sta profilando è un processo ancora più insidioso, perché prende di mira direttamente l'individuo sulla base di correlazioni statistiche elaborate dagli algoritmi. Il problema emerge in modo particolarmente evidente nei settori delle assicurazioni, del credito, del lavoro e del marketing, dove l'analisi predittiva non si limita più a valutazioni generali, ma costruisce profili dettagliati, influenzando in modo diretto l'accesso a beni e servizi essenziali. L'uso sempre più pervasivo dell'IA per stabilire tariffe assicurative, concedere prestiti, selezionare candidati per un impiego o personalizzare le offerte commerciali, crea una rete di esclusione silenziosa, in cui chi viene penalizzato non sempre ha la possibilità di comprendere o contestare le ragioni della propria esclusione.

L'analisi si concentra poi sul funzionamento degli algoritmi e sul modo in cui essi riproducono e amplificano i *bias*. Ad esempio, se i dati utilizzati nell'addestramento del modello algoritmico contengono pregiudizi, questi vengono assimilati e replicati in modo sistematico, aggravando ulteriormente le disuguaglianze. L'idea che l'intelligenza artificiale sia uno strumento neutrale e imparziale si rivela così un'illusione pericolosa: gli algoritmi non sono entità autonome, ma il risultato di scelte umane che determinano cosa misurare, come classificare le informazioni e quali criteri adottare per prendere decisioni. Affidarsi ciecamente a tali strumenti senza un adeguato controllo significa legittimare un sistema in cui la tecnologia diventa il veicolo di discriminazioni sempre più raffinate e invisibili.

Particolare attenzione è dedicata alle implicazioni di questo fenomeno per i diritti fondamentali, in particolare per la privacy, la protezione dei dati personali, l'autonomia decisionale e l'accesso equo alle opportunità. La progressiva automatizzazione dei processi decisionali rischia di privare gli individui del controllo sulle proprie scelte, riducendoli a semplici insiemi di dati analizzati da sistemi opachi e inaccessibili. L'assenza di trasparenza e accountability impedisce non solo di individuare eventuali distorsioni nei modelli predittivi, ma anche di opporsi a decisioni ingiuste. La questione non è solo tecnica, ma eminentemente politica: lasciare che questi strumenti operino senza una regolamentazione adeguata equivale ad accettare un modello di società in cui il potere decisionale si concentra nelle mani di pochi attori privati, con conseguenze potenzialmente devastanti per la giustizia sociale.

L'ultima parte del saggio è dedicata alla necessità di un intervento regolatorio che garantisca maggiore giustizia nell'uso dell'intelligenza artificiale. Il problema non è l'innovazione in sé, ma il modo in cui essa viene implementata e governata. È necessario affermare con forza che il progresso tecnologico non può giustificare la violazione dei diritti fondamentali, né può diventare un pretesto per legittimare nuove forme di esclusione sociale. L'assenza di un quadro normativo adeguato rischia di lasciare spazio a dinamiche di potere incontrollate, in cui pochi attori dominano il mercato delle informazioni e impongono regole opache e inaccessibili.

La battaglia contro la discriminazione di precisione non è solo una sfida tecnica, ma una questione di giustizia e democrazia. Rifiutare l'idea di un futuro governato da algoritmi incontrollabili significa riaffermare la centralità dell'essere umano e la necessità di difendere con determinazione i principi di uguaglianza e libertà.

2. Verso una discriminazione “di precisione”

L'adozione dei più recenti sistemi di intelligenza artificiale ha dato origine a nuove e insidiose forme di discriminazione, che possiamo definire come “discriminazioni di precisione”. A differenza delle discriminazioni tradizionali, che si basano sull'appartenenza a categorie generali come razza, genere, età, religione o disabilità, la discriminazione di precisione non opera su gruppi ampi e identificabili, ma su individui specifici, sfruttando

la mole di dati personali oggi disponibile per prendere decisioni mirate, calibrate su ogni singolo soggetto. Il risultato è una forma di esclusione ancora più subdola e difficile da individuare, perché non si manifesta attraverso trattamenti palesemente differenziati, ma tramite scelte algoritmiche “invisibili” che penalizzano alcuni individui sulla base di correlazioni statistiche, senza che essi possano rendersene conto o difendersi.

L’espressione “discriminazione di precisione” nasce per sottolineare come questa logica sia il riflesso di un modello più ampio che permea diversi settori: la medicina di precisione utilizza dati genetici e comportamentali per personalizzare i trattamenti sanitari; l’agricoltura di precisione impiega analisi dettagliate su suolo e clima per ottimizzare le rese; il marketing di precisione studia i comportamenti dei consumatori per plasmare strategie pubblicitarie su misura; l’ingegneria di precisione lavora con tolleranze minime, ovvero margini di errore ridottissimi, misurabili in micrometri o nanometri, per garantire massima accuratezza nella produzione industriale; la chirurgia di precisione consente interventi mirati con minime ripercussioni sui tessuti; la logistica di precisione ottimizza trasporti e stoccaggio con algoritmi predittivi.

Se in questi ambiti la precisione è un valore positivo, nel caso delle discriminazioni il discorso cambia radicalmente. Qui, l’intelligenza artificiale non sta ottimizzando il benessere, ma sta perfezionando l’ingiustizia, rendendo la discriminazione più raffinata e meno individuabile. Chi viene penalizzato da un algoritmo non sempre sa di esserlo né può contestarlo: l’accesso a un mutuo, la selezione per un impiego, il costo di una polizza assicurativa possono essere influenzati da variabili apparentemente neutre che, senza alcuna giustificazione esplicita, escludono o penalizzano determinati individui. Ciò rappresenta una nuova forma di disegualianza sistematica, invisibile e incontrastabile, che sfida le tradizionali nozioni di giustizia e *accountability*. Alimentati da una mole crescente di dati personali, gli algoritmi di intelligenza artificiale, elaborano profili estremamente dettagliati che vanno ben oltre le categorie identitarie classiche, prendendo di mira scelte, abitudini e stili di vita delle persone. La discriminazione non si basa più su tratti evidenti e riconoscibili, ma su correlazioni invisibili, sulle minime variazioni nei comportamenti quotidiani¹.

¹ Sul tema della discriminazione digitale, la lettura è ormai sterminata. Segnalo qui almeno: S.U. Noble, *Algorithms of Oppression. How Search Engines Reinforce Racism*, NYU University Press, New York, 2018; J. Kleinberg, J. Ludwig, S. Mullainathan, C.R. Sunstein, *Discrimination in the age of algorithm*, in *Journal of Legal Analysis*, 10, 2018,

Questo meccanismo è già in atto in diversi settori chiave. Nel campo delle assicurazioni, ad esempio, le compagnie analizzano i dati relativi alla salute, alle abitudini di consumo e persino ai comportamenti di guida per personalizzare i premi assicurativi. In teoria, questa strategia potrebbe essere presentata come una forma di valutazione "oggettiva", ma in realtà si traduce in un meccanismo di esclusione sistematica: chi ha stili di vita ritenuti "a rischio" (come fumatori, persone con condizioni mediche specifiche o residenti in determinate aree) viene penalizzato con costi più elevati, senza un effettivo margine di negoziazione o di tutela.

Nel settore finanziario, il problema si acuisce. Gli algoritmi che valutano l'accesso a prestiti e carte di credito non si limitano più a esaminare lo storico finanziario, ma tracciano ogni aspetto del comportamento economico, dalla frequenza degli acquisti alla tipologia dei negozi frequentati, fino agli orari di spesa. Si crea così un circolo vizioso per cui chi parte da una posizione di svantaggio economico viene progressivamente spinto ai margini, con tassi di interesse più elevati e minori possibilità di accesso al credito.

Anche nel mondo del lavoro, l'uso delle tecnologie di selezione automatizzata del personale rischia di normalizzare una nuova forma di discriminazione. I software di *screening* utilizzati da molte aziende per valutare i candidati non si basano solo sulle competenze e sull'esperienza, ma su elementi opachi che possono includere gli account social, le interazioni sulla rete, le attività online e la condivisione di contenuti. Questo significa che un candidato può essere scartato senza nemmeno sapere il perché, solo perché il suo profilo non si adatta a un modello predefinito dall'algoritmo.

