



LA TRASFORMAZIONE DIGITALE DELLO SPAZIO EUROPEO

Cod.: P24002

Napoli, Castel Capuano, 16 gennaio 2024

Relatore: Avv. Roberto Arcella

Rischi, prevenzione e responsabilità derivanti dall'utilizzo dell'intelligenza artificiale generativa: nesso di causalità ed *explainability* nei *large language model* gestione del rischio nella consultazione delle banche dati giurisprudenziali

L'intelligenza artificiale non è uno strumento che il giurista attento scopre solo negli ultimi tempi. Di nuovo oggi c'è che è stata resa accessibile alla generalità dei consociati un'intelligenza artificiale di tipo generativo ma, in realtà, è da anni che si fa uso di un diverso tipo di intelligenza artificiale, che è quella basata su regole deterministiche. L'IA deterministica si basa su algoritmi che seguono regole e logiche predefinite, produce risultati consistenti e prevedibili per lo stesso *input*, in quanto non è dotata di capacità di apprendimento o adattamento. Parliamo, in altri termini, dell'utilizzo dei computer e dei software che seguono una logica condizionale semplice per eseguire azioni specificate, secondo il meccanismo "*If This Then That*", utilizzati principalmente per automatizzare una varietà di processi e attività, riducendo la necessità di intervento umano. Questa logica si basa su eventi, chiamati "trigger" (grilletto). Il "trigger" (il *se questo*) è l'evento che innesca

l'azione. Quando il *trigger* viene attivato, il software esegue l'azione specificata (*l'allora quello*).

Non a caso, l'ordinamento giuridico ha reagito, non da ieri, con norme a tutela dei consociati finalizzate a preservarli dal rischio ed a garantire la risarcibilità dei danni derivanti dall'utilizzo di siffatti strumenti. Senza voler entrare nel dettaglio e passare in rassegna tali provvedimenti, basti pensare alla prima legge a tutela della riservatezza dei dati personali, che risale al 1996 (n. 675 del 31/12/1996), ed alle relative evoluzioni avute con il D.lgs. 196/2003 fino ad arrivare al GDPR ed ai rimaneggiamenti della normativa nazionale che ne sono seguiti. O, ancora, alle norme in materia di responsabilità del produttore, vale a dire al d.p.r. 224/1998, che recepì la Direttiva CEE 374/1985, norme poi trasfuse negli articoli 114 e seguenti del codice del consumo, con la previsione della c.d. responsabilità presunta la quale, come noto, a differenza di quella oggettiva, prescinde dall'accertamento della colpa del produttore ma non anche dalla dimostrazione dell'esistenza di un difetto del prodotto, del danno e, soprattutto, del collegamento causale tra l'uno e l'altro. Risarcimento danni da perdita di dati, da violazione del diritto alla riservatezza, da malfunzionamento di software rappresentano le ipotesi più frequenti nella casistica.

Ovviamente, l'intelligenza artificiale generativa ha aperto nuovi scenari, sui quali l'Unione Europea ha già proposto o adottato alcuni regolamenti, alcuni che centrano altri che lambiscono l'argomento e che fanno parte di un vero e proprio "pacchetto", sui quali non mi soffermerò se non per un brevissimo cenno:

- il **Digital Service Act (DSA)**: Il Digital Service Act (Regolamento UE 2022/2065) mira a modernizzare il quadro normativo per i servizi digitali. Esso si concentra sulla regolamentazione delle piattaforme online, stabilendo regole chiare per garantire la sicurezza e la trasparenza delle attività online. Prevede una serie di obblighi per i fornitori di servizi digitali, in particolare per le grandi piattaforme, che includono la gestione delle informazioni illegali, la protezione dei consumatori e la promozione di un ambiente online sicuro;
- il **Digital Market Act (DMA)**: Il Regolamento, che fa parte del "pacchetto" DSA, è stato adottato dal Parlamento Europeo e dal Consiglio il 14

settembre 2022 ed è stato pubblicato nella Gazzetta Ufficiale il 12 ottobre 2022. Si concentra sulla promozione di una maggiore concorrenza nel settore dei servizi digitali e stabilisce regole per prevenire comportamenti anticoncorrenziali da parte dei cosiddetti "*gatekeeper*", ovvero grandi piattaforme che svolgono un ruolo cruciale nei mercati digitali.

