



ORDINE DEGLI
AVVOCATI DI MILANO

Talk to the
FUTURE

INTELLIGENZA ARTIFICIALE E AVVOCATURA

Dispensa del Corso di Alta Formazione
dell'Ordine degli Avvocati di Milano
e dell'Associazione Studi Legali Associati (2024-25)

A cura di **Giuseppe Vaciago**
Prefazione di **Antonino La Lumia**

In coorganizzazione con:



Intelligenza Artificiale e Avvocatura
Dispensa del Corso di Alta Formazione dell'Ordine degli Avvocati di Milano e
dell'Associazione Studi Legali Associati (2024-25)

Indice

Indice degli autori	5
Prefazione	7
Capitolo I.....	10
L'intelligenza Artificiale e avvocatura: prospettive e sfide attuali.....	10
1. Introduzione	10
2. Governare l'Innovazione Tecnologica nell'Avvocatura -Avv. Antonino La Lumia ..	11
3. IA, Regolamentazione e Diritti Fondamentali: la Prospettiva Costituzionale - Prof.ssa Marilisa D'Amico.....	13
4. Dentro il Motore dell'IA Generativa: Comprendere la Tecnica per Governarla - Dott. Fabrizio Zanette	17
5. Quattro Storie di Rischio: l'IA come Minaccia e la Difesa del Cliente - Matteo Flora	22
Capitolo II.....	29
Intelligenza Artificiale e professione legale	29
1. Introduzione	29
2. Il bicchiere è sempre pieno: capire l'IA per sfruttarla al meglio - Dott. Gianluca Gilardi	29
2. Il prompting come asset strategico: tecniche di programmazione conversazionale - Ing. Francesco Marinuzzi.....	32
3. L'IA in azione: regole di prompting e sperimentazioni pratiche - Dott. Luca Andreola.....	35
4. L'IA in azione: dimostrazioni pratiche e implicazioni operative - Avv. Martina Elisa Mauloni	37
Capitolo III.....	41
Intelligenza Artificiale e Studi professionali	41

1. Introduzione	41
2. Manuale d'uso del Prompting - Avv. Tommaso Ricci	41
3. Intelligenza artificiale e sistema giudiziario - Prof. Amedeo Santosuosso	45
4. I profili etici e deontologici dell'Intelligenza Artificiale nella professione legale - Prof. Giovanni Ziccardi	47
5. L'intelligenza artificiale nel mondo accademico - Prof.ssa Claudia Sandei	49
Capitolo IV	53
Soluzioni AI per Avvocati nel panorama Legal Tech Italiano	53
1. Introduzione	53
2. Keplera – La piattaforma “all-in-one” per la gestione documentale legale – Dott. Marco Barone.....	53
3. Tiresia – L'IA “super-collaboratore” per il Diritto Tributario – Dott.ssa Marta Elmi Camossi e Avv. Sebastiano Stufano.....	55
4. Gayadeed AI Sphere – L'“AI Hub” sicuro per i dati sensibili – Dott. Egidio Casati..	57
5. Lexroom AI – Dott. Mauro Borselli	59
6. Kopjra– Dott. Tommaso Grotto.....	61
7. Conclusioni – Diversità di soluzioni e consigli operativi per gli avvocati	64
Capitolo V	67
Intelligenza Artificiale e Deontologia Forense: profili etici e responsabilità	67
1. Introduzione	67
2. Demistificare l'IA: funzionamento, rischi ed etica – Dott. Alessandro Vitale	67
3. Responsabilità professionale e doveri deontologici nell'era dell'IA - Avv. Giorgio Treglia.....	70
4. Rischi “invisibili” nell'uso quotidiano dell'IA e strategie di prevenzione - Avv. Valentina Gariboldi	72
5. Verso una regolamentazione dell'IA: il DDL professioni e la sperimentazione nel Tribunale di Catania - Avv. Elio Guarnaccia	75
Capitolo VI	78
L'Intelligenza Artificiale e l'implementazione negli Studi Legali	78
1. Introduzione	78
2. Rimescolare le Carte: l'IA come opportunità per Studi di ogni dimensione - Avv. Giuseppe Vaciago	78

3. Navigare il Mercato Legal Tech: Scelta, Adozione e Rischi nell'Uso dell'IA - Avv. Massimo Sterpi.....	80
4. Make or Buy? Modelli, Costi e Strumenti per una Scelta Consapevole - Dott. Gianluca Gilardi	83
5. L'Esperienza di uno Studio Internazionale: Contesto, Casi d'Uso e Giurisprudenza - Avv. Andrea Tuninetti	85
6. La Filosofia Augmented vs. Automation: un Percorso Strutturato per l'Adozione dell'IA - Avv. Benedetto Lonato.....	89
7. Conclusioni – Distrararsi in un futuro sempre più tecnologicamente ampio	92
 Capitolo VII	94
Innovazione tecnologica nello Studio Legale: persone, processi e AI in azione	94
1. Introduzione	94
2. L’impatto dell’IA generativa in un grande studio legale - Dott. Alejandro Perez. ..	94
3. Superare le barriere al cambiamento: psicologia e strategie per l’innovazione - Ing. Mauro Baldoni.....	98
4. Dal workflow allo strumento: un approccio pratico all’innovazione - Avv. Antonio Giuseppe Di Pietro	102
5. Un approccio “terapeutico” al legal tech: Persone, Processi, Tecnologia – Dott. Marco Mendola.....	105
6. Dall’approccio al metodo - Dott. Gianluca Gilardi.....	108
 Capitolo VIII	111
AI Act e Professione Legale: rischi, trasparenza e strategie di adattamento	111
1. Introduzione	111
2. Coordinate di Lettura dell’AI Act: Rischi, Trasparenza e Paradossi Normativi - Avv. Giuseppe Vaciago	111
3. Miti e Realtà dell’IA: Inquadramento dell’AI Act e della Responsabilità - Prof. Alessandro Mantelero.....	114
4. L’Implementazione dell’AI Act: Governance, Consapevolezza e Ruolo del Giurista - Prof. Edoardo Raffiotta	117
5. L’AI Act nella Pratica Forense: Obblighi, Responsabilità e Scenari Futuri - Avv. Luciano Quarta	120
6. Governance dell’IA nello Studio Legale: Rischi e Opportunità - Prof. Pierluigi Perri	124

7. Conclusioni – Non solo AI Act, ma governance dell’intelligenza artificiale	128
Capitolo IX.....	130
Intelligenza Artificiale tra Privacy e Copyright	130
1. Introduzione	130
2. Rischi, Opportunità e il “Peccato Originale” della Privacy nell’IA - Avv. Giuseppe Vaciago	130
3. GDPR e AI Act: un’interazione complessa - Avv. Lucio Scudiero	132
4. Opere generate dall’IA: tutela, piattaforme e contenziosi - Avv. Lucia Maggi	135
5. Il prompt e la paternità dell’opera: prospettiva USA - Avv. Cristiano Bacchini	139
6. Conclusioni – Implicazioni pratiche per gli avvocati.....	141

Indice degli autori

Dott. Luca Andreola – *Technical Editor per la Legal Business Unit di Lefebvre Giuffrè*
Avv. Cristiano Bacchini – *Coordinatore Commissione Ip Anti Trust dell'Ordine degli Avvocati di Milano*
Ing. Mauro Baldoni – *CEO and Founder Beat 2 Impact*
Dott. Marco Barone – *Co-fondatore della startup legal tech Keplera*
Dott. Mauro Borselli – *Co-ideatore della piattaforma di ricerca giuridica Lexroom AI*
Dott.ssa Marta Elmi Camossi – *Junior Product Manager presso 2Legal Tech*
Dott. Egidio Casati – *Founder & Co-CEO di NymLab*
Prof. Marilisa D'Amico – *Professoressa Ordinaria di Diritto Costituzionale, Università degli Studi di Milano*
Avv. Antonio Giuseppe Di Pietro – *General Counsel di RSM Italy e co-fondatore di 36Brains*
Prof. Matteo Flora – *Founder di The Fool e Professore a contratto all'Università di Pavia*
Avv. Valentina Gariboldi – *Innovation Manager presso Linklaters Italia*
Dott. Gianluca Gilardi – *CEO di LT42, esperto in diritto e tecnologia*
Avv. Elio Guarnaccia – *Consigliere dell'Ordine degli Avvocati di Catania e Founding Partner di FIL Studio Legale*
Dott. Tommaso Grotto – *CEO e co-fondatore di Kopjra*
Avv. Antonino La Lumia – *Presidente dell'Ordine degli Avvocati di Milano*
Avv. Benedetto Lonato – *Equity Partner di LCA Studio Legale e Docente all'Università di Genova*
Avv. Lucia Maggi – *CEO di 42 Law Firm*
Prof. Alessandro Mantelero – *Professore Associato di Law & Technology al Politecnico di Torino*
Ing. Francesco Marinuzzi – *Consigliere Tesoriere della Fondazione Ordine Ingegneri di Roma*
Avv. Martina Mauloni – *Avvocato presso 42 Law Firm*
Dott. Marco Mendola – *Legal Technologist presso TLT LLP (UK)*
Dott. Alejandro Perez – *Head of Knowledge Management & Legal Tech presso Chiomenti*
Prof. Pierluigi Perri – *Professore Aggregato di Informatica Giuridica, Università degli Studi di Milano e Of Counsel di Chiomenti*
Avv. Luciano Quarta – *Managing Partner di Quarta & Partners*
Prof. Edoardo Raffiotta – *Professore Associato di Diritto Costituzionale all'Università di Milano-Bicocca e Of Counsel LCA Studio Legale*
Avv. Tommaso Ricci – *Senior Lawyer e Legal Tech Ambassador presso DLA Piper*
Prof. Claudia Sandei – *Professoressa Ordinaria di Diritto Commerciale, Università di Padova*
Prof. Amedeo Santosuosso – *Professore di Law, Science and New Technologies presso IUSS Pavia*
Avv. Lucio Scudiero – *Counsel presso Legance – Avvocati Associati*
Avv. Massimo Sterpi – *Partner e Co-chair IP & TMT presso Gianni & Origoni*
Avv. Sebastiano Stufaro – *Founding Partner di Stufano Legal e Chairman di 2Legal Tech*
Avv. Giorgio Treglia – *Consigliere dell'Ordine degli Avvocati di Milano e Docente all'Università Statale di Milano*
Avv. Andrea Tuninetti – *Counsel presso Clifford Chance*
Avv. Giuseppe Vaciago – *Founder di 42 Law Firm e Docente al Politecnico di Torino*