Ancora: il settore del marketing digitale dimostra come la discriminazione di precisione possa manipolare profondamente la percezione e le scelte delle persone. Il celebre caso della catena di supermercati Target, negli Stati Uniti, è emblematico: analizzando i pattern di acquisto, l'azienda è riuscita a identificare donne incinte prima ancora che

pp. 113-174; C.R. Sunstein, *Algorithms, correcting biases*, in *Social Research*, 86, 2, 2019, pp. 499-511; R. Xenedis, L. Senden, *EU Non-Discrimination Law in the Era of Artificial Intelligence: Mapping the Challenges of Algorithmic Discrimination*, in U. Bernitz et al. (eds.), *General Principles of EU law and the EU Digital Order*, Wolters Kluwer, Alphen aan den Rijn, 2020, pp. 151-182; R. Nunn, *Discrimination in the Age of Algorithms*, in W. Barfield (ed.), *The Cambridge Handbook of the Law of Algorithms*, Cambridge, Cambridge University Press, 2020, pp. 182-198. Per quanto riguarda la letteratura in lingua italiana: E. Falletti, *Discriminazione algoritmica: una prospettiva comparata*, Torino, Giappichelli, 2022; B.G. Bello, *(In)giustizie digitali. Un itinerario su tecnologie e diritti*, Pisa, Pacini, 2023, in particolare pp. 73-86; F. Casa, *Il filosofo del diritto e le discriminazioni digitali*, in *Ordines. Per un sapere interdisciplinare delle istituzioni europee*, 2, 2023, pp. 228-246.

queste lo avessero comunicato a familiari e amici, inviando loro offerte personalizzate per prodotti legati alla gravidanza². Questo episodio dimostra non solo il potere predittivo dell'intelligenza artificiale, ma anche la sua capacità di violare la privacy in modi che i consumatori stessi non riescono a controllare.

Questi esempi confermano come la discriminazione di precisione non sia solo un'evoluzione della discriminazione tradizionale, ma una sua forma raffinata, invisibile e sistemica, che mette a rischio principi fondamentali come l'uguaglianza e l'autodeterminazione.

La pervasività dell'IA nella vita quotidiana sta creando gerarchie di accesso e opportunità che non sono più basate su criteri trasparenti, ma su decisioni automatizzate, difficilmente contestabili. Davanti a questa trasformazione, è necessario un intervento urgente per garantire che le tecnologie digitali non diventino strumenti di esclusione e di concentrazione del potere nelle mani di pochi, ma che siano sottoposte a un controllo democratico e a una regolamentazione rigorosa che metta al centro i diritti delle persone. Le tecnologie digitali, se lasciate senza adeguata regolamentazione e controllo, rischiano di diventare strumenti di discriminazione sistematica anziché di progresso.

Per arginare questo rischio, è necessario ripensare l'intero paradigma dell'intelligenza artificiale, trovando un equilibrio tra l'efficienza delle nuove tecnologie e la tutela dei diritti e delle libertà individuali. Gli algoritmi riproducono fedelmente i dati e i modelli su cui vengono addestrati, ma se questi dati contengono *bias*, essi non faranno altro che amplificare pregiudizi e disuguaglianze. La discriminazione algoritmica è il prodotto di scelte umane, di un sistema che privilegia certe logiche rispetto ad altre.

L'idea che gli algoritmi siano strumenti infallibili e imparziali è una pericolosa illusione. Se addestrati su dati distorti o parziali, assorbiranno e riprodurranno le stesse discriminazioni del mondo reale. Tuttavia, a differenza delle macchine, gli esseri umani hanno il potere – e la responsabilità – di intervenire. Possono scegliere di correggere le distorsioni, di sviluppare modelli più equi, di rendere i processi decisionali trasparenti e contestabili. Soprattutto, possono decidere di non affidarsi ciecamente agli algoritmi, ma di riaffermare il primato del controllo umano sulla tecnologia.

² Cfr. K. Hill, *How Target figured out a teen girl was pregnant before her father did*, in *Forbes*, Feb. 16, 2012.

3. Sistema giudiziario. L'esemplarità del caso Loomis

Nel contesto contemporaneo, l'uso degli algoritmi si sta diffondendo rapidamente, investendo anche i sistemi giudiziari. Tuttavia, la loro applicazione in questo ambito pone questioni critiche, soprattutto in relazione alla giustizia e ai diritti fondamentali³.

Uno degli utilizzi più controversi riguarda la previsione della probabilità che un individuo commetta reati futuri. Sebbene gli strumenti algoritmici siano spesso presentati come più efficienti e precisi rispetto al giudizio umano, numerosi studi hanno dimostrato che non solo possono perpetuare, ma addirittura amplificare le disuguaglianze sociali⁴. La pretesa di oggettività matematica maschera il rischio concreto di decisioni discriminatorie, basate su modelli che riflettono i pregiudizi preesistenti nel sistema giudiziario e nella società. Affidarsi a tali strumenti senza un adeguato controllo critico significa legittimare un'automazione dell'ingiustizia, che rafforza stereotipi e penalizza in modo sistematico gruppi già vulnerabili⁵.

La spinta verso la sostituzione dei giudici umani con algoritmi si fonda su tre motivazioni principali: ridurre il carico di lavoro degli esseri umani, migliorare l'efficienza del sistema giudiziario e garantire una maggiore certezza del diritto. Tuttavia, prima ancora di discutere le modalità di implementazione di tale cambiamento, è cruciale chiedersi se questo obiettivo sia realmente auspicabile. L'efficienza, spesso intesa come rapidità e contenimento dei costi, è un valore che non può essere disgiunto dalla complessità del ruolo giudiziario, che non si esaurisce nella mera applicazione meccanica di disposizioni normative, ma richiede interpretazione, contestualizzazione e capacità di discernere il peso delle circostanze specifiche. La pretesa di sostituire la deliberazione umana con la logica algoritmica rischia di ridurre la giustizia a un esercizio puramente computazionale, cancellando la dimensione etica e discrezionale che caratterizza il giudizio umano⁶.

I giudici, in quanto esseri umani, possono certo incorrere in decisioni incoerenti, influenzate da fattori estranei alla sostanza del caso. L'idea di affidarsi agli algoritmi nasce

³ A.G. Ferguson, *The rise of big data policing: surveillance, race, and the future of law enforcement*, New York University Press, New York, 2020.

⁴ R. Richardson, J. Schultz, K. Crawford, *Dirty Data, Bad Predictions: How Civil Rights Violations Impact Police Data, Predictive Policing Systems, and Justice*, in *New York University Law Review*, 94, 2019, pp. 15-55.

⁵ Sul punto: S.G. Mayson, *Bias In, Bias Out*, in *The Yale Law Journal*, 128, 8, 2019, pp. 2218-2300.

⁶ J. Kleinberg, L. Himabindu, L. Jure, L. Jens, M. Sendhil, *Human Decisions and Machine Predictions*, in *Quarterly Journal of Economics*, 133, 1, 2018, pp. 237-293.

proprio dalla volontà di eliminare queste distorsioni, garantendo maggiore coerenza e neutralità nelle sentenze. Tuttavia, i risultati ottenuti finora smentiscono questa promessa: gli algoritmi, invece di correggere i pregiudizi, spesso li amplificano⁷.

Gli algoritmi sono strumenti estremamente sensibili alle scelte compiute nelle varie fasi di sviluppo: errori nella selezione dei dati, nella definizione delle variabili o nell'addestramento del modello possono portare a risultati distorti e discriminatori. Uno degli aspetti più critici riguarda la qualità dei dati utilizzati per il loro addestramento, che spesso riflettono le disuguaglianze e i pregiudizi già radicati nei sistemi giuridici e sociali. Inoltre, la rigidità degli algoritmi impedisce loro di cogliere tutte le variabili contestuali che un giudice umano sarebbe in grado di considerare, producendo decisioni meccaniche e prive di quella sensibilità che il giudizio richiede.