- il **Digital Governance Act** (DGA – Regolamento UE 2022/868 del 30 maggio 2022) è parte della strategia dell'UE per il digitale e si concentra sulla promozione dell'utilizzo dei dati nel settore pubblico. Entrato in vigore il 24 settembre 2023, Mira a facilitare il riutilizzo dei dati del settore pubblico e a migliorare il modo in cui i dati vengono condivisi e gestiti all'interno dell'UE, includendo anche disposizioni sulla l'interoperabilità dei dati tra enti pubblici e tra il settore pubblico e quello privato e sul relativo riutilizzo.

- il **Data Act** è, infine, un regolamento recentissimo (13 dicembre 2023, n. 2854/2023), che si propone di regolamentare il flusso di dati all'interno dell'Unione Europea e tra questa e altri paesi. Si concentra sulla condivisione e l'utilizzo dei dati, in particolare quelli generati da prodotti e servizi connessi. Esso mira a garantire che i dati possano essere accessibili e utilizzati in modo equo, promuovendo l'innovazione e la creazione di nuovi servizi.

Con specifico riferimento a responsabilità e danni, assai significativa è anche la **Risoluzione del Parlamento europeo del 16 febbraio 2017** recante raccomandazioni alla Commissione concernenti norme di diritto civile sulla robotica, che riconosce quindi le enormi potenzialità sia le sfide significative presentate dallo sviluppo della robotica e dell'intelligenza artificiale, sollecitando un approccio regolamentare che bilanci innovazione, etica, e sicurezza. In essa viene suggerito un sistema di assicurazione obbligatoria, accompagnato da un fondo di risarcimento, e viene altresì proposta l'individuazione di una sorta di *status* giuridico dei robot avanzati. Si propone, infine, l'adozione di un codice etico per ingegneri e comitati di ricerca in robotica, suggerendo un impegno condiviso a livello internazionale per sviluppare le relative soluzioni.

Del pari, va segnalata la "**Risoluzione del Parlamento europeo del 20 ottobre 2020 recante raccomandazioni alla Commissione su un regime di responsabilità civile per l'intelligenza artificiale (2020/2014)**". In essa si enfatizza l'importanza di un quadro giuridico che infonda fiducia nella sicurezza, affidabilità e coerenza di prodotti e servizi tecnologici, inclusi quelli basati sull'IA e, in tale cornice, si sottolinea la necessità di un approccio equilibrato che protegga le potenziali vittime di danni e, allo stesso tempo, non costituisca per i produttori un freno allo sviluppo di nuove tecnologie. In tale contesto, di centrale importanza è la definizione (e la figura) dell' *operatore* in ambito IA, che include sia l'operatore di front-end (l'utente finale del sistema) sia l'operatore di back-end (il fornitore, vale a dire chi controlla o mantiene il sistema); al riguardo, si afferma che la sussistenza della responsabilità sia dell'operatore di front-end che di quella dell'operatore del back-end e si suggerisce che la stessa dovrebbe essere proporzionale al grado di controllo e benefici derivanti dall'uso del sistema di IA. Coerentemente con la risoluzione del 16 febbraio 2017 sulla robotica, anche tale provvedimento sottolinea l'importanza del riconoscimento di risarcimenti effettivamente soddisfattivi, anche per gli eventuali danni non patrimoniali e si raccomanda anche l'istituzione di forme di assicurazione obbligatoria.

Da ultimo, l'**AI Act**. Si tratta della proposta di Regolamento del Parlamento europeo e del Consiglio recante norme armonizzate sull'intelligenza artificiale, che rappresenta uno dei primi tentativi a livello globale di regolamentare questa tecnologia emergente. La normativa si caratterizza per essere *risk-based*, vale a dire che essa disciplina la materia con una classificazione dei sistemi di IA in base al livello di rischio associato al loro uso. Viene stabilita una distinzione tra sistemi a rischio "inaccettabile", "elevato", "limitato" o "minimo":

- **Rischio Inaccettabile:** include sistemi di IA che violano i diritti fondamentali o sono utilizzati in modi che sono considerati manipolativi o ingiusti (es. software di sorveglianza di massa o sistemi di IA che impiegano tecniche di *social scoring* da parte dei governi).
- **Rischio Elevato:** vi rientrano sistemi di IA utilizzati in ambiti critici come la sanità, i trasporti, la giustizia e taluni aspetti del governo. Tali sistemi devono soddisfare requisiti rigorosi in termini di trasparenza, *explainability*, sicurezza e supervisione.