Dott. Alessandro Vitale – *Founder della startup di AI conversazionale Conversate*

Dott. Fabrizio Zanette – *Esperto di sistemi di IA generativa per il settore legale*

Prof. Giovanni Ziccardi – *Professore Ordinario di Informatica Giuridica, Università degli Studi di Milano*

Prefazione – Avv. Antonino La Lumia

Viviamo in un'epoca di trasformazioni profonde, in cui l'intelligenza artificiale (IA) sta già ridefinendo processi sociali e professionali – incluso l'esercizio del diritto. La nostra responsabilità, come giuristi, è doppia: da un lato rispondere in modo propositivo all'innovazione, dall'altro garantire che la rivoluzione tecnologica non intacchi mai i principi di umanità e i fondamenti etici e deontologici della nostra professione. L'Ordine degli Avvocati di Milano ha colto per tempo questa sfida, promuovendo un approccio di "fiducia attiva" verso l'IA, nella convinzione che essa non sia una minaccia da cui difendersi, bensì un'evoluzione da governare con consapevolezza. Governare il cambiamento significa rimanere saldamente al timone: l'avvocato deve guidare la tecnologia, mantenendo al centro la propria competenza, la propria etica e – soprattutto – la centralità dell'uomo nelle decisioni. Solo attraverso lo studio e l'aggiornamento continui possiamo utilizzare l'IA come strumento servente e cooperativo, integrato nella nostra attività senza permettere che la macchina sostituisca mai il giudizio umano o ne comprometta l'autonomia. In quest'ottica, innovazione tecnologica e tradizione giuridica non sono in contrasto, ma devono procedere insieme, perché l'una esalti – e non eroda – i valori dell'altra.

È con questa visione che l'Ordine milanese, in collaborazione con l'Associazione Studi Legali Associati, ha promosso un corso di alta formazione dedicato all'intelligenza artificiale applicata al mondo forense. In 10 lezioni, 37 autorevoli docenti hanno accompagnato centinaia di colleghi attraverso le molteplici sfaccettature di IA e diritto. Questa dispensa nasce da quel percorso: ne raccoglie i contributi e li sviluppa in forma organica, offrendo uno sguardo d'insieme completo e approfondito sul tema, un dialogo tra accademia, professione legale e mondo tecnologico. Pagina dopo pagina, emergono voci e prospettive diverse che convergono tutte su un punto fondamentale: per l'avvocatura, comprendere e governare l'IA è una necessità ineludibile. Nei capitoli iniziali vengono delineati il contesto e le sfide attuali: illustri costituzionalisti, informatici ed esperti legali tratteggiano le opportunità offerte dall'IA e i rischi connessi, evidenziando come l'approccio europeo – dal GDPR fino all'AI Act – punti a incanalare lo sviluppo tecnologico entro cornici di garanzie e diritti fondamentali. Sin dall'apertura si ribadisce come la chiave per affrontare la rivoluzione digitale risieda nella formazione: solo un avvocato consapevole del funzionamento dei sistemi di IA, dei loro limiti e delle implicazioni etiche, potrà infatti sfruttarne il potenziale senza esserne sopraffatto. Di qui l'insistenza sul coltivare una vera e propria "AI literacy" per la categoria forense, ossia un'alfabetizzazione diffusa che permetta a ogni avvocato di dialogare con tecnologie avanzate, mantenendo spirito critico e indipendenza di giudizio.

In questo senso, il volume offre non solo nozioni tecniche di base – ad esempio spiegando in termini chiari come funzionano i modelli generativi o il nuovo prompting conversazionale – ma anche spunti pratici su come integrare queste soluzioni nella quotidianità professionale. Vengono presentati casi d'uso concreti e sperimentazioni sul campo: dalla redazione di atti con l'ausilio di assistenti virtuali, all'automazione di attività

ripetitive, fino all'impiego strategico di strumenti di legal tech per la ricerca giurisprudenziale o la gestione documentale. Queste pagine raffigurano l'IA "in azione", nei suoi risvolti più pratici e operativi, fornendo al lettore esempi tangibili di come la tecnologia possa affiancare l'avvocato nel lavoro di ogni giorno. Allo stesso tempo, molta attenzione viene riservata ai profili etico-deontologici. L'introduzione di sistemi di intelligenza artificiale negli studi legali impone infatti nuove riflessioni su responsabilità professionale, tutela della privacy dei dati dei clienti e salvaguardia del segreto professionale. Gli autori sottolineano a più riprese come affidarsi acriticamente a un sistema automatizzato significherebbe tradire la funzione critica e creativa del giurista. Occorre evitare due nemici insidiosi – la pigrizia e l'assuefazione intellettuale – che possono derivare dall'uso indiscriminato di soluzioni automatiche. Al contrario, l'avvocato deve mantenere vivo lo sforzo intellettuale, usando l'IA come ausilio per migliorare efficienza e precisione, ma senza mai abdicare al proprio ruolo di decisore consapevole. La tecnologia va quindi dominata e inserita in un quadro di responsabilità personale: è sempre l'uomo che deve restare "al timone", soprattutto di fronte ai dilemmi etici che l'uso degli algoritmi inevitabilmente pone. In quest'ottica di utilizzo responsabile, il volume discute anche i principi-guida da adottare. Viene ricordato, ad esempio, come il Consiglio dell'Ordine abbia elaborato HOROS, una Carta dei Principi per l'uso consapevole dell'IA in ambito forense (prima in Italia nel suo genere): tale documento individua alcuni pilastri fondamentali – competenza, trasparenza verso i clienti e costante controllo umano sugli output – che ogni avvocato dovrebbe rispettare quando integra strumenti di IA nella propria attività. Sono principi di buon senso e deontologia, ma quanto mai imprescindibili nell'era digitale, per evitare di sacrificare la tutela dei diritti sull'altare dell'efficienza. Solo rispettando questi parametri, l'innovazione potrà dirsi davvero sostenibile e al servizio della giustizia. Uno degli aspetti più originali del percorso formativo raccolto in questo testo è la visione strategica sull'organizzazione dello studio legale. Non si parla solo di tecnologia in astratto, ma di come gestire il cambiamento all'interno delle nostre strutture professionali. Dai contributi emerge chiaramente che l'adozione dell'IA richiede anche un'innovazione culturale: occorre preparare le persone, ripensare i processi e superare le resistenze al cambiamento. Vengono condivise esperienze su come grandi studi internazionali stiano già sfruttando l'IA generativa e su come perfino realtà più piccole possano trarne beneficio, adottando un approccio graduale e consapevole. Si discute se sia meglio sviluppare soluzioni "in house" o affidarsi a fornitori esterni (make or buy), come orientarsi nel mercato delle tecnologie legali e come valutare rischi e benefici di ciascuna innovazione prima di integrarla. In queste pagine troviamo anche riflessioni di taglio manageriale e psicologico: strategie per superare le barriere al cambiamento, importanza di coinvolgere i professionisti più giovani e di formare competenze interne, fino a suggerire una sorta di *approccio "terapeutico"* al legal tech, in cui persone, processi e tecnologia evolvano di pari passo. Si delinea così un quadro in cui l'IA diventa catalizzatore di miglioramento organizzativo, purché vi sia una visione chiara e condivisa dello scopo: non utilizzare l'IA per moda, ma introdurla dove può davvero liberare risorse e aggiungere valore al servizio reso al cliente. L'ultimo tratto del volume è dedicato allo scenario regolatorio e ad alcune questioni giuridiche emergenti. Un intero capitolo