Nonostante queste criticità, in alcuni contesti sperimentali gli algoritmi hanno dimostrato di poter ridurre le disparità, ma solo quando sviluppati con estrema attenzione e monitorati costantemente. Un esempio positivo è il sistema di rilascio pre-processuale adottato a New York, dove un algoritmo ha contribuito ad aumentare i tassi di rilascio per tutti i gruppi razziali, riducendo al contempo il divario tra di essi. Questo dimostra che la tecnologia, se progettata in modo corretto e con adeguate garanzie di equità e trasparenza, potrebbe avere un ruolo nella lotta alle disuguaglianze. Tuttavia, l'uso degli algoritmi nel sistema giudiziario non può essere accettato acriticamente né considerato una panacea: senza controlli stringenti, il rischio è quello di automatizzare e rafforzare le ingiustizie, anziché correggerle⁸.

Occorre sottolineare che le decisioni giudiziarie non si riducono a semplici operazioni logico-deduttive, ma possiedono una dimensione aletica, ossia veritativa, che implica un'attività interpretativa complessa. Il diritto non è una scienza esatta e l'interpretazione giuridica non può essere ridotta a una mera applicazione automatica di regole predefinite: ogni decisione richiede di stabilire la gerarchia e il peso dei criteri interpretativi, tenendo conto del contesto normativo, dei principi generali e della specificità del caso concreto.

⁷ J. Ludwig, S. Mullainathan, *Fragile Algorithms and Fallible Decision-Makers: Lessons from the Justice System*, in *The Journal of Economic Perspectives*, 35, 4, 2021, pp. 71-96.

⁸ M. Zilka et al., *The Progression of Disparities within the Criminal Justice System: Differential Enforcement and Risk Assessment Instruments*, in *Proceedings of the 2023 ACM Conference on Fairness, Accountability, and Transparency*, Association for Computing Machinery, New York, 2023, pp. 1553-1569.

Se la precisione matematica dei sistemi di intelligenza artificiale può sembrare allettante, l'idea che possano garantire un'imparzialità assoluta è un'illusione. L'aspirazione a un giudice-robot trascura il fatto che il diritto non è solo una questione di deduzione logica, ma di interpretazione e di equilibrio tra principi spesso in tensione tra loro. La pretesa di sostituire il giudizio umano con un calcolo algoritmico rischia così di nascondere, dietro la presunta neutralità della macchina, i pregiudizi e le limitazioni di chi ne ha progettato il funzionamento. Il rischio di errore algoritmico non può essere eliminato, poiché nessuna tecnologia è infallibile. Gli algoritmi dipendono da hardware e software che possono presentare malfunzionamenti, da dati di input che possono essere incompleti, distorti o errati, e da criteri di progettazione che riflettono scelte umane inevitabilmente fallibili.

L'idea che la decisione algoritmica sia intrinsecamente più affidabile di quella umana ignora il fatto che il potere decisionale non viene trasferito alla macchina, ma semplicemente riposizionato, passando da giudici e operatori del diritto ai programmatori, ai data scientist e alle aziende che sviluppano i sistemi. La questione della responsabilità diventa quindi cruciale: chi risponde di un errore giudiziario quando la decisione è presa da un algoritmo? Il trasferimento dell'autorità decisionale dall'uomo alla tecnologia non elimina i pregiudizi, ma li rende più difficili da individuare e contestare. Un sistema di giustizia che accetta in modo passivo il giudizio algoritmico smette di essere giustizia: diventa un meccanismo opaco di automazione dell'ingiustizia⁹.

Uno dei problemi più gravi riguarda l'impatto razziale delle valutazioni algoritmiche. I sistemi predittivi utilizzati dalle forze dell'ordine si basano su variabili che, pur non includendo esplicitamente la razza, risultano fortemente correlate ad essa. Fattori come la storia criminale, il quartiere di residenza o la frequenza degli arresti passati vengono impiegati per costruire profili di rischio che, lungi dall'essere neutrali, perpetuano le discriminazioni strutturali già presenti nel sistema giudiziario. L'illusione di un processo decisionale "oggettivo" nasconde il fatto che questi strumenti non fanno altro che proiettare nel futuro le ingiustizie del passato.

I tentativi di correggere tali distorsioni si sono dimostrati inefficaci. Eliminare dai dati i fattori direttamente correlati con la razza, ricalibrare gli algoritmi per ottenere previsioni più eque tra i diversi gruppi, o persino rinunciare del tutto ai metodi algoritmici non ri-

⁹ D. Yeung, I. Khan, N. Kalra, O. Osoba, *Identifying Systemic Bias in the Acquisition of Machine Learning Decision Aids for Law Enforcement Applications*, RAND Corporation, Santa Monica (CA), 2021.

solgono il problema alla radice. In una società segnata da profonde disuguaglianze razziali, qualsiasi modello predittivo finisce inevitabilmente per riprodurre e amplificare tali squilibri. Il rischio algoritmico non è un'anomalia del sistema, ma una manifestazione visibile delle ingiustizie insite nel concetto stesso di predizione. Questa consapevolezza costringe a guardare oltre la questione tecnica e ad affrontare la radice del problema: un sistema di giustizia che, invece di correggere il passato, lo codifica e lo cristallizza nel futuro.

Un altro problema cruciale nell'utilizzo degli algoritmi nelle decisioni giudiziarie riguarda la gestione di set di dati limitati e di variabili estremamente complesse. Gli algoritmi richiedono una quantità significativa di dati per produrre previsioni affidabili, ma la scarsità di informazioni disponibili spesso compromette la loro capacità predittiva. Ad esempio, il numero di recidivi che commettono reati gravi dopo il rilascio dal carcere è relativamente basso, rendendo difficile costruire modelli statistici realmente accurati. Inoltre, fattori personali e contestuali, come il supporto familiare o la partecipazione a reti sociali e comunità religiose, influiscono in modo significativo sulla probabilità di recidiva, ma sono difficili da integrare in un sistema algoritmico. In questi casi, è essenziale che il giudizio umano possa intervenire, valutando tali elementi in modo causale e integrando la previsione algoritmica con una decisione più ampia e contestualizzata, che tenga conto delle reali condizioni di vita della persona interessata¹⁰.

Una proposta particolarmente interessante riguarda l'idea di impiegare una "competizione" tra algoritmi per individuare i modelli predittivi più efficaci e affidabili. Questo approccio, già utilizzato in altri ambiti dell'intelligenza artificiale, potrebbe consentire di confrontare diversi sistemi in termini di accuratezza, equità e capacità di generalizzazione, riducendo il rischio di *bias* e distorsioni. Tuttavia, affinché una simile competizione possa davvero migliorare la qualità degli strumenti predittivi, è essenziale stabilire criteri rigorosi che non si limitino alla mera performance numerica, ma includano valutazioni sulla trasparenza, sulla spiegabilità delle decisioni e sull'impatto etico e sociale. In assenza di tali garanzie, il rischio è che la competizione si riduca a un semplice confronto tecnico senza affrontare le reali criticità dell'uso degli algoritmi nel sistema giudiziario¹¹.

¹⁰ M.A. Wojcik, *Machine-learned bias? Algorithmic decision making and access to criminal justice*, in *Legal Information Management*, 20, 2, 2020, p. 99.

¹¹ S. Levmore, F. Fagan, *Competing Algorithms for Law: Sentencing, Admissions, and Employment*, in *The University of Chicago Law Review*, 88, 2, pp. 367-412.

Nel dibattito sull'uso degli algoritmi nel sistema giudiziario, il caso Loomis rappresenta un esempio emblematico delle criticità di questi strumenti¹². Nel 2013, Eric Loomis fu arrestato mentre si trovava alla guida di un'auto precedentemente utilizzata in una spauratoria. Durante il processo, il giudice fece ricorso al risultato proposto da COMPAS (*Correctional Offender Management Profiling for Alternative Sanctions*), un algoritmo predittivo impiegato per valutare il rischio di recidiva. Il sistema analizzava una serie di variabili, tra cui il contesto socioeconomico e il passato giudiziario dell'imputato, per stimare la probabilità che potesse commettere nuovi reati in futuro.

Questo caso ha sollevato questioni fondamentali sulla trasparenza e sull'equità dell'uso degli algoritmi in ambito giudiziario. Uno degli aspetti più critici riguardava l'impossibilità, per Loomis e i suoi avvocati, di contestare direttamente il funzionamento dell'algoritmo, poiché il modello di COMPAS era proprietario e non accessibile. Loomis ha quindi contestato l'uso di COMPAS, sostenendo che gli fosse stato negato l'accesso alle modalità di calcolo del punteggio assegnatogli, integrandosi così una violazione del suo diritto al giusto processo.