- **Rischio Limitato:** per questa categoria è enfatizzato il requisito della trasparenza e vi rientrano, ad esempio i *chatbot*, per i quali gli utenti dovrebbero essere informati che stanno interagendo con un sistema di IA.
- **Rischio Minimo:** sono generalmente quelli che presentano un basso potenziale di danno o impatto sui diritti e le libertà individuali (es. IA utilizzate in giochi, applicazioni di arte generativa o per creare playlist musicali personalizzate, algoritmi impiegati per filtrare i contenuti indesiderati, come lo spam nelle e-mail, purché non riguardino la moderazione dei contenuti su larga scala che potrebbe influenzare la libertà di espressione; software che automatizza compiti ripetitivi e di routine in contesti aziendali, come la gestione di fatture o la programmazione di appuntamenti; sistemi che analizzano grandi set di dati per identificare tendenze di mercato o preferenze dei consumatori, senza prendere decisioni automatiche che hanno un impatto significativo sugli individui).

Per limitarci al tema in esame, vale a dire all'utilizzo dell'IA in ambito giustizia che si colloca nel livello di "*rischio elevato*", l'AI Act stabilisce una serie di requisiti stringenti, che includono la necessità di una valutazione del rischio, elevate garanzie di sicurezza, documentazione dettagliata, trasparenza e informazioni agli utenti, supervisione umana e misure di gestione del rischio. Si richiede che gli utenti siano informati quando interagiscono con un sistema di IA (ad eccezione di casi eccezionali, come le indagini penali), in modo che possano prendere decisioni consapevoli, e sono previsti sorveglianza e controllo del mercato per garantire la conformità con le regole. Quanto alle sanzioni in caso di violazioni, la proposta di regolamento è molto severa, potendo le stesse giungere fino al 6% del fatturato globale annuo dell'impresa. In altre parole, le regole e i requisiti imposti ai vari sistemi di IA sono proporzionati al livello di rischio che questi sistemi presentano per i diritti e le libertà delle persone.

Tale approccio regolatorio, come si coglie già ad un primo e sommario esame, risulta profondamente diverso rispetto a quello scelto nell'ambito del GDPR, che, sebbene incoraggi una valutazione del rischio in determinate situazioni, consta invece di una regolamentazione

uniforme (*One-Size-Fits-All*), che applica quindi le stesse regole a tutti i trattamenti di dati personali, indipendentemente dal livello di rischio.

Risk-based è anche l'approccio adoperato in materia di automazione dei veicoli con guida autonoma (ADAS, *Advanced Driver Assistance Systems*). Nello specifico, seppur nell'ambito di un unico contesto di mercato vengono individuate diverse categorie di rischio. Tale approccio muove da una classificazione operata dalla *Society of Automotive Engineers* (SAE) basata sul livello di automazione dei veicoli vale a dire, guardando l'altra faccia della medaglia, sul grado maggiore o minore di interazione umana con la macchina (analogamente a quanto stabilito nella Risoluzione del 20 ottobre 2020), sicché sono stati definiti sei livelli di automazione, da zero a cinque:

- **Livello 0 - Nessuna Automazione:** Il veicolo non ha sistemi di automazione. Il guidatore ha il controllo completo del veicolo in ogni momento.
- **Livello 1 - Assistenza alla Guida:** Il veicolo può offrire supporto in specifiche funzioni, come la regolazione della velocità (cruise control) o il mantenimento della corsia, ma il guidatore deve rimanere attivamente coinvolto nel processo di guida.
- **Livello 2 - Automazione Parziale:** Il veicolo può assumere il controllo di alcune funzioni, come accelerazione, frenata e sterzata, ma il guidatore deve restare vigile e pronto a intervenire in ogni momento.
- **Livello 3 - Automazione Condizionata:** Il veicolo può gestire tutte le funzioni di guida in specifiche condizioni (come in autostrada) senza l'intervento del guidatore. Tuttavia, il guidatore deve essere pronto a riprendere il controllo quando richiesto dal sistema.
- **Livello 4 - Alta Automazione:** Il veicolo può gestire tutte le funzioni di guida in specifiche condizioni senza alcun intervento umano. A differenza del livello 3, il veicolo in condizioni di livello 4 può continuare a operare in sicurezza anche se il guidatore non risponde alla richiesta di riprendere il controllo.
- **Livello 5 - Automazione Completa:** Il veicolo può gestire tutte le funzioni di guida in tutte le condizioni e su tutti i tipi di strada. Non è necessario alcun intervento umano, e il veicolo è progettato per operare senza un guidatore umano.