affronta l'AI Act europeo, analizzando rischi, paradossi normativi e obblighi che deriveranno dalla futura disciplina. Questa parte fornisce al lettore una bussola per orientarsi tra classificazioni di rischio, requisiti di trasparenza degli algoritmi e possibili impatti sull'attività forense, aiutandoci a capire come il legislatore intende governare l'IA affinché resti uno strumento affidabile e sicuro. Infine, quasi a chiudere il cerchio, il volume esplora due ambiti dove tecnologia e diritto si incontrano: la privacy e il copyright. Si ragiona sulle implicazioni in materia di protezione dei dati personali quando si utilizzano sistemi di IA – tema quantomai sensibile nell'era dei big data – e si affronta la spinosa questione delle opere generate dall'IA: possono essere tutelate dal diritto d'autore? Di chi è la paternità di un contenuto creato da una macchina? Argomenti un tempo relegati alla fantascienza giuridica stanno rapidamente diventando problemi concreti, e questa dispensa offre alcune chiavi di lettura preziose per iniziare a orientarsi. Senza pretendere di dare risposte definitive, gli autori ci propongono le domande giuste e delineano gli scenari che si aprono, dai possibili contenziosi sul fronte del copyright alle interazioni tra AI Act e GDPR nell'amministrazione della giustizia. Anche sotto questi profili, il messaggio di fondo è chiaro: l'avvocato deve farsi trovare pronto, con mente aperta ma anche con solide conoscenze, per interpretare le norme esistenti in modo adattivo e guidare l'evoluzione delle nuove. In conclusione, l'auspicio è che questo volume possa fungere da bussola e ispirazione per ogni avvocato di fronte alle sfide dell'era digitale. Le pagine che seguono non offrono ricette miracolose – in un campo così in divenire sarebbe impossibile – ma forniscono una mappa aggiornata di ciò che già oggi è realtà e di ciò che si prospetta all'orizzonte. La sfida, per gli avvocati di oggi e di domani, non è respingere l'IA, ma farla propria mantenendo saldi i valori fondanti della professione. L'innovazione tecnologica, infatti, deve tradursi anzitutto in un'innovazione culturale: introdurre strumenti nuovi non basta se non evolve anche la mentalità con cui li utilizziamo. Curiosità verso la tecnologia, spirito critico, etica e senso di responsabilità dovranno procedere di pari passo. Solo in questo modo l'IA potrà integrarsi nel mondo legale in armonia con i principi di giustizia e con la centralità della persona, che restano il faro del diritto. È una sfida che, come Ordine, abbiamo voluto fortemente raccogliere e condividere, accompagnando gli avvocati con strumenti formativi adeguati. La "Carta dei Principi" citata ne è un esempio concreto, così come il progetto Talk to the Future che fa da sfondo: iniziative pensate per promuovere un'alfabetizzazione continua e un osservatorio permanente sulle soluzioni di IA dedicate al sistema giustizia.

Siamo fermamente convinti che solo tramite conoscenza e confronto la categoria forense potrà governare il cambiamento, rimanendo fedele alla propria missione. Ringrazio quindi tutti coloro che hanno contribuito a quest'opera – i relatori, il curatore Avv. Giuseppe Vaciago e i partner di ASLA – e mi auguro che la lettura stimoli in ciascuno di noi nuove idee, consapevolezze e, soprattutto, la voglia di essere protagonisti attivi di questa rivoluzione. In fondo, l'obiettivo è uno solo: fare in modo che, anche nell'era degli algoritmi, il diritto rimanga profondamente umano.

Capitolo I

L'intelligenza Artificiale e avvocatura: prospettive e sfide attuali

1. Introduzione

L'avvento dell'intelligenza artificiale (IA) sta trasformando in profondità la società e, con essa, la pratica del diritto. In un contesto in rapido mutamento, gli avvocati si trovano di fronte a nuove opportunità ma anche a dilemmi etici e strategici. L'Ordine degli Avvocati di Milano ha colto questa sfida promuovendo un percorso formativo dedicato all'IA per i professionisti forensi. Nella lezione inaugurale di tale percorso, diverse voci autorevoli hanno offerto prospettive complementari sul rapporto tra IA e professione legale, delineando un quadro organico delle questioni più rilevanti.

In apertura, l'**Avv. Antonino La Lumia**, presidente dell'Ordine milanese, ha tracciato la visione strategica dell'avvocatura di fronte all'innovazione tecnologica, insistendo sulla necessità di **governare il cambiamento** senza smarrire la centralità dell'uomo. A seguire, la **Prof.ssa Marilisa D'Amico**, costituzionalista, ha inquadrato la regolamentazione europea dell'IA (in primis l'AI Act) nella cornice dei **diritti fondamentali**, illustrando l'approccio basato sul rischio e il ruolo cruciale del giurista nella tutela di questi valori. Il **Dott. Fabrizio Zanette**, esperto nello sviluppo di sistemi di IA ad alto rischio, ha poi aperto il "cofano" della tecnologia, spiegando in termini chiari il funzionamento dei modelli generativi e delle nuove architetture (come i sistemi a **agenti**), così da collegare le radici tecniche ai problemi giuridici emergenti. Infine, **Matteo Flora**, esperto di sicurezza informatica, ha proposto una serie di casi concreti – **quattro storie** apparentemente lontane dalle aule di tribunale – per mostrare come l'IA possa diventare non solo uno strumento di supporto, ma anche una potenziale **arma** capace di esporre i clienti degli avvocati a rischi inediti.

Nei capitoli che seguono si ripercorrono in forma narrativa e ragionata i principali contributi di ciascun relatore. Ciascuna sezione mette in luce i temi chiave emersi, le implicazioni per la professione forense e gli spunti strategici o filosofici offerti, componendo un mosaico coerente delle sfide e delle opportunità che attendono gli avvocati nell'era dell'IA.

2. Governare l’Innovazione Tecnologica nell’Avvocatura -Avv. Antonino La Lumia

Sin dalle battute iniziali del suo intervento, l’avv. Antonino La Lumia – in qualità di presidente dell’Ordine degli Avvocati di Milano – ha voluto trasmettere un messaggio di fiducia attiva nei confronti dell’innovazione. L’intelligenza artificiale, afferma La Lumia, **non è una minaccia** da cui difendersi, bensì un’**evoluzione da governare**. Questo concetto di “governo” implica che la categoria forense deve assumere un ruolo propositivo: non subire passivamente il cambiamento tecnologico, ma guidarlo mettendo al centro la propria competenza ed etica professionale. La chiave per riuscirci è una sola: **formazione e conoscenza**. Solo acquisendo competenze solide e un patrimonio aggiornato di conoscenza tecnologica l’avvocato può utilizzare l’IA come uno **strumento servente e cooperativo**, integrato nella sua attività, senza però permettere che la macchina **sostituisca mai l’uomo** nelle scelte di valore. In quest’ottica, l’IA va vista come un supporto evoluto per automatizzare compiti ripetitivi o non strategici, liberando tempo e risorse, mentre **l’intuito, la creatività, la strategia e la decisione finale devono restare saldamente in mano all’avvocato**. La Lumia avverte: la centralità della decisione umana è un principio non negoziabile, che distingue un uso etico dell’IA da una delega inammissibile di responsabilità alla macchina.

Accanto all’entusiasmo per le potenzialità, La Lumia richiama dunque la necessità di affrontare **una sfida etica e culturale**. L’avvocatura, spiega, deve vincere due nemici insidiosi: la pigrizia e l’assuefazione intellettuale. Affidarsi acriticamente alle risposte di un sistema automatizzato, magari facendosi sedurre dalla comodità, significherebbe tradire la funzione critica e creativa del giurista. Al contrario, l’avvocato deve mantenere vivo lo sforzo intellettuale, usando l’IA solo come ausilio per migliorare efficienza e precisione nelle attività di basso valore aggiunto, ma **senza abdicare al proprio ruolo di decisore consapevole**. La tecnologia, in altri termini, va dominata e inserita in un quadro di **responsabilità personale**: è l’uomo – con la sua coscienza professionale – che deve restare al timone, soprattutto di fronte ai dilemmi etici che l’uso degli algoritmi inevitabilmente pone.