La Corte Suprema del Wisconsin ha stabilito che COMPAS poteva essere impiegato, purché non fosse l'unico elemento determinante nella sentenza. Tuttavia, questa decisione non ha risolto le più ampie preoccupazioni legate all'uso degli algoritmi nel sistema giudiziario. Il caso ha evidenziato una questione cruciale: la delega crescente delle decisioni giuridiche a strumenti algoritmici rischia di erodere le garanzie fondamentali del processo equo, rendendo la giustizia sempre più dipendente da sistemi opachi e inaccessibili alla sua messa in discussione.

Il sistema giuridico non è statico, ma si adatta alle trasformazioni sociali, economiche e culturali: un algoritmo, per quanto preciso, non possiede la capacità di innovare o di reinterpretare il diritto in risposta a nuove esigenze, limitandosi a riprodurre schemi consolidati. Questo comporterebbe il pericolo di una stagnazione giurisprudenziale, in cui il passato vincola rigidamente il futuro, impedendo qualsiasi sviluppo normativo progressista.

Anche qualora gli algoritmi potessero operare con assoluta esattezza, la loro capacità decisionale resterebbe confinata entro i limiti dei dati e delle istruzioni ricevute. Il diritto,

¹² Su questo *leading case* sono numerosissimi i contributi sviluppati da studiosi e studiose di tutto il mondo. Mi limito qui a ricordare l'importante contributo italiano A. Simoncini, S. Suweis, *Il cambio di paradigma nell'intelligenza artificiale e il suo impatto sul diritto costituzionale*, in *Rivista di filosofia del diritto*, VIII, fasc. 1, 2019, pp. 87-106.

però, è fatto anche di concetti indeterminati, di margini di discrezionalità, di casi eccezionali che sfuggono a una classificazione rigida. La flessibilità e la creatività del giudice umano sono elementi imprescindibili per affrontare le sfumature del ragionamento giuridico e adattare le norme a situazioni impreviste. L'aspirazione a decisioni perfettamente prevedibili si scontra con l'incertezza intrinseca del mondo reale, dove il diritto non si applica meccanicamente, ma si costruisce attraverso un confronto critico con la realtà.

Di fronte a questa complessità, la fiducia e la legittimazione di un giudice umano, con tutti i suoi limiti, restano preferibili alla presunta infallibilità di un giudice-robot. La giustizia non è solo una questione di calcolo e precisione, ma di equità, umanità e capacità di comprendere il valore delle singole esistenze. Ridurre il processo giuridico a un algoritmo significherebbe svuotarlo della sua dimensione più profonda, trasformandolo in un esercizio di mera applicazione automatica, privo di sensibilità e di coscienza critica.

4. Tutela della salute

Nel settore sanitario, l'adozione dell'IA ha determinato progressi significativi, migliorando la diagnosi, la ricerca farmacologica e la gestione delle malattie¹³. In particolare, il suo impiego in radiologia ha reso possibile l'analisi avanzata di immagini diagnostiche, come raggi X, TAC e risonanze magnetiche, consentendo l'individuazione di patologie con una precisione straordinaria. Questa evoluzione non solo accelera il processo diagnostico, ma contribuisce a ridurre il margine di errore umano, garantendo una maggiore affidabilità nelle valutazioni mediche. Tuttavia, sebbene queste innovazioni offrano potenzialità straordinarie, è necessario interrogarsi su come bilanciare il loro impatto con la necessità di preservare la centralità del giudizio clinico umano, evitando di delegare in modo acritico decisioni cruciali a modelli predittivi che, per quanto sofisticati, rimangono strumenti privi di discernimento etico e autonomia critica.

L'intelligenza artificiale sta rivoluzionando anche la ricerca farmacologica, accelerando in modo significativo il processo di sviluppo di nuovi farmaci. Attraverso l'analisi avanzata di dati molecolari, è in grado di identificare composti farmacologicamente attivi e di

¹³ Su questi profili, a titolo esemplificativo: D. Schönberger, *Artificial intelligence in healthcare: A critical analysis of the legal and ethical implications*, in *International Journal of Law and Information Technology*, 27, 2, 2019.

individuare nuove strategie terapeutiche, comprese soluzioni per malattie che fino a poco tempo fa erano considerate incurabili.

Anche nell'epidemiologia, l'IA sta giocando un ruolo cruciale. Grazie alla capacità di elaborare enormi quantità di dati, è possibile individuare modelli e tendenze che facilitano l'identificazione dei fattori di rischio e la prevenzione delle malattie. L'intelligenza artificiale può migliorare la capacità di rilevare precocemente le epidemie e di valutare l'efficacia degli interventi di salute pubblica, fornendo strumenti preziosi per la gestione delle crisi sanitarie.

Un altro settore in cui l'IA sta mostrando il suo potenziale è la formazione medica. Le simulazioni chirurgiche avanzate permettono agli studenti di medicina e ai chirurghi in formazione di esercitarsi in ambienti virtuali che riproducono fedelmente le condizioni reali di una sala operatoria. Questo approccio migliora le competenze pratiche e la fiducia degli operatori sanitari, riducendo al contempo i rischi per i pazienti e contribuendo alla formazione di professionisti altamente qualificati.

L'oncologia è forse uno dei campi in cui l'intelligenza artificiale sta apportando i cambiamenti più significativi. Attraverso l'analisi di dati clinici e di immagini diagnostiche, l'IA consente di individuare segni precoci di tumori con una precisione senza precedenti. Inoltre, la capacità di elaborare profili genetici permette di suggerire terapie personalizzate, aumentando le possibilità di successo dei trattamenti e migliorando le percentuali di sopravvivenza.

Il principale interrogativo che emerge dall'adozione dell'intelligenza artificiale nell'assistenza sanitaria riguarda la capacità dell'attuale quadro normativo di garantire una protezione efficace contro le discriminazioni automatizzate¹⁴. In particolare, è essenziale chiedersi se il sistema di tutela anti-discriminatoria dell'Unione Europea sia adeguato a proteggere i pazienti dall'uso di algoritmi che potrebbero perpetuare o amplificare disuguaglianze preesistenti.

Per affrontare questa problematica è necessario analizzare la discriminazione nell'ambito sanitario sotto tre prospettive interconnesse¹⁵. Il profilo *sociale* evidenzia come le

¹⁴ M.A. Wójcik, *Algorithmic Discrimination in Health Care: An EU Law Perspective*, in *Health and Human Rights*, 24, 1, 2022, pp. 93-104.

¹⁵ S. Hoffman, A. Podgurski, *Artificial intelligence and discrimination in health care*, in *Yale Journal of Health Policy, Law, and Ethics*, 19, 3, 2020, pp. 1-8; I. G. Cohen, R. Amarasingham, A. Shah, et al., *The legal and ethical concerns*

nuove tecnologie possano riprodurre schemi di esclusione già presenti nei sistemi di assistenza, penalizzando in particolare gruppi vulnerabili; il profilo *giuridico*, pone invece la questione della trasparenza e della responsabilità nell'uso degli strumenti algoritmici, interrogandosi su come garantire il rispetto dei principi di equità e giustizia nel processo decisionale sanitario; infine, il profilo *tecnologico*, che impone una riflessione sulla progettazione degli algoritmi e sulla necessità di introdurre controlli rigorosi per evitare che i modelli predittivi amplifichino *bias* strutturali.