Senza entrare nel merito delle varie ipotesi di inquadramento giuridico che sono state proposte, principalmente nel Regno Unito, è interessante osservare come taluni dispositivi ADAS siano in procinto di diventare obbligatori in Europa già dal luglio 2024, in virtù del Regolamento UE 2019/2144: si tratta degli ISA (*Intelligent speed assistant*), della scatola nera, dell'interfaccia di installazione di dispositivi di tipo *alcolock*, dell'avviso della disattenzione e della stanchezza del conducente, dell'avviso avanzato di distrazione del conducente, della segnalazione di arresto di emergenza, del rilevamento in retromarcia e della frenata automatica d'emergenza, il che dimostra, per un verso, l'elevato grado di fiducia nutrito verso tali tecnologie ma, per altro verso, pone il problema delle tutele giuridiche in termini di estrema attualità, quantomeno rispetto alle tecnologie quali quelle illustrate che si collocano dal terzo livello SAE in su e che, quindi, non si affidano completamente all'IA e richiedono, invece, un certo grado di interazione umana.

Sul tema delle tutele, un contributo notevole era stato fornito dal *Report Liability for AI and other digital technologies*, redatto nel 2019 dal Gruppo di Esperti sulla Responsabilità e le Nuove Tecnologie Formazione delle Nuove Tecnologie. Il rapporto giunge alla conclusione che i regimi di responsabilità in vigore negli Stati membri garantiscono una protezione adeguata di base alle vittime, il cui danno sia causato dal funzionamento di tali nuove tecnologie. Si dà atto tuttavia, che le caratteristiche specifiche di queste tecnologie e delle loro applicazioni - complessità, modifica tramite aggiornamenti o autoapprendimento durante il funzionamento, limitata prevedibilità e vulnerabilità alle minacce della sicurezza informatica - possono rendere più difficile garantire il risarcimento dei danni patiti per fatto riconducibile al relativo uso, e vengono conseguentemente suggeriti determinati adeguamenti ai regimi di responsabilità dell'UE e nazionali, in termini di alleggerimento dell'onere probatorio, di statuizioni di obblighi di trasparenza, di registrazione delle attività compiute dalla macchina e di diritto di accesso ai dati registrati, all'obbligo di istituire forme di assicurazione obbligatoria per la responsabilità civile e, infine, alle previsioni che la distruzione dei dati della vittima debba essere sempre considerato un danno risarcibile.

Sul tema, vanno anche ricordate le **due proposte di direttive** del Parlamento e del Consiglio sulla responsabilità per danno da prodotti difettosi ("Proposta PLD") e quella

sull'adeguamento delle norme sulla responsabilità civile non contrattuale all'intelligenza artificiale ("Proposta AILD"). La prima (**PLD - Product Liability Directive**) è un aggiornamento della Direttiva sulla responsabilità dei prodotti difettosi. L'obiettivo è di modernizzare le norme esistenti per tenere conto delle peculiarità dei prodotti basati sull'IA e di altre tecnologie avanzate. Le modifiche proposte mirano a garantire che i consumatori siano adeguatamente tutelati dai danni causati da prodotti difettosi, e contengono anche un aggiornamento delle definizioni di "difetto" e "danno" per riflettere le nuove realtà tecnologiche. Su di essa si è espresso positivamente, l'11 ottobre 2023, il Garante europeo per la protezione dei dati, auspicando che le relative proposte si applichino in tutti i casi di danni causati da un sistema di IA, indipendentemente dalla sua classificazione come ad alto rischio o non ad alto rischio. La Proposta AILD (**AI Liability Directive**) si concentra, invece, sull'adeguamento delle norme sulla responsabilità civile extracontrattuale in relazione all'IA. Il suo scopo è fornire un quadro normativo chiaro per la responsabilità civile derivante dall'uso di sistemi di IA, che include la definizione di chi è responsabile in caso di danni causati da sistemi di IA, sia in termini di responsabilità diretta (produttori e fornitori) sia indiretta (per esempio, utilizzatori di IA).