Proprio per indicare una rotta sicura nell’adozione dell’IA, il Presidente milanese ricorda che l’Ordine ha elaborato una **“Carta dei Principi”** per l’uso consapevole dell’intelligenza artificiale, approvata sul finire del 2024. Questo documento valoriale individua tre pilastri fondamentali che ogni avvocato dovrebbe rispettare nell’integrare l’IA nella propria pratica quotidiana. **Primo**, la **competenza**: non è ammissibile utilizzare uno strumento tecnologico senza comprenderne a fondo il funzionamento, i limiti e le implicazioni. Il dovere di competenza impone all’avvocato uno studio continuo: esattamente come non si può applicare una legge senza conoscerla, non si può impiegare un software di IA senza averne almeno una cognizione generale e aver valutato l’affidabilità delle sue fonti e dei suoi algoritmi. **Secondo**, la **trasparenza** verso i clienti: il professionista deve comunicare con chiarezza se e come sta utilizzando sistemi di intelligenza artificiale nello svolgimento dell’incarico. Questo per garantire fiducia e per permettere al cliente di comprendere in che modo viene gestito il suo caso,

evitando opacità che potrebbero minare il rapporto fiduciario. **Terzo, il controllo** (o “presidio” umano): ogni output generato dall’IA deve essere vagliato dall’avvocato, il quale rimane responsabile ultimo delle scelte compiute. La supervisione costante dell’uomo è indispensabile per intercettare eventuali errori o distorsioni dell’IA e per assicurare che lo strumento rimanga al servizio dei fini perseguiti, senza deviazioni indesiderate.

Un altro tema toccato da La Lumia riguarda la **tutela dei dati** e la responsabilità professionale collegata. Nell’era digitale, si dice spesso che i dati sono il “nuovo petrolio”: un asset preziosissimo, da proteggere con estrema cura. L’avvocato, per il ruolo sociale che riveste, custodisce quotidianamente informazioni sensibili e riservate dei propri assistiti, informazioni che attengono alla sfera più intima delle persone e delle imprese. La **vita virtuale** dei clienti – fatta di comunicazioni elettroniche, documenti digitali, tracce lasciate online – è ormai parte integrante (e sempre più estesa) della loro esistenza. Di conseguenza, proteggere questi dati da accessi non autorizzati, perdite o utilizzi indebiti non è solo un adempimento tecnico, ma diventa un **dovere deontologico primario**. La Lumia sottolinea che ogni innovazione tecnologica introdotta in uno studio legale deve essere valutata anche sotto questo profilo: non possiamo sacrificare la riservatezza e la sicurezza dei dati sull’altare dell’efficienza. Ad esempio, se si decide di impiegare un servizio di cloud o uno strumento di IA che elabora dati dei clienti, bisogna assicurarsi che rispetti le normative sulla privacy e che adotti misure di sicurezza adeguate, perché la responsabilità ultima ricade comunque sull’avvocato. In sintesi, **innovare** sì, ma senza mai abbassare la guardia sulla protezione delle informazioni dei propri assistiti.

In conclusione, dalla visione di La Lumia emerge un principio guida netto: **l’innovazione tecnologica deve tradursi anzitutto in un’innovazione culturale**. Gli avvocati del 2025 e degli anni a venire hanno il dovere morale e professionale di far evolvere la professione con lungimiranza, pensando alle prossime generazioni. Non basta introdurre strumenti nuovi se non cambia la mentalità con cui si affrontano. Occorre coltivare la **curiosità intellettuale** verso la tecnologia, ma unita a spirito critico, etica e senso di responsabilità. Solo in questo modo l’intelligenza artificiale potrà essere integrata nel mondo legale in armonia con i principi di giustizia e con la **centralità della persona**, che restano il faro del diritto. La sfida, in definitiva, non è respingere l’IA, ma farla propria mantenendo i valori fondanti della professione: una sfida che l’Ordine di Milano invita ad accogliere, fornendo agli avvocati gli strumenti formativi per riuscirci.

3. IA, Regolamentazione e Diritti Fondamentali: la Prospettiva Costituzionale - Prof.ssa Marilisa D'Amico

L'intervento della Professoressa Marilisa D'Amico ha offerto un ampio sguardo **costituzionale** sul tema dell'intelligenza artificiale, inserendo la regolamentazione emergente (in particolare l'AI Act europeo) nel contesto dei valori e dei diritti fondamentali tutelati dall'ordinamento. D'Amico parte da una constatazione di fondo: stiamo entrando in un **nuovo modello di società digitale**, in cui l'IA agisce come fattore di trasformazione sociale pervasiva. Questo scenario pone sfide inedite alla protezione dei diritti umani, richiedendo un riadattamento degli strumenti giuridici e delle garanzie. **Perché regolamentare l'IA?** – domanda retoricamente la costituzionalista – Non certo per frenare il progresso, ma per **incanalarlo all'interno dei principi fondamentali** che le nostre democrazie si sono date, primo fra tutti la dignità della persona. L'idea cardine è che la tecnologia deve rimanere **sotto controllo umano** e servire lo sviluppo della società senza erodere i diritti e le libertà individuali. In questo senso, l'approccio europeo alla rivoluzione digitale si distingue nettamente da altri modelli globali: **non laissez-faire**, non abbandono della tecnologia a se stessa, ma uno **sviluppo regolato**. L'Europa – ricorda D'Amico – ha scommesso su una strategia in cui regole chiare e valori condivisi possano trasformare i rischi dell'IA in opportunità, creando fiducia nei sistemi e quindi favorendone anche la diffusione economica.

Per rendere concreti questi concetti, D'Amico analizza dapprima il **doppio volto dell'IA**: da un lato le opportunità enormi, dall'altro i rischi reali. Sul fronte delle **opportunità**, l'IA promette benefici in numerosi campi: diagnosi medica più rapida e precisa, trasporti più sicuri ed efficienti (si pensi ai veicoli a guida autonoma), automazione industriale che migliora produttività e sicurezza sul lavoro, gestione avanzata del rischio finanziario e così via. Queste innovazioni potrebbero accrescere il benessere generale, salvare vite, ridurre gli sprechi e aprire nuovi mercati. Tuttavia, queste stesse tecnologie portano con sé **rischi importanti**, soprattutto quando operate senza adeguati controlli. Il rischio maggiore, sottolinea D'Amico, **non è tanto nella tecnologia in sé**, quanto nel suo **agire in modo incontrollato e non trasparente**, sfuggendo alla nostra capacità di tracciarne effetti e criteri decisionali. Vengono citati alcuni studi emblematici che hanno messo in luce **effetti discriminatori** nei primi sistemi di IA adottati: ad esempio, i sistemi di **riconoscimento facciale** che commettevano errori molto più frequenti nel riconoscere volti di donne o di persone con pelle scura, riflettendo i bias nei dati di addestramento (dataset non abbastanza diversificati e team di sviluppo poco eterogenei); oppure i controversi sistemi di **giustizia predittiva** utilizzati in ambito penale negli Stati Uniti, criticati perché tendevano a penalizzare maggiormente determinate minoranze, rivelando pregiudizi razziali latenti nei modelli matematici.

Questi esempi sono la punta dell'iceberg di una sfida più ampia: l'impatto dell'IA può coinvolgere **un'intera gamma di diritti fondamentali**. Non c'è solo in gioco il principio di uguaglianza (messo in pericolo da bias e discriminazioni algoritmiche). D'Amico elenca almeno altri tre settori critici: il **diritto alla salute**, che può essere leso sia direttamente (ad esempio, errori diagnostici di un'IA medica) sia in modo meno ovvio (pensiamo al rischio di dipendenza psicologica da sistemi di assistenza virtuale o alla manipolazione delle abitudini degli utenti tramite algoritmi pervasivi); il **diritto al lavoro**, minacciato da

un'eventuale sostituzione di intere categorie professionali se l'automazione procede senza correttivi e senza politiche di riconversione delle competenze; infine il **diritto di accesso alla giustizia** e, più in generale, la tutela di libertà civili come la libertà di espressione: un massiccio ricorso ad algoritmi nelle piattaforme online può influenzare quali informazioni circolano e quali no, oppure, nel campo legale, sistemi predittivi potrebbero un giorno orientare decisioni processuali, col rischio di opacizzare le basi su cui tali decisioni sono prese. In breve, la rivoluzione dell'IA obbliga i giuristi a un delicato lavoro di **bilanciamento**: massimizzare i benefici senza tradire i principi di base dell'ordinamento.

D'Amico quindi passa a illustrare come l'**Unione Europea** abbia deciso di affrontare questa sfida sul piano normativo, incarnando proprio quella filosofia di "sviluppo regolato" menzionata. Il fulcro dell'azione europea è l'**AI Act**, il regolamento UE 2024/1689, che mira a dettare regole comuni per l'uso dell'IA nel mercato unico, secondo un approccio **basato sul rischio**. La professoressa evidenzia che la definizione di "sistema di intelligenza artificiale" adottata dall'AI Act è volutamente **ampia e tecnologicamente neutra**, così da abbracciare anche le evoluzioni future e impedire facili scappatoie: il legislatore europeo vuole evitare di inseguire la tecnologia in ritardo, preferendo un inquadramento che resti valido nonostante il rapido progresso tecnico.