Per quanto riguarda il primo ambito, occorre ricordare che nel 2013 l'Agenzia per i Diritti Fondamentali ha pubblicato un rapporto sulle disuguaglianze nell'accesso e nella qualità dell'assistenza sanitaria in alcuni Stati membri dell'UE. Lo studio si è concentrato su tre gruppi vulnerabili: donne, anziani e giovani con disabilità intellettiva. Questi gruppi affrontano spesso "discriminazioni multiple", sia "additive" (discriminazione simultanea per vari motivi) sia "intersezionali" (sinergia unica di più motivi di discriminazione). Ad esempio, una persona gay con disabilità può essere discriminata per la sua disabilità e per il proprio orientamento sessuale, mentre le donne appartenenti a minoranze etniche hanno esperienze diverse rispetto agli uomini della stessa etnia o alle donne bianche. Il rapporto evidenzia che la discriminazione può essere "diretta", con negazione dell'accesso paritario all'assistenza sanitaria, o "indiretta", con trattamenti che non considerano le esigenze specifiche dei pazienti¹⁶. Molti intervistati non hanno denunciato la discriminazione a causa della mancanza di conoscenza delle procedure di ricorso, difficoltà nel dimostrare le accuse, sfiducia nel processo di reclamo e paura di ritorsioni. Il rapporto ha anche rilevato che molti professionisti sanitari hanno una comprensione insufficiente del concetto di discriminazione e, sebbene consapevoli delle barriere esistenti, esitano a etichettarle, appunto, come discriminazione. Solo pochi professionisti sono stati in grado di spiegare la discriminazione multipla e offrire soluzioni¹⁷.

Per quanto riguarda il profilo *giuridico*, abbiamo già ricordato che i temi dell'uguaglianza e della non discriminazione sono trattati sia nelle fonti primarie sia in quelle secondarie

that arise from using complex predictive analytics in health care, in *Health Affairs*, 33, 7, 2014, pp. 1139-1140; W.N. Price II, *Artificial intelligence in health care: Applications and legal implications*, in *SciTech Lawyer*, 14, 2017, p. 10.

¹⁶ M. Samorani, L. G. Blount, *Machine learning and medical appointment scheduling: creating and perpetuating inequalities in access to health care*, in *American Journal of Public Health*, 110, 4, 2020, p. 440.

¹⁷ A. Rajkomar, M. Hardt, M. D. Howell, et al., *Ensuring fairness in machine learning to advance health equity*, in *Annals of Internal Medicine*, 169, 12, 2018, pp. 866-872.

rie del diritto dell'UE. La Corte di Giustizia Europea ha confermato che la non discriminazione è un principio generale del diritto dell'UE nel caso Mangold¹⁸. La Carta Europea protegge i diritti di tutti ad accedere all'assistenza sanitaria preventiva e ai trattamenti medici (art. 35) e include una disposizione antidiscriminatoria aperta (art. 21). L'articolo 19 del TFUE consente al Consiglio europeo, con il consenso del Parlamento europeo, di adottare misure per combattere la discriminazione basata su sesso, razza, origine etnica, religione, convinzioni personali, disabilità, età o orientamento sessuale.

In relazione all'accesso all'assistenza sanitaria, si applicano due direttive UE principali: la Direttiva 2000/43/CE sull'uguaglianza razziale e la Direttiva 2004/113/CE sui beni e servizi. Queste direttive proibiscono la discriminazione basata su razza, origine etnica e sesso nel contesto dell'assistenza sanitaria, applicandosi sia alla discriminazione diretta che indiretta nel settore pubblico e privato. Nessuna delle due direttive protegge esplicitamente dalla discriminazione multipla, anche se la Direttiva sull'uguaglianza razziale vi fa riferimento nel preambolo. Entrambe le direttive invertono l'onere della prova, richiedendo al convenuto di dimostrare che non è stata commessa discriminazione se il ricorrente presenta una prova *prima facie*.

Giacomo Di Federico¹⁹ ha identificato tre problemi principali con la legge antidiscriminazione dell'UE nell'assistenza sanitaria: 1) le direttive non proibiscono la discriminazione basata su religione, sesso, disabilità, età o orientamento sessuale; 2) la possibilità di presentare richieste di discriminazione multipla è limitata; 3) l'attuazione delle direttive varia tra gli Stati membri, creando disparità nella protezione contro la discriminazione. Inoltre, gli organismi di parità affrontano difficoltà dovute al basso numero di reclami, problemi di raccolta delle prove e mancanza di risorse e competenze. Le disposizioni antidiscriminatorie delle direttive e dell'articolo 21 della Carta Europea si applicano solo quando la situazione rientra nel campo di applicazione del diritto dell'UE, limitando le protezioni offerte ai pazienti. L'UE ha competenza condivisa nella regolamentazione della circolazione di beni e servizi medici e delle preoccupazioni di sicurezza in materia di salute pubblica, ma l'organizzazione dei sistemi sanitari nazionali resta di competenza degli Stati membri, portando a livelli diseguali di protezione contro la discriminazione sanitaria in tutta l'UE.

¹⁸ Causa C-144/04 EU:C:2005:709, Werner Mangold contro Rüdiger Helm.

¹⁹ G. Di Federico, *Access to healthcare in the European Union: Are EU patients (effectively) protected against discriminatory practices?*, in L.S. Rossi, F. Casolari, *The principle of equality in EU law*, Springer Publishing, New York, 2017, pp. 229-232.

Per quanto riguarda, infine, il profilo *tecnologico*, dato che la normativa antidiscriminatoria dell'UE non affronta adeguatamente la discriminazione nei confronti dei pazienti, l'introduzione dell'IA aggiunge nuove preoccupazioni. L'IA può sia aiutare a risolvere pratiche discriminatorie esistenti sia creare nuovi modelli di discriminazione difficili da individuare e affrontare.

Una delle preoccupazioni più rilevanti legate agli algoritmi di apprendimento automatico è la loro imprevedibilità e la difficoltà di spiegare il processo decisionale che li guida. Il caso di IBM Watson è emblematico: impiegato in ambito medico per supportare la diagnosi e la scelta dei trattamenti, questo sistema utilizza tecniche avanzate di apprendimento automatico, ma la sua dipendenza dai dati lo rende spesso opaco e poco comprensibile, anche per gli stessi specialisti. L'incapacità di tracciare con precisione il percorso logico che porta a una raccomandazione clinica solleva interrogativi sulla sua affidabilità e sulla possibilità di correggere eventuali errori o distorsioni²⁰.

Strumenti di supporto decisionale come IBM Watson possono fornire un valido aiuto agli operatori sanitari, integrando informazioni provenienti da un'enorme quantità di dati e migliorando la capacità di personalizzare le cure. L'IA, infatti, è in grado di promuovere un approccio alla medicina più mirato, adattando i trattamenti alle esigenze specifiche dei singoli pazienti. Tuttavia, l'efficacia di queste tecnologie dipende dalla loro trasparenza e dalla possibilità, per i professionisti della salute, di comprendere e verificare le decisioni suggerite, evitando di affidarsi acriticamente a sistemi il cui funzionamento rimane, in larga parte, una *black box*.

Altro elemento fondamentale da considerare è il ricorso agli algoritmi che categorizzano i pazienti in base alla "razza". Il concetto stesso di "differenza razziale" è complesso e spesso frainteso. Se da un lato alcune condizioni genetiche possono essere più frequenti in specifiche popolazioni – come la beta-talassemia nelle popolazioni mediterranee o l'anemia falciforme tra le persone di origine africana – la maggior parte delle differenze in termini di salute derivano da fattori socio-economici, piuttosto che biologici. L'accesso alle cure, le condizioni di vita, la qualità dell'alimentazione e l'esposizione a fattori ambientali giocano un ruolo determinante nello stato di salute di

²⁰ F. Lagioia, G. Contissa, *The strange case of Dr. Watson: Liability implications of evidence-based decision support systems in the health care*, in A. Santosuosso, C. A. Redi (eds.), *Law, Sciences and New Technologies*, Pavia University Press, Pavia, 2020.

un individuo, e ridurre la questione a una categorizzazione razziale rischia di rafforzare stereotipi anziché affrontare le reali cause delle disuguaglianze sanitarie. L'uso di algoritmi che incorporano variabili razziali deve quindi essere valutato con estrema cautela. Il rischio è quello di naturalizzare differenze che sono in larga parte costruite socialmente, con il risultato di perpetuare trattamenti differenziati ingiustificati, anziché migliorare l'accesso equo alle cure.