Il fattore comune delle tutele in predicato di nascita dall'AI Act è rappresentato quindi, in primo luogo, dai requisiti di trasparenza e di *explainability* dei modelli. E da un primo e sommario esame, appare evidente che neanche tali requisiti rappresentano una novità, connettendosi essi all'esigenza di ricostruire il rapporto causale tra un *input* ed un *output* in termini di trasparenza non solo *ex ante* ma anche *ex post*. Tornando al contesto della responsabilità presunta del produttore (nello specifico di sistemi di IA), quello della spiegabilità dei processi rappresenta evidentemente un punto focale dell'indagine giuridica, perché, se un processo non è spiegabile, sarà evidentemente più complesso ricostruire il rapporto causale tra difetto del software e l'eventuale danno conseguenza.

V'è poi da considerare il fattore umano e la relativa interazione con i sistemi di IA che può determinare un'interruzione del rapporto causale, come emerge chiaramente anche dalla Risoluzione 2020/2014 del 20 ottobre 2020.

Il quadro normativo a livello unionale, per quanto magmatico, consente già, conclusivamente, di delineare un perimetro abbastanza chiaro entro il quale gravitano le questioni sul tema della responsabilità del produttore di sistemi AI.

Fatta questa premessa, farei qualche riflessione sul tema di estrema attualità nell'ambito dell'informatica giuridica italiana: **l'intelligenza artificiale applicata alle banche dati della giurisprudenza di merito.**

Al riguardo, giova muovere dal rilievo che, a quanto è dato sapere, la banca dati della giurisprudenza di merito raccoglie gli *abstract* e i provvedimenti civili (sentenze, decreti e ordinanze) provenienti dal Sistema Informatico del Settore Civile (SICI) dal 01/01/2016 ad oggi e, quindi, un patrimonio di poco meno di quattro milioni di pronunce. Le relative funzionalità sembrano essere, quanto alla BDR (banca dati riservata in uso ai magistrati), quella di ricerca giurisprudenziale con linguaggio naturale basata su ChatGPT-3, dell'estrazione di lemmi e la *summarization*, che impegna un processo mediante il quale il modello genera sostanzialmente un riassunto conciso e informativo di un testo più lungo. Tale ultima funzionalità implica diverse abilità e tecniche computazionali, tra cui:

- **Comprensione del Testo:** Il modello deve prima comprendere il testo originale, identificando le idee principali, i temi chiave e le informazioni rilevanti.
- **Identificazione delle Informazioni Essenziali:** Successivamente, il LLM seleziona le parti più importanti del testo, distingue tra informazioni cruciali e dettagli secondari.
- **Riformulazione e Condensazione:** Il modello riformula le informazioni essenziali in un formato più breve, mantenendo l'essenza e la coerenza del testo originale, con la combinazione di diverse frasi o concetti in una formulazione più sintetica.
- **Mantenimento della Coerenza e della Logica:** Durante il processo di riassunto il modello mantiene una narrazione logica e coerente, assicurando che il riassunto sia comprensibile e fedele al testo originale.
- **Stile e Adattabilità:** lo stile del riassunto viene adattato al contesto o alle preferenze dell'utente, come ad esempio un riassunto formale per un contesto accademico o un riassunto più colloquiale per una conversazione informale.

Va a questo punto evidenziato che esiste un legame diretto tra la capacità di *summarization* dei Large Language Models (LLM) e la loro abilità di interpretare gli input espressi con linguaggio naturale, nel senso che la seconda presuppone la prima e viceversa. Tra l'altro, sia la *summarization* che l'interpretazione del linguaggio naturale si basano sull'apprendimento automatico e sul continuo miglioramento del modello tramite l'esposizione a nuovi dati e feedback, che permettono ai LLM di affinare la loro capacità di comprendere e processare il linguaggio naturale nel tempo.

In tale quadro, si è posta estrema attenzione alla tutela della riservatezza, grazie ai medesimi strumenti che consentono la pseudonimizzazione, per taluni versi eccessiva (vengono talora erroneamente oscurate anche le date di riferimenti giurisprudenziali richiamati), delle sentenze.

La riflessione che va fatta sul tema, per quanto concerne i LLM e con specifico riferimento alle banche dati giurisprudenziali è, a mio modo di vedere, duplice: a) la prima attiene ai limiti entro i quali può operare il requisito della trasparenza e dell'*explainability* di siffatti modelli; b) la seconda riguarda la collocazione di siffatta tecnologia all'interno di una delle categorie di rischio delineate dall'AI Act.