Il cuore dell'AI Act consiste nella **classificazione dei sistemi di IA per livello di rischio**, con regole differenziate. D'Amico spiega queste categorie in modo chiaro. **Al livello più alto** vi sono i sistemi considerati a "**rischio inaccettabile**": si tratta di applicazioni dell'IA talmente pericolose per i diritti o la sicurezza che il regolamento **ne vieta direttamente l'uso** (art. 5). Esempi citati: i sistemi di **social scoring** sul modello cinese (valutazione massiva delle persone in base ai loro comportamenti o caratteristiche, con effetti sulla vita sociale), le tecniche di manipolazione subliminale dei comportamenti che possano causare danni (ovvero IA che influenzano le scelte di una persona al di fuori della sua consapevolezza), oppure certi utilizzi invasivi del **riconoscimento biometrico a distanza** in tempo reale in luoghi pubblici. Su queste pratiche l'Europa ha già preso posizione netta di divieto, alcuni dei quali – sottolinea D'Amico – sono entrati in vigore dal febbraio 2025. **Al secondo livello** troviamo i sistemi ad **alto rischio** (art. 6): qui rientrano quelle applicazioni che, per contesto d'uso, possono incidere fortemente sulla vita delle persone e pertanto **sono ammesse** nel mercato **ma sottoposte a requisiti rigidissimi** e a procedure di valutazione della conformità prima di poter essere immesse in servizio. Esempi di settori ad alto rischio includono la sanità (un'IA che aiuta nella diagnostica medica, ad esempio), la giustizia (sistemi utilizzati in ambito giudiziario o dalle forze dell'ordine), l'istruzione (si pensi ad algoritmi che valutano esami o profilano studenti) e l'accesso a servizi essenziali o benefici pubblici (ad esempio sistemi che decidono sull'ottenimento di un prestito, di un sussidio, ecc.). In tali casi, l'AI Act prescrive obblighi come trasparenza, tracciabilità, documentazione tecnica rigorosa, qualità dei dati, sorveglianza umana e, come vedremo a breve, **valutazione di impatto sui diritti fondamentali**. **Al terzo livello** vi sono i sistemi a **rischio limitato**, per i quali la legge non impone barriere all'accesso al mercato ma richiede **obblighi di trasparenza** in situazioni specifiche. Ad esempio, i cosiddetti **chatbot** dovranno dichiarare esplicitamente all'utente di essere un'entità artificiale e non un essere umano, così che la persona sia sempre consapevole di interagire con una macchina. Analogamente, se un contenuto

(video, immagine) è generato dall'IA e può essere scambiato per autentico, potrebbe essere obbligatorio segnalare che si tratta di un deepfake o di un'opera artificiale, per prevenire inganni. **Infine**, la grande maggioranza dei sistemi – quelli a **rischio minimo o basso** – non sarà soggetta a requisiti particolari; restano ovviamente i vincoli generali (come il rispetto del diritto civile o penale in caso di danni, e della normativa privacy), ma l'AI Act non aggiunge oneri specifici poiché li considera usi ordinari e poco impattanti (si pensi a un filtro antispam nell'email, o a un sistema di raccomandazione di film su una piattaforma streaming).

D'Amico pone particolare enfasi su un aspetto innovativo dell'AI Act che evidenzia il **ruolo cruciale del giurista: l'obbligo di valutazione d'impatto sui diritti fondamentali** (art. 27 della proposta di regolamento). Per tutti i sistemi di IA classificati come "alto rischio", infatti, non basterà rispettare requisiti tecnici: sarà necessario condurre, prima della loro messa in uso, un'analisi approfondita di come quel sistema possa incidere sui diritti e le libertà garantiti dalla **Carta dei Diritti Fondamentali dell'Unione Europea**. Questo sposta in primo piano le competenze proprie dei giuristi, chiamati a collaborare con gli sviluppatori tecnologici. D'Amico offre un esempio concreto per spiegare la logica di questa valutazione: supponiamo che un'azienda debba scegliere tra due diversi algoritmi di IA per ottimizzare un certo processo produttivo; uno dei due algoritmi è più efficiente (più veloce o preciso nello svolgere il compito) ma **ha un maggiore impatto ambientale** in termini di consumo energetico, l'altro invece è più lento ma "verde", cioè ha un impatto trascurabile sull'ambiente. Ebbene, applicando l'art. 37 della Carta UE (che riconosce il **diritto alla tutela dell'ambiente**), la valutazione d'impatto dovrebbe orientare la scelta verso la soluzione meno inquinante, poiché la sostenibilità ambientale è un valore da tutelare. Solo qualora entrassero in gioco altri diritti di rango superiore – ad esempio il diritto alla vita o alla salute, che potrebbero esigere la massima velocità se parliamo di un algoritmo usato in un dispositivo medico salvavita – si potrebbe giustificare la scelta dell'algoritmo più impattante. Questo esempio illumina la natura **multidisciplinare** e "a bilanciamento" della valutazione d'impatto: non è una mera checklist burocratica, ma un processo che richiede di identificare i possibili effetti dell'IA su vari diritti (dalla privacy, alla non discriminazione, alla dignità, alla sicurezza, ecc.), di prevedere misure per mitigare i rischi, di documentare il tutto in un rapporto e, soprattutto, di aggiornare l'analisi se cambiano le condizioni o emergono nuovi rischi. È un lavoro continuo, nel quale l'avvocato e il giurista sono chiamati a svolgere un compito di **guardiani dei diritti** in sede di innovazione tecnologica.

Verso la conclusione del suo intervento, la Prof.ssa D'Amico allarga lo sguardo alla **sfida strategica** che l'Europa – e l'Italia in particolare – hanno di fronte in materia di IA, facendo riferimento al recente *Rapporto Draghi* sulla competitività europea¹. Questo rapporto, elaborato nel 2023 sotto la guida di Mario Draghi, individua nell'**innovazione tecnologica diffusa** uno dei fattori cruciali per colmare il divario economico e tecnologico che separa l'Europa da altre potenze mondiali. In tale prospettiva, sottolinea D'Amico, un ruolo di primo piano spetta all'**Università** e al sistema formativo. Il Rapporto Draghi evidenzia che i sistemi di istruzione superiore devono diventare più reattivi e allineati alle nuove esigenze di competenze: l'accademia deve farsi motore di un'alfabetizzazione avanzata nel campo dell'IA e del digitale, formando giovani professionisti (non solo informatici, ma anche giuristi, manager, funzionari pubblici)

dotati di un solido “**pacchetto di competenze**” sull’intelligenza artificiale. Ciò significa, ad esempio, inserire nei percorsi formativi degli avvocati elementi di diritto delle nuove tecnologie, di etica dell’IA, ma anche rudimenti sul funzionamento degli algoritmi e sulla gestione dei dati. L’obiettivo è creare una generazione di giuristi che non abbiano paura della tecnologia e che sappiano interagire proficuamente con essa, contribuendo a uno sviluppo economico competitivo ma compatibile con i valori europei.

Parallelamente, D’Amico lancia un messaggio preciso alla **professione legale**: gli avvocati devono accettare la sfida e diventare **protagonisti dell’innovazione**. Se l’Europa vuole davvero trasformare la regolamentazione in un fattore di competitività – creando ad esempio un mercato unico per sistemi di IA sicuri, affidabili e rispettosi dei diritti – c’è bisogno di professionisti consapevoli e preparati, capaci di assistere imprese e istituzioni in questo percorso. L’avvocatura italiana, in particolare, ha l’opportunità di posizionarsi all’avanguardia: D’Amico cita proprio l’Ordine di Milano come **esempio virtuoso**, un’istituzione che si è mossa con tempestività (anche attraverso iniziative formative come il corso in cui questa lezione si inserisce) per colmare il gap di competenze e preparare i suoi iscritti al cambiamento in atto. Iniziative simili, se replicate su larga scala, possono aiutare il Paese a essere non solo ricettore passivo di regole europee, ma anche attore capace di contribuire all’attuazione efficace di quelle regole e allo sviluppo di una **cultura dell’IA** improntata ai nostri principi costituzionali. In definitiva, la prospettiva tracciata dalla Prof.ssa D’Amico è quella di un’innovazione governata dalla conoscenza e dalla responsabilità: università e professioni forensi collaborano per promuovere un ecosistema in cui l’IA sia strumento di progresso economico e sociale, ma sempre **a misura d’uomo** e sotto lo sguardo vigile del diritto.