Le sfide poste dall'IA sollevano interrogativi sull'efficacia dell'attuale quadro normativo antidiscriminatorio dell'UE, che già fatica a garantire una protezione adeguata contro la discriminazione multipla e strutturale. Gli algoritmi non discriminano attraverso categorie tradizionali, ma introducono nuovi schemi di esclusione, rendendo obsolete le protezioni offerte dalla legislazione vigente. Questo aspetto impone una riflessione sulla necessità di aggiornare il diritto antidiscriminatorio europeo, che, nelle sue versioni attuali (ad esempio, la Direttiva 2000/43/CE sulla parità di trattamento indipendentemente dalla razza o dall'origine etnica e la Direttiva 2000/78/CE sull'occupazione e le condizioni di lavoro), non prevede strumenti specifici per affrontare le distorsioni prodotte dagli algoritmi.

In conclusione, la discriminazione algoritmica nell'assistenza sanitaria non è solo una questione tecnologica, ma un fenomeno che perpetua le disuguaglianze sociali e viola diritti fondamentali, mettendo in pericolo la salute e la vita delle persone più vulnerabili. La normativa antidiscriminatoria dell'UE, così com'è concepita oggi, appare inadeguata a rispondere alle sfide poste dall'IA, rendendo necessario un ripensamento complessivo della sua capacità di protezione in un'era di automazione crescente.

5. Settore finanziario, bancario e assicurativo

L'inclusione del settore finanziario, bancario e assicurativo tra quelli ad alto rischio nella Risoluzione 2020/2012 del Parlamento europeo segna un passaggio cruciale nella regolamentazione dell'uso dell'intelligenza artificiale nelle aree critiche dell'economia. Questa scelta riflette una crescente consapevolezza delle implicazioni etiche e giuridiche dell'IA in ambito finanziario e la necessità di sviluppare meccanismi di controllo che ne garantiscano un utilizzo equo e trasparente.

L'introduzione dell'IA nei settori finanziari ha aperto nuove opportunità, ma ha anche sollevato questioni delicate, soprattutto in relazione alla valutazione del rischio creditizio. In passato, le decisioni di concessione del credito si basavano su parametri tradizionali come il reddito e la storia creditizia. Con l'impiego dell'IA, l'analisi si è fatta molto più sofisticata: gli algoritmi possono elaborare una quantità senza precedenti di informazioni, includendo non solo dati finanziari, ma anche elementi legati ai comportamenti economici individuali, ai pagamenti passati e persino a variabili non finanziarie, come l'attività online e le interazioni sociali.

Questo approccio permette valutazioni del rischio più dettagliate e, in teoria, potrebbe ridurre il pericolo di concedere credito a soggetti non solvibili. Tuttavia, la sua applicazione pone problemi di trasparenza e discriminazione, poiché i criteri utilizzati dagli algoritmi non sono sempre espliciti e verificabili. Inoltre, l'inclusione di dati personali non strettamente finanziari solleva questioni sulla privacy e sul possibile rafforzamento di *bias* preesistenti, penalizzando gruppi già svantaggiati. Se non regolamentata con attenzione, l'automazione del rischio creditizio potrebbe quindi accentuare le disuguaglianze nell'accesso ai servizi finanziari, trasformando il credito da strumento di inclusione a ulteriore barriera sociale.

Un esempio concreto di questa trasformazione è rappresentato dal credit scoring algoritmico, un sistema basato sull'IA per calcolare il punteggio di credito di un individuo. Questo metodo consente una valutazione più approfondita del rischio e può individuare schemi che sfuggirebbero a un'analisi manuale. Tuttavia, l'uso di tali strumenti solleva questioni cruciali in termini di trasparenza e responsabilità, poiché i consumatori spesso non hanno accesso o comprensione dei criteri con cui vengono prese le decisioni che li riguardano.

Oltre alla valutazione del merito creditizio, l'IA è sempre più impiegata nella rilevazione delle frodi finanziarie, nella gestione degli investimenti, nella previsione delle dinamiche di mercato e nella personalizzazione dei servizi finanziari.

L'Europa ha riconosciuto la necessità di affrontare queste sfide etiche e giuridiche attraverso interventi normativi specifici, cercando di bilanciare l'efficienza dell'IA con la tutela dei diritti fondamentali e delle garanzie di sicurezza. Il quadro regolatorio si pone l'obiettivo di rendere le istituzioni finanziarie responsabili nell'uso dell'IA, introducendo

obblighi di trasparenza e meccanismi di tutela per i consumatori. Tuttavia, la velocità con cui queste tecnologie si evolvono richiede un costante aggiornamento normativo per evitare che il potere decisionale delegato agli algoritmi si traduca anche in questo campo in nuove forme di discriminazione e opacità.

L'inclusione del settore finanziario tra quelli ad alto rischio evidenzia la complessità delle sfide etiche e giuridiche poste dall'intelligenza artificiale in questo contesto. È essenziale trovare un equilibrio tra l'efficienza e la precisione offerte dall'IA e la tutela dei diritti fondamentali, assicurandosi che l'uso di queste tecnologie rispetti standard normativi e principi etici condivisi. La definizione di un quadro regolatorio adeguato è imprescindibile per affrontare queste sfide in modo efficace e sostenibile.

Uno degli aspetti più complessi è, ancora una volta, la spiegabilità delle decisioni algoritmiche: gli intermediari finanziari devono essere in grado di comunicare in modo chiaro e comprensibile come vengono prese le decisioni, soprattutto quando incidono sull'accesso a servizi essenziali come il credito. La mancanza di trasparenza in questo ambito compromette la fiducia dei clienti e può tradursi in una violazione del principio di equità.

Un'ulteriore area di preoccupazione riguarda la protezione dei dati personali. La necessità di trattare i dati in modo confidenziale e sicuro si riferisce alla garanzia che tali informazioni siano accessibili solo a soggetti autorizzati e utilizzate esclusivamente per scopi legittimi. Questo principio, che ha una dimensione sia tecnica sia giuridica, è alla base della regolamentazione sulla privacy e protezione dei dati, come il GDPR nell'Unione Europea, e richiede standard rigorosi per prevenire abusi e violazioni della riservatezza.

Infine, vorrei ricordare che l'uso dell'intelligenza artificiale nel *rating* bancario rappresenta una trasformazione particolarmente significativa nel settore finanziario, rendendo le valutazioni creditizie più rapide ed efficienti. Tuttavia, l'integrazione degli algoritmi in questi processi solleva interrogativi cruciali in termini di privacy, protezione dei dati e discriminazione algoritmica, richiedendo un'attenta regolamentazione per bilanciare innovazione e diritti fondamentali.

Com'è noto, l'articolo 22 del GDPR affronta direttamente il tema della profilazione automatizzata, stabilendo che nessuna decisione che abbia un impatto giuridico significativo su una persona possa essere presa esclusivamente da un algoritmo, senza l'intervento umano. Questo aspetto è particolarmente rilevante per il *rating* bancario, in cui molte

decisioni creditizie vengono automatizzate sulla base di dati finanziari personali. Tuttavia, se da un lato l'uso dell'IA consente una valutazione più accurata della capacità di rimborso, dall'altro esiste il rischio di opacità nelle decisioni, discriminazioni indirette e perdita di controllo sulle modalità con cui gli algoritmi determinano l'accesso al credito.

Il considerando 71 del GDPR evidenzia la necessità di regole trasparenti nei processi decisionali automatizzati, affinché i soggetti coinvolti possano comprendere i criteri adottati per le valutazioni. Tuttavia, il bilanciamento tra la necessità di garantire una valutazione creditizia efficace e il rispetto del diritto alla privacy rimane una sfida aperta. Il trattamento dei dati personali nel *rating* bancario avviene su larga scala e spesso coinvolge informazioni sensibili, rendendo essenziale garantire la sicurezza e la conformità ai principi normativi. Se un individuo ha avuto difficoltà finanziarie in passato, l'IA potrebbe suggerire condizioni di prestito più sfavorevoli, creando un effetto di auto-rafforzamento della vulnerabilità finanziaria. Inoltre, poiché gli algoritmi apprendono dai dati storici, potrebbero replicare e perpetuare pregiudizi sistemici già esistenti, penalizzando gruppi socialmente svantaggiati. Questo fenomeno solleva la necessità di una regolamentazione più stringente sull'uso dell'IA nel settore finanziario, imponendo trasparenza, audit delle decisioni algoritmiche e misure di mitigazione del rischio di *bias*.