Quanto alla spiegabilità, nell'ambito dei Large Language Models (LLM) come GPT-3 o GPT-4 essi sono principalmente dovuti alla complessità e alla natura opaca delle reti neurali su cui questi modelli sono basati. I LLM sono infatti costituiti da milioni o miliardi di parametri. Questa enorme complessità rende estremamente difficile per gli umani comprendere completamente come le decisioni specifiche vengano prese o come vengano generate le risposte. Inoltre, i processi interni di un LLM sono spesso descritti come una "scatola nera": anche se possiamo osservare gli input e gli output, i processi intermedi che portano a un determinato output non sono direttamente osservabili o comprensibili. V'è, inoltre, un profilo di assenza di ragionamento causale: i LLM non "ragionano" nel senso umano del termine; piuttosto, essi generano risposte basandosi su pattern statistici appresi durante il training, per il che non c'è un percorso logico o causale che può essere facilmente spiegato o seguito. L'antidoto ai limiti alla spiegabilità di tali modelli può essere, tuttavia, individuato nella consapevolezza dell'utente, soprattutto quanto alla natura dei LLM,

doendosi prendere preventivamente coscienza che i risultati (*output*) possono essere viziati, oltre che dai *bias* connaturati alla progettazione del modello ed a quelli relativi alla scelta del dataset, anche da vere e proprie “allucinazioni” nelle quali possono incorrere i sistemi di AI basati sul linguaggio naturale.

La seconda perplessità attiene invece alla collocazione della ricerca giurisprudenziale all'interno di una delle categorie di rischio delineate dall'AI Act. È infatti noto che il settore della giustizia si colloca nella categoria di "rischio elevato", e ciò è dovuto al potenziale impatto significativo che i sistemi di intelligenza artificiale possono avere sui diritti fondamentali delle persone quando utilizzati in ambito giudiziario. Si ritiene infatti, giustamente, che l'utilizzo di IA nel settore giudiziario può influenzare diritti fondamentali, come il diritto ad un giusto processo, la privacy, la non discriminazione e la protezione dei dati personali. Mi sia consentito di dubitare che la ricerca giurisprudenziale con strumenti di IA debba, senza se e senza ma, essere collocata nel livello di rischio elevato per il solo fatto che essa utilizza un ampio dataset di provvedimenti giudiziari e che il modello di IA adoperato consenta un'interazione con il materiale di conoscenza attraverso il linguaggio naturale. In altri termini, tale attività non rientra, a mio modo di vedere, necessariamente nella categoria di rischio elevato secondo l'AI Act dell'Unione Europea per diversi motivi:

- 1) I sistemi di IA utilizzati nella ricerca giurisprudenziale forniscono tipicamente supporto informativo e non prendono decisioni autonome. Non si tratta, in altri termini, di giustizia predittiva, mentre il ruolo dell'IA è principalmente quello di assistere gli operatori del diritto nell'individuazione di precedenti rilevanti;
- 2) Nella ricerca giurisprudenziale, gli avvocati ed i giudici utilizzano gli strumenti basati su IA come mero ausilio;
- 3) Da quanto detto ai punti che precedono, diventano centrali la figura non solo dell'operatore “fornitore” del servizio di back-end (la DGSIA), ma anche quella dell'operatore front-end (il magistrato, l'avvocato – se e quando sarà ammesso alla consultazione della BDR che formulano i *prompt*), con le conseguenziali ricadute in termini di responsabilità;

4) Mentre l'uso di IA in altre aree del diritto (come nella predizione delle decisioni o nel *profiling* degli imputati) può avere un impatto diretto sui diritti fondamentali delle persone, la ricerca giurisprudenziale tramite IA si concentra principalmente sull'accesso e sull'analisi delle informazioni.

Conclusivamente, e con specifico riferimento all'utilizzo di strumenti di IA soprattutto nel "dominio giustizia", appare indispensabile un approccio olistico alla responsabilità, che contempli un'attenta valutazione dei potenziali rischi e delle implicazioni etiche legate all'utilizzo dell'IA. In questo contesto, il ruolo della consapevolezza diventa fondamentale. Gli operatori devono essere pienamente informati non solo sulle capacità e i limiti dei sistemi di IA che utilizzano. La crescente integrazione dell'IA in svariati settori impone un'attenta riflessione sull'importanza della formazione degli operatori. La conoscenza approfondita dei principi di funzionamento dell'IA, delle sue potenzialità e dei suoi limiti, è essenziale per garantire un utilizzo responsabile e sicuro di queste tecnologie.

Roberto Arcella

Avvocato, Consigliere dell'Ordine degli Avvocati di Napoli

Componente del Gruppo di Lavoro della FIIF-CNF