4. Dentro il Motore dell'IA Generativa: Comprendere la Tecnica per Governarla - Dott. Fabrizio Zanette

Dopo aver affrontato visioni e principi generali, la lezione inaugurale è entrata nel vivo degli aspetti tecnici con l'intervento del dott. Fabrizio Zanette. Forte della sua esperienza nello sviluppo di prodotti di IA ad alto rischio, Zanette si è posto un obiettivo chiaro: **demistificare** la tecnologia dell'IA generativa, spiegando ai giuristi "come funziona davvero" un sistema come ChatGPT. Questa opera di "**apertura del cofano**" tecnologico è fondamentale – sostiene Zanette – perché molti problemi giuridici emergenti (dalla proprietà intellettuale dei modelli, alla responsabilità per i loro output) non possono essere affrontati consapevolmente se non si ha almeno un'idea di cosa accada all'interno di queste "scatole nere". In altre parole, il relatore mira a fornire ai professionisti legali **coordinate tecniche di base**, in modo accessibile ma accurato, affinché possano dialogare con specialisti e clienti sul tema e prendere decisioni informate nell'uso dell'IA.

Zanette inizia con un'affermazione controintuitiva ma illuminante: un **Large Language Model (LLM)** – il tipo di IA generativa alla base di molti strumenti moderni – **non "capisce" veramente il significato** di ciò che dice. Per quanto possa sembrare intelligente e coerente nelle risposte, il modello non ha coscienza né intenzionalità; il suo funzionamento si basa su un principio semplicissimo: **predire la prossima parola**. Più precisamente, date una sequenza di parole (o frammenti di parole, detti *token*) già fornite, l'IA calcola quale sia il token successivo più probabile sulla base delle statistiche apprese durante il suo addestramento. Questo processo di previsione avviene ripetutamente, parola dopo parola, fino a formare frasi e paragrafi completi. L'**intelligenza apparente** che noi percepiamo – la sensazione che la macchina risponda con logica, pertinenza e coerenza – è in realtà una **proprietà emergente** di questo meccanismo puramente probabilistico, non il frutto di una comprensione reale. Ciò significa, ad esempio, che se il modello genera una soluzione a un caso giuridico o scrive una clausola contrattuale, non lo fa perché "sa" il diritto o "ragiona" come un avvocato: sta semplicemente proseguendo la stringa di testo in modo statisticamente plausibile. Questo concetto, spiega Zanette, è fondamentale per l'avvocato: aiuta a non mitizzare l'IA attribuendole qualità umane che non possiede, e a mantenere sempre un occhio critico sugli output, sapendo che **verosimiglianza non equivale a verità o correttezza giuridica**. Un modello generativo può produrre documenti dall'apparenza ineccepibile, ma contenenti errori sostanziali o riferimenti del tutto inventati (le cosiddette "allucinazioni" dell'IA). Comprendere la logica limitata del *predictive text* è dunque il primo passo per esercitare quella vigilanza necessaria di cui parlava anche La Lumia.

Dopo aver illustrato il principio base, Zanette passa a descrivere l'**anatomia di un LLM**, evidenziando i componenti tecnici chiave e collegandoli ai problemi legali. Un modello come ChatGPT è essenzialmente una **rete neurale**, ossia una struttura matematica composta da migliaia (anzi, miliardi) di nodi interconnessi. I **nodi** sono piccoli elaboratori che eseguono calcoli elementari sulle informazioni in ingresso; le **connessioni** tra nodi

trasportano segnali e ognuna di esse è associata a un **peso** numerico. Proprio questi **pesi** rappresentano il cuore del funzionamento: determinano come l'informazione viene trasformata ad ogni passaggio nella rete. In fase di utilizzo (inferenza), l'input fornito dall'utente attraversa strati e strati di nodi e connessioni, e i pesi modulano il segnale finché alla fine viene prodotta una distribuzione di probabilità sulle possibili "prossime parole" da generare. Ma da dove vengono i pesi? Zanette spiega il processo di **addestramento (training)**: inizialmente la rete neurale è come un cervello vuoto, i pesi sono impostati a valori casuali e il modello non sa fare praticamente nulla. Lo si "educa" fornendogli enormi quantità di dati (di solito testi raccolti da Internet, libri, archivi digitali) e facendogli predire parole, confrontando poi le predizioni con quelle corrette. Attraverso algoritmi di ottimizzazione, i pesi vengono aggiustati progressivamente ad ogni errore, finché il modello non diventa abilissimo nel predire la parola successiva in qualsiasi contesto linguistico. Questo processo richiede risorse computazionali gigantesche e ha come prodotto finale proprio quell'insieme di miliardi di pesi ottimizzati che costituiscono, in definitiva, la "conoscenza" del modello.

È a questo punto che Zanette collega la spiegazione tecnica alle **questioni giuridiche** che ne scaturiscono, definendo quella che chiama la "radice" di molti problemi legali dell'IA. Infatti, se l'addestramento utilizza **"tutto Internet"** come base di conoscenza, ci si deve chiedere: **è lecito usare quei dati?** Molti dei testi online sono opere protette da diritto d'autore; il modello, leggendole, ne incorpora in un certo senso informazioni e schemi. Le leggi sul copyright sono messe alla prova da questa nuova situazione: alcuni sostengono che il *data scraping* per addestrare l'IA sia un uso consentito (magari come eccezione per ricerca scientifica o per fair use), altri lo considerano una violazione su larga scala dei diritti degli autori. In parallelo, si pone il tema della **proprietà dei pesi**: una volta terminato l'addestramento, chi "possiede" questo modello addestrato? Da un lato c'è l'azienda che ha sviluppato l'architettura e svolto il training (investendo in hardware e know-how); dall'altro c'è la comunità di autori dei dati di partenza, senza i quali quei pesi non avrebbero assunto i valori che hanno. I pesi sono frutto sia dell'opera umana originaria (i testi) sia dell'elaborazione algoritmica: si possono considerare un nuovo prodotto creativo, di proprietà esclusiva di chi ha creato l'IA, o un'opera derivata collettiva contenente in filigrana conoscenza altrui? La questione è aperta e complessa, con implicazioni anche sulla **proprietà dell'output** generato: se i pesi inglobano conoscenze altrui, il testo prodotto dal modello a chi appartiene? All'azienda titolare del modello, all'utente che l'ha richiesto, o in parte anche agli autori dei contenuti originali che hanno informato il modello stesso? E ancora: in caso di danno provocato da un output (si pensi a un consiglio errato dato da un'IA che causi un pregiudizio economico), **chi è responsabile?** L'utente che ha utilizzato lo strumento, lo sviluppatore del modello, o nessuno dei due perché magari l'IA è considerata solo un mezzo? Questi interrogativi, mostra Zanette, non possono essere risolti senza capire come funziona internamente l'IA, perché concetti tradizionali del diritto devono essere riadattati: ad esempio, stabilire se c'è una violazione di copyright richiede di comprendere se l'output dell'IA è una *rielaborazione* di opere protette o qualcosa di nuovo; valutare la responsabilità

richiede di capire il grado di autonomia/imprevedibilità del sistema rispetto al controllo umano.

Un altro elemento tecnico chiarito nel seminario è il **funzionamento della “memoria” dell’IA** e il ruolo del **contesto** nelle interazioni con modelli come ChatGPT. Zanette spiega che questi sistemi non conservano una memoria permanente delle conversazioni passate con gli utenti (per ragioni sia tecniche che di privacy); ciò che ricordano è limitato a ciò che gli viene fornito di volta in volta come **contesto**. In pratica, ogni volta che poniamo una domanda a un LLM, gli inviamo non solo la domanda ma anche, spesso, l’intera cronologia rilevante del dialogo corrente (ad esempio, i nostri messaggi precedenti e le sue risposte, fino a riempire il limite di *context* che il modello può elaborare). Tutto questo testo funge da memoria temporanea su cui il modello basa la prossima predizione. Ma **non appena la conversazione termina**, e addirittura al completamente di ogni risposta, il modello “dimentica” ciò che è stato detto, perché quei dati non vengono automaticamente immagazzinati nel sistema (a meno che non siano salvati per addestramenti futuri, ma questo è un altro processo). Per i giuristi, comprendere questo meccanismo è importante anche in ottica di **tutela della riservatezza**: se inserisco informazioni riservate di un cliente in una chat con l’IA, devo sapere che quelle informazioni viaggiano al modello e potrebbero essere trattate dai server del fornitore del servizio; inoltre, ad ogni nuova sessione, l’IA non sa nulla del mio cliente a meno che io non glielo fornisca di nuovo. Questa “smemoratezza” progettuale dell’IA è una garanzia (nessuna profilazione occulta oltre il contesto fornito) ma anche un limite (bisogna reintrodurre ogni volta i dati pertinenti). Zanette evidenzia che l’utente deve imparare a gestire bene il contesto: fornire all’IA **prompt** ben formulati e completi (“prompt engineering”), includendo le informazioni necessarie, evitare di condividere dati che non si vuole escano dal proprio controllo, e al contempo fornire tutti i dati necessari a produrre una risposta informata e di qualità (“context engineering”). Inoltre, questo aspetto si collega a temi di **governance dei dati**: alcune aziende potrebbero scegliere di usare versioni locali dell’IA, configurate per avere accesso ai propri dati interni, o customizzati (“fine tuning”) per far sì che il modello “ricordi” procedure o documenti aziendali rilevanti senza doverli inserire ogni volta, mantenendo però il controllo su dove questi dati risiedono.