Per garantire un equilibrio tra innovazione e diritti fondamentali, è essenziale sviluppare norme specifiche che regolino l'uso degli algoritmi nel credito, imponendo obblighi di spiegabilità e *accountability* agli intermediari finanziari. Solo attraverso un adeguato quadro normativo sarà possibile evitare che il progresso tecnologico si traduca in nuove forme di esclusione economica, garantendo un sistema finanziario equo e accessibile a tutti.

6. Lavoro e selezione del personale

Negli ultimi decenni, i progressi tecnologici hanno rivoluzionato molti settori lavorativi, portando con sé benefici ma anche problemi e gravi preoccupazioni²¹.

Le tecnologie di IA possono eseguire compiti sino ad oggi affidati ad esseri umani con una velocità e una precisione che superano quelle umane. Ad esempio, sistemi come Kira

²¹ Sul tema un'ottima introduzione è il volume A. Aloisi, V. De Stefano, *Il tuo capo è un algoritmo. Contro il lavoro disumano*, Laterza, Roma-Bari, 2020.

Systems e Blue J Legal nel settore legale automatizzano la previsione dei risultati legali, permettendo agli avvocati di concentrarsi su decisioni più complesse. Kira Systems analizza i contratti e identifica clausole rilevanti in tempi significativamente ridotti rispetto a quelli necessari agli avvocati umani; Blue J Legal, invece, prevede gli esiti delle dispute fiscali basandosi su precedenti giudiziari, migliorando la precisione delle previsioni e riducendo il tempo dedicato alla ricerca giuridica.

L'IA può anche aumentare la produttività del lavoro umano in compiti decisionali. Nei casi in cui essa fornisce previsioni, gli esseri umani possono prendere decisioni più informate e precise. Nel settore medico, per esempio, gli strumenti di IA come quelli sviluppati da Zebra Medical Vision aiutano i radiologi a individuare anomalie nelle immagini mediche con maggiore precisione, migliorando le diagnosi senza sostituire completamente i radiologi. Questo non solo aumenta l'efficienza, ma può anche aumentare la domanda di lavoro in altri ambiti medici. L'IA può così creare anche nuovi compiti riducendo l'incertezza e rendendo economicamente fattibili attività precedentemente impossibili. Un esempio è dato dalla gestione preventiva dei veicoli, dove le tecnologie di previsione consentono di programmare la manutenzione prima che si verifichino guasti, migliorando l'efficienza operativa e riducendo i costi. In alcuni settori, come quello della guida commerciale, l'IA può aumentare l'efficienza ma anche erodere il vantaggio competitivo dei lavoratori umani (ad esempio, le app di navigazione basate su IA come Waze hanno ridotto l'importanza dell'esperienza dei tassisti di Londra, rendendo accessibili a chiunque le conoscenze sulle migliori rotte della città).

L'automazione e l'intelligenza artificiale stanno dunque trasformando profondamente il mondo del lavoro, con un impatto che varia a seconda delle professioni e dei settori. Se da un lato alcune mansioni potrebbero vedere una riduzione della domanda, dall'altro si registra un incremento della produttività e la creazione di nuovi compiti. Tuttavia, al centro del dibattito vi è il fatto che i sistemi automatizzati raccolgono e quantificano una vasta gamma di dati sui lavoratori, comprese le loro abitudini, le caratteristiche personali e persino informazioni biometriche sensibili, come livelli di stress e salute. Questi dati alimentano i sistemi di apprendimento automatico, che vengono utilizzati dai datori di lavoro per prendere decisioni sulle assunzioni, valutare le prestazioni, prevedere problemi sul posto di lavoro e persino determinare la retribuzione²².

²² P. Kim, *Data-Driven Discrimination at Work*, in *William & Mary Law Review*, 48, 2017, pp. 857-936.

Molti sono gli studi che hanno sollevato preoccupazioni riguardo alle crescenti limitazioni della privacy e dell'autonomia dei lavoratori, al potenziale che la discriminazione a livello sociale possa infiltrarsi nei sistemi di apprendimento automatico, e alla generale mancanza di trasparenza riguardo alle regole del posto di lavoro²³.

Particolare rilevanza ha assunto il ruolo dell'IA nella selezione del personale²⁴. Le organizzazioni multinazionali stanno utilizzando massicciamente soluzioni di IA in questo ambito, sollevando una lunga serie di questioni in termini di equità, trasparenza ed etica. I software di IA sono in grado di scansionare enormi volumi di applicazioni in brevissimo tempo, eliminando i candidati che non soddisfano i requisiti minimi e classificando i migliori tra loro, facendo risparmiare tempo e denaro alle aziende. Avendo a disposizione modelli predittivi, l'IA può valutare le capacità e le esperienze dei candidati in modo molto preciso, assegnando ai candidati punteggi riferibili alla possibile performance, feedback psicologici e indici di affidabilità²⁵.

Occorre però guardare l'altra faccia della medaglia: la scelta attuata sulla base dei filtri di IA non riflette necessariamente un'opinione bilanciata e obiettiva. Se i dati storici mostrano una preferenza per un determinato gruppo demografico a discapito di altri, un algoritmo di selezione del personale potrebbe inconsapevolmente imparare a replicare tale discriminazione, escludendo candidati con caratteristiche diverse sulla base di correlazioni statistiche. Questo non solo cristallizza pregiudizi sistemici già esistenti, ma rende più difficile individuarli e correggerli, poiché vengono presentati come il risultato di un processo tecnico e imparziale²⁶.

L'automazione delle attività associate alla selezione comporta una totale disumanizzazione, cancellando il coinvolgimento umano nel processo, e dunque mortificando le qualità relazionali e di adattamento culturale dei candidati. Inoltre, i candidati non vengono informati in modo trasparente su come i loro dati vengono utilizzati e in che modo ven-

²³ S. Borelli, M. Ranieri, *La discriminazione nel lavoro autonomo. Riflessioni a partire dall'algoritmo Frank*, in *Labour & Law Issues*, 7, 1, 2021, pp. I.18-I.47.

²⁴ A. Kelly-Lyth, *Challenging Biased Hiring Algorithms*, in *Oxford Journal of Legal Studies*, 41, 4, 2021, pp. 899-928.

²⁵ A. Köchling-M. C. Wehner, *Discriminated by an algorithm: a systematic review of discrimination and fairness by algorithmic decision-making in the context of HR recruitment and HR development*, in *Business Research*, 13, 3, 2020, pp. 795-848.

²⁶ M. Peruzzi, *Intelligenza artificiale e lavoro. Uno studio su poteri datoriali e tecniche di tutela*, Giappichelli, Torino, 2023.

gono prese le decisioni. Gli esiti ottenuti dagli algoritmi di IA risultano così “de-responsabilizzanti”. L’organizzazione, del resto, non sempre è in grado di spiegare e chiarire i risultati ottenuti tramite i propri sistemi di IA.

L’uso dell’IA nella selezione del personale, dunque, offre un numero di opportunità per migliorare l’efficienza nel processo di assunzione. Tuttavia, comporta diversi tipi di sfide insieme a gravi questioni etiche. Le aziende dovrebbero agire in modo equilibrato, ponendo l’equità, la trasparenza e il rispetto della privacy al centro della loro attività. Ma, ovviamente, non ci si può illudere che imbocchino questa strada senza la spinta (gentile o meno) di sanzioni, positive (come incentivi o sgravi fiscali) o negative che siano.

Secondo Veena Dubal²⁷, le tecnologie di estrazione dei dati e i sistemi decisionali algoritmici stanno trasformando l’esperienza lavorativa in particolare per le forze lavoro a basso reddito e delle minoranze razziali. Queste tecnologie non solo monitorano costantemente i lavoratori, ma determinano anche la loro retribuzione in modo imprevedibile e spesso ingiusto. Sotto quest’ultimo profilo va il fenomeno noto come «discriminazione salariale algoritmica»²⁸.

Questo termine descrive una pratica in cui i lavoratori vengono pagati con salari orari diversi, calcolati utilizzando formule variabili basate su dati dettagliati riguardanti il comportamento individuale, la domanda e l’offerta, e altri fattori.