Con i fondamenti chiariti, Zanette passa infine a illustrare le **evoluzioni più recenti ed avanzate** nel campo dei modelli di IA, come a dire: “ora che sappiamo come funziona un LLM di base, vediamo come lo si sta potenziando per superarne i limiti”. Vengono così introdotti tre concetti chiave: **RAG, Tool Use e Agenti**. La sigla **RAG** sta per *Retrieval-Augmented Generation*, ovvero “generazione aumentata tramite recupero di informazioni”. In pratica, anziché affidarsi solo alla conoscenza interna del modello (che è limitata al suo training e potrebbe essere obsoleta, ad esempio i modelli addestrati nel 2021 non sanno nulla di eventi successivi), si dota l’IA della capacità di **interrogare fonti esterne** durante la conversazione. Così, spiega Zanette, se il modello “si accorge” – in base al prompt – di non avere dati sufficienti per rispondere a una domanda specifica, può effettuare una ricerca in un database aziendale, o su internet, o su un archivio di

documenti fornito dall'utente, ottenere l'informazione mancante e "aumentare" il proprio contesto includendo tale informazione prima di generare la risposta finale. Per esempio, in ambito legale, un LLM integrato con RAG potrebbe, di fronte a una domanda su una sentenza del 2023, andare a cercare il testo di quella sentenza in una banca dati giuridica e poi usarlo per fornire una risposta accurata e aggiornata, cosa che altrimenti non potrebbe fare se il suo training risalisce a prima del 2023. Questa tecnica riduce l'obsolescenza del modello e permette un controllo sulle fonti (si possono limitare le fonti a repertori affidabili, mitigando il problema delle allucinazioni).

Il **Tool Use** (uso di strumenti) è un'altra frontiera interessante: l'idea qui è di dare all'IA la capacità di **compiere azioni esterne** qualora servano per rispondere meglio all'utente. Un caso emblematico, riportato da Zanette, riguarda i calcoli matematici: un LLM standard proverà a risolvere un'operazione complessa predicendone il risultato, non "calcolandolo", e per questo spesso sbaglierà; invece, un LLM dotato di *tool use* potrà, di fronte a una richiesta di calcolo, chiamare un'apposita **calcolatrice software** esterna e restituire il risultato esatto dopo averlo da lei ottenuto. Lo stesso vale per altre operazioni: consultare un calendario, eseguire una traduzione con un motore specializzato, compiere una simulazione, ecc. Dal punto di vista dell'utilizzo professionale, questo apre scenari molto potenti: l'IA diventa un **centro di comando** che può orchestrare vari servizi – ad esempio generare un grafico, interrogare una banca dati finanziaria, inviare una mail – integrando i risultati in una risposta coerente all'utente. L'integrazione di questi tool viene facilitata da standard emergenti, come l'MCP (Model Context Protocol), che consentono agli utilizzatori consapevoli di configurarne facilmente l'utilizzo, estendendo le capacità dell'IA. Per un avvocato, ciò potrebbe significare avere un assistente virtuale che non solo redige bozze, ma magari effettua ricerche giurisprudenziali in banche dati autorizzate, scarica i documenti pertinenti, li riassume e compone un memorandum, il tutto su richiesta e sotto supervisione.

Il passo ulteriore, e attuale frontiera della ricerca, è quello degli **Agenti IA, ossia di LLM dotati di "tool" e di autonomia decisionale su come raggiungere gli obiettivi**, e dell'**Orchestrazione**. Zanette invita a immaginare un sistema dove esistono **molti modelli specializzati**, ognuno un "agente" addestrato o configurato per un compito particolare (ad esempio: un agente esperto in diritto societario, un altro specializzato in linguaggio contabile, un agente abilissimo a pianificare step di un progetto, ecc.). Sopra questi agenti specializzati opera un **agente orchestratore**, ossia un agente specializzato nel coordinare l'attività di sotto-agenti, che ha il compito di scomporre un problema complesso in sottoproblemi, assegnare ciascun compito all'agente più adatto, raccogliere i risultati e combinarli in un output finale coerente. In pratica, non si interagisce più con un singolo modello monolitico, ma con una **squadra di intelligenze artificiali** coordinate. L'esempio descritto è calzante: si chiede al sistema di **redigere un atto di citazione** completo. Un LLM tradizionale genererebbe il testo basandosi su ciò che "ha imparato" e sulle istruzioni ricevute, rischiando errori. Un sistema ad agenti, invece, farebbe qualcosa di diverso: l'agente orchestratore potrebbe esplicitare il suo

piano (“per preparare l’atto di citazione ho bisogno di: 1) individuare i riferimenti normativi e giurisprudenziali rilevanti; 2) ottenere i fatti di causa e organizzarli; 3) scrivere le sezioni in linguaggio appropriato”). Quindi attiverebbe un agente “ricercatore legale” per trovare norme e sentenze pertinenti, un agente “estrattore” per raccogliere i dettagli del caso forniti dal cliente, magari un agente “scrittore” per buttar giù una bozza di atto, e così via, **verificando** passo passo i contributi. Il risultato finale sarebbe una bozza di altissima qualità, ottenuta però con un processo trasparente e modulare, dove l’avvocato potrebbe intervenire per correggere la rotta in ogni sotto-passaggio. Zanette spiega che queste architetture multi-agente stanno diventando realtà nei laboratori più avanzati e rappresentano la prossima dimensione della **collaborazione uomo-macchina**. Non avremo solo un “oracolo” a cui chiedere e che risponde, ma un team di **partner** composto da più specialisti virtuali.

Quali sono le implicazioni di tutto ciò per la professione forense? Zanette invita a riflettere su come evolverà il ruolo dell’avvocato. Con IA sempre più performanti e integrate nei flussi di lavoro, il legale potrà delegare alle macchine non solo compiti semplici (come già avviene, ad esempio, per la ricerca di documenti o la verifica di dati numerici), ma anche **operazioni articolate e strategiche**. Ad esempio, un’analisi due diligence su una mole enorme di contratti potrà essere condotta in prima battuta da agenti IA, che estraggono clausole, segnalano anomalie, confrontano testi; la stesura di un complesso contratto internazionale potrà essere coadiuvata da agenti che suggeriscono formulazioni basate su precedenti e simili. In questo scenario, l’avvocato **non è più solo un produttore di contenuti giuridici**, ma diventa piuttosto un **architetto del sistema** ed un **supervisore critico**. Architetto, perché dovrà saper scegliere quali strumenti e agenti integrare ed impiegare, come configurarli, quali fonti autorizzare e come impostare i parametri di lavoro dell’IA a seconda del caso concreto. Supervisore, perché il controllo finale spetta comunque a lui: dovrà verificare i risultati, individuarne i punti deboli, garantire che il prodotto giuridico finale sia accurato e conforme alle norme e agli interessi del cliente. Si tratta di un’evoluzione del mestiere, non di una sua negazione. Il **valore aggiunto** dell’avvocato in futuro consisterà ancor più nella visione d’insieme, nella capacità di **orchestrare competenze umane e artificiali**, assicurando qualità e affidabilità.

Zanette conclude con un messaggio che si ricollega idealmente a quello di La Lumia: per poter svolgere questo nuovo ruolo, l’avvocato **non può permettersi di ignorare il funzionamento** degli strumenti che utilizza. La comprensione (anche se non specialistica, ma concettuale) di come operano gli algoritmi, di quali sono i loro punti deboli e i loro punti di forza, diventa una **necessità**, non più un vezzo per tecnofili. Solo comprendendo i meccanismi tecnici il giurista potrà davvero **governare la tecnologia** – adattandola ai propri scopi, prevenendone gli abusi – e **sfruttarne appieno le potenzialità** in modo sicuro e responsabile. In sintesi, l’intervento di Zanette è un invito alla **consapevolezza tecnica** come parte integrante del bagaglio del moderno avvocato: perché dietro ogni questione legale legata all’IA (che sia la responsabilità per un errore del software, la privacy dei dati usati, la validità di un contratto generato

automaticamente, ecc.) c'è una realtà tecnologica da capire. La tecnica, diritto e etica, in questo campo, viaggiano fianco a fianco.