A differenza delle forme più tradizionali di pagamento variabile, la discriminazione salariale algoritmica – praticata attraverso i “bonus” e le schede di valutazione di Amazon o i sistemi di allocazione del lavoro di Uber – deriva dalla pratica della “discriminazione dei prezzi”, in cui ai singoli consumatori viene addebitato tanto quanto una compagnia determina che potrebbero essere disposti a pagare. Come pratica di gestione del lavoro, la discriminazione salariale algoritmica consente alle aziende di personalizzare e differenziare i salari per i lavoratori in modi a loro sconosciuti, pagando loro per comportarsi in modi che l’azienda desidera, dopo aver individuato la soglia minima che presumibilmente i lavoratori sono disposti ad accettare²⁹.

Data l’asimmetria informativa tra lavoratori e aziende, le aziende possono calcolare le esatte tariffe salariali necessarie per incentivare i comportamenti desiderati, mentre i lavoratori possono solo cercare di indovinare come le aziende determinano i loro salari.

²⁷ V. Dubal, *On algorithmic Wage Discrimination*, in *Columbia Law Review*, 123, fasc. 7, 2023, pp. 1929-1992.

²⁸ *Ibidem*.

²⁹ *Ibidem*.

La discriminazione salariale algoritmica crea un mercato del lavoro in cui le persone che svolgono lo stesso lavoro, con le stesse competenze, per la stessa azienda, nello stesso momento, possono ricevere salari orari diversi. I salari digitalmente personalizzati sono spesso determinati attraverso sistemi oscuri e complessi che rendono quasi impossibile per i lavoratori prevedere o comprendere la loro retribuzione costantemente mutevole e frequentemente decrescente³⁰.

La discriminazione salariale algoritmica solleva gravi preoccupazioni giuridiche ed etiche. Storicamente, le leggi sui salari sono state progettate per garantire equità e prevedibilità nella retribuzione dei lavoratori. Tuttavia, gli algoritmi che personalizzano i salari per specifici lavoratori e momenti temporali violano sia lo spirito che la lettera di queste leggi. Tali pratiche contravvengono ai principi di parità di retribuzione per lavoro uguale, creando disparità salariali significative tra lavoratori che svolgono compiti simili.

I lavoratori che subiscono la discriminazione salariale algoritmica riportano esperienze di ingiustizia “calcolatoria” e manipolazione. I sistemi di pagamento algoritmici, spesso opachi e complessi, rendono difficile per i lavoratori prevedere o comprendere le loro retribuzioni, che possono variare drasticamente in base a fattori fuori dal loro controllo, creando un ambiente di lavoro instabile e stressante, paragonabile al gioco d’azzardo, dove i lavoratori sono costantemente in balia delle fluttuazioni algoritmiche.

La trasparenza dei dati e la leggibilità degli algoritmi, pur essendo passi importanti, non sono sufficienti a risolvere i problemi fondamentali. È necessaria un’azione legislativa chiara che vieti queste pratiche, proteggendo i lavoratori dalla manipolazione economica e garantendo loro retribuzioni eque e prevedibili. Sebbene attualmente prevalente nei settori del lavoro *on-demand* e della *gig economy*, la discriminazione salariale algoritmica ha il potenziale di diffondersi in altri ambiti lavorativi, tra cui la sanità, l’ingegneria, il commercio al dettaglio e la ristorazione.

Se non affrontato, questo fenomeno potrebbe normalizzare nuove norme culturali di compensazione per il lavoro a basso salario, esacerbando le disuguaglianze economiche esistenti e minando ulteriormente la stabilità economica delle fasce più vulnerabili della società.

³⁰ Cfr. A. Shapiro, *Dynamic exploits: calculative asymmetries in the on-demand economy*, in *New Technology, Work and Employment*, 35, 2, 2020, pp. 162-177.

7. Conclusioni

La discriminazione di precisione è l'ultima evoluzione dell'ingiustizia sociale, un fenomeno che sfrutta il potenziale delle tecnologie per approfondire le disuguaglianze invece di combatterle. Il mito dell'intelligenza artificiale come strumento neutrale e oggettivo è la grande menzogna del capitalismo digitale, che impone strumenti di esclusione sociale con l'apparenza dell'efficienza. Dietro ogni algoritmo c'è un'ideologia, un modello di società che sceglie chi merita opportunità e chi deve essere lasciato indietro. E il sistema che sta emergendo è chiarissimo: un mondo in cui il potere economico e tecnologico concentra ancora di più le risorse nelle mani di pochi, mentre le classi popolari subiscono l'automatizzazione delle disuguaglianze.

Nel sistema giudiziario, l'uso degli algoritmi per determinare la pericolosità sociale o il rischio di recidiva non è altro che una nuova forma di repressione classista e razzista, travestita da innovazione. Non è una coincidenza se gli strumenti predittivi della giustizia colpiscono in maniera sproporzionata le fasce più deboli, i lavoratori precari, le comunità marginalizzate. La pretesa di un'oggettività algoritmica è un inganno: quello che accade, in realtà, è la riproduzione su scala automatica degli stessi meccanismi discriminatori che il potere ha sempre usato per opprimere. Il caso Loomis dimostra che l'IA non elimina il pregiudizio, ma lo trasforma in sentenza, in condanna automatizzata, senza appello e senza possibilità di difesa.

Nel settore sanitario, il quadro è altrettanto chiaro: l'intelligenza artificiale non sta democratizzando l'accesso alla salute, ma sta rafforzando il privilegio di classe. Le tecnologie avanzate sono un lusso per chi può permetterselo, mentre per le classi lavoratrici e per chi vive ai margini la sanità diventa sempre più inaccessibile, sempre più stratificata, sempre più escludente. Gli algoritmi di diagnosi e trattamento, se non regolati, escludono chi non rientra nei modelli statistici dominanti, negano cure a chi non ha uno storico sanitario digitalizzato, replicano nei dati medici le disuguaglianze sociali e razziali. Il rischio è che l'automazione diventi uno strumento di selezione sociale, un nuovo dispositivo di controllo che garantisce vita e salute solo a chi rientra nei parametri della produttività e della conformità economica.

Il settore finanziario e del credito mostra con ancora più brutalità l'essenza di questa discriminazione. Gli algoritmi bancari e assicurativi non sono strumenti di valutazione, ma dispositivi di segregazione economica, che decidono chi può accedere alle risorse e chi deve restare schiacciato sotto il peso dell'esclusione finanziaria. I modelli predittivi replicano la storia delle disuguaglianze di classe: se sei nato povero, resterai povero. Se provieni da un quartiere svantaggiato, sarai un soggetto a rischio. Se il tuo lavoro non è considerato "sicuro", ti verranno negati prestiti e opportunità. E così la miseria diventa ereditaria, cementata nelle logiche di un sistema che conferisce alla finanza il potere di decidere il futuro delle persone.

Nel mondo del lavoro, l'IA viene venduta come un'opportunità per migliorare i processi, ma la verità è che sta distruggendo la dignità del lavoro e rafforzando il dominio del capitale sulla forza-lavoro. Gli algoritmi di selezione e monitoraggio aziendale non solo eliminano la contrattazione collettiva, ma impongono un modello di sorveglianza permanente e di disumanizzazione del lavoratore. Le macchine non negoziano, non tengono conto della complessità della vita, non concedono tregue. Chi non si adatta al ritmo imposto viene scartato, senza spiegazioni, senza possibilità di replica. La precarizzazione diventa norma, il ricatto occupazionale si intensifica.

La regolamentazione dell'intelligenza artificiale non può limitarsi a un *lifting* normativo, non può essere un semplice ritocco "estetico" al quadro giuridico esistente: deve essere un gesto capace di arginare la primazia del capitale digitale sull'esser umano. Servono controllo pubblico e collettivo sulle tecnologie, accesso trasparente ai modelli decisionali, diritto di contestazione e revisione delle decisioni algoritmiche: serve il rispetto dei principi giuridici ed etici che hanno animato quella che Norberto Bobbio chiamava l'età dei diritti.

Solo così l'intelligenza artificiale potrà essere non uno strumento di oppressione, ma un mezzo per emancipare le classi più deboli, redistribuire ricchezza, ridurre le disuguaglianze invece di amplificarle. Non si tratta soltanto di "migliorare" gli algoritmi, ma di ricostituirne il senso di una tecnologia al servizio dell'umanità.