5. Quattro Storie di Rischio: l'IA come Minaccia e la Difesa del Cliente - Matteo Flora

L'ultimo intervento offre un cambio di prospettiva netto, spingendo la riflessione oltre le mura dello studio legale e dentro la realtà concreta in cui operano i clienti. Matteo Flora, esperto di sicurezza informatica e di gestione del rischio cibernetico, sceglie infatti un taglio **narrativo e provocatorio**: invece di parlare in astratto di minacce digitali, racconta **“quattro storie che non c'entrano niente con un tribunale”** – come lui stesso le definisce – ma che illustrano efficacemente altrettanti rischi legati all'uso dell'IA nella società odierna. Flora esordisce ricordando brevemente il *noto caso* del Tribunale di Firenze, dove un avvocato, fidandosi ciecamente di un software di intelligenza artificiale per la preparazione di un atto, si era ritrovato a citare precedenti giurisprudenziali **inesistenti**². Quell'episodio, pur eclatante, viene però volutamente accantonato da Flora: l'attenzione non va posta sull'errore del singolo professionista, quanto piuttosto sul contesto in cui i **clienti degli avvocati** (aziende e cittadini) si trovano a interagire con sistemi di IA e a subirne potenzialmente le conseguenze. In altre parole, Flora sposta il punto di vista: dal legale che usa male uno strumento (caso che serve comunque da monito) al mondo esterno dove l'IA può colpire i nostri assistiti, generando **vulnerabilità nuove** che l'avvocato deve imparare a prevenire e contrastare.

Ecco dunque la prima vicenda: **una storia di manipolazione dell'opinione pubblica e di rischio reputazionale**. Flora ci fa immaginare lo scenario di un'azienda farmaceutica che lancia sul mercato un nuovo farmaco, frutto di anni di ricerca e investimento. Questo farmaco potrebbe salvare vite o migliorare la qualità della cura per una certa malattia. Tuttavia, un gruppo di **attivisti ostili** (o forse un'azienda concorrente priva di scrupoli) decide di orchestrare una campagna per distruggerne la reputazione. Come? Utilizzando l'IA su vasta scala. In passato, per diffondere disinformazione occorre tempo, persone, magari budget elevati; oggi **un singolo individuo**, se dotato degli strumenti giusti, può generare in poche ore **migliaia di contenuti malevoli** a costo quasi zero. Nel nostro scenario, l'attacco si svolge così: compaiono improvvisamente decine di **articoli di blog e finti reportage** che descrivono effetti collaterali terribili (ma falsi) causati dal farmaco; nei forum di pazienti e sui social network proliferano **recensioni negative** e testimonianze di persone (inventate) che denunciano l'inefficacia o la pericolosità del medicinale; perfino dei **video deepfake** vengono diffusi, in cui sedicenti medici o pazienti – con volti e voci assolutamente credibili – affermano di aver constatato gravi problemi legati al prodotto. Tutto questo contenuto tossico, generato dall'IA (testi scritti in linguaggio naturale convincente, volti e voci sintetiche indistinguibili dal reale), viene poi amplificato tramite bot sui social media, ottenendo un effetto virale. **L'impatto?** In brevissimo tempo la fiducia del pubblico nel farmaco (e nell'azienda produttrice) è devastata. Medici e pazienti, bombardati da queste notizie, iniziano a dubitare e ad evitare il prodotto; i regolatori aprono un'indagine; le azioni dell'azienda crollano in

Borsa. Il danno economico e di immagine è enorme, e tutto per **una campagna di disinformazione orchestrata digitalmente**.

Flora chiede: quale lezione ne trae un avvocato? La **difesa** degli interessi di un cliente, oggi, **non si gioca più soltanto nelle aule giudiziarie**. Un'azienda può subire attacchi letali fuori dal tribunale, nella piazza virtuale dell'opinione pubblica, e avrà bisogno di protezione anche lì. Questo significa che l'avvocato deve allargare le proprie competenze o, quanto meno, **collaborare con altri specialisti** (esperti di cybersecurity, di gestione della crisi, di comunicazione) per offrire al cliente una tutela a 360 gradi. Ad esempio, potrebbe diventare parte del ruolo dell'avvocato d'impresa consigliare il cliente su come monitorare costantemente il web alla ricerca di campagne diffamatorie o fake news relative all'azienda (ci sono strumenti di IA che fanno anche questo, segnalando picchi anomali di menzioni negative); predisporre in anticipo un piano di **crisis management** legale-comunicativo per reagire prontamente a ondate di disinformazione (smentite ufficiali, diffide legali verso i responsabili identificati, comunicazione trasparente con il pubblico); e, non ultimo, conoscere le **dinamiche tecniche** di queste minacce (come funzionano i bot, come si riconosce un deepfake) per poter istruire efficacemente la propria squadra di risposta. In sintesi, nella prima storia Flora ci mostra l'IA come **arma di manipolazione di massa**, evidenziando che la **reputazione** – un bene intangibile ma cruciale per persone e società – oggi può essere presa di mira con strumenti nuovi. L'avvocato deve quindi farsi trovare pronto a difenderla su nuovi fronti, integrando le competenze giuridiche tradizionali con la comprensione dei meccanismi della disinformazione online.

La seconda storia affronta un tipo di minaccia diverso, più subdolo e **truffaldino**: la **"truffa del finto CEO"** potenziata dall'IA, un caso di *voice phishing* estremamente sofisticato. Lo scenario dipinto da Flora è il seguente: un'impiegata dell'ufficio amministrativo di una società riceve una telefonata urgente. Dall'altro capo del filo c'è, apparentemente, il **direttore generale (CEO)** dell'azienda: riconosce perfettamente la voce del suo capo, la stessa intonazione e accento. Il "CEO" le spiega in modo agitato che è in corso un'operazione finanziaria segretissima e di vitale importanza, da concludere entro pochi minuti, e le ordina di predisporre immediatamente un bonifico di una grossa somma su un certo conto estero, garantendole che tutti i documenti arriveranno a breve ma che intanto bisogna agire senza indugio. Immaginiamo la pressione psicologica su quella dipendente: ha di fronte la voce inconfondibile del suo amministratore delegato, che le chiede di fare presto e di non parlarne con nessuno per ragioni di riservatezza. **La verosimiglianza è totale**, ma in realtà si tratta di una messinscena criminale. Grazie alle tecniche di **clonazione vocale**, un malintenzionato ha registrato (o trovato online) qualche intervento pubblico del vero CEO, ha addestrato un modello di sintesi vocale per replicarne timbro e modo di parlare, e ora sta telefonando fingendosi lui. L'impiegata, purtroppo, cade nel tranello: esegue il bonifico come richiestole. I soldi finiscono su un conto controllato dai truffatori e vengono rapidamente spostati altrove, rendendo difficile recuperarli. L'azienda subisce un danno finanziario diretto ingente, e solo dopo qualche ora si scopre che il vero CEO era all'oscuro di tutto.

Anche questa vicenda, che purtroppo non è solo teoria ma ricalca episodi realmente accaduti a livello internazionale, offre spunti di riflessione importanti per gli studi legali. **Qual è la lezione per l'avvocato?** Se si occupa di diritto penale, di crimini informatici o di compliance aziendale, deve rendersi conto che l'IA ha portato la **classica frode del CEO** a un livello esponenzialmente più insidioso. Un tempo, email sgrammaticate o accenti strani potevano tradire il tentativo di phishing; oggi i truffatori dispongono di **impersonificazioni vocali perfette**. Diventa quindi fondamentale, per i consulenti legali di imprese, **sensibilizzare e predisporre misure preventive**. Flora suggerisce la necessità di implementare procedure di verifica interna a più fattori per ogni operazione straordinaria: ad esempio, stabilire che nessun pagamento oltre una certa soglia possa avvenire senza un doppio controllo, magari con un contatto diretto (di persona o videochiamata) con il dirigente coinvolto, o comunque mai sulla base della sola voce al telefono. L'**autenticazione vocale** non può più essere ritenuta affidabile come un tempo. Un avvocato potrebbe suggerire di adottare codici di sicurezza interni o parole d'ordine per validare comunicazioni d'emergenza, oppure di diffondere linee guida aziendali che dicano: "Se ricevi istruzioni anomale e urgenti, verifica sempre tramite un secondo canale". In caso di truffa avvenuta, l'avvocato dovrà poi assistere l'azienda nelle indagini, nel dialogo con le banche (per tentare blocchi dei fondi), e forse anche in possibili profili di responsabilità interna (se l'impiegato ha agito con troppa leggerezza, l'azienda potrebbe domandarsi come evitare errori simili in futuro). Insomma, questa storia mostra l'IA come **strumento di inganno** capace di superare barriere percettive che davamo per scontate, e insegna che la **sicurezza aziendale** va ripensata includendo queste nuove minacce: è compito anche dei legali partecipare a questa azione, aggiornando protocolli e formazione del personale con l'evolversi della tecnologia.

La terza storia narrata da Flora ha tinte quasi da spy story e riguarda il **furto d'identità e la creazione di persone fittizie**. Immaginiamo – dice l'esperto – che un individuo o un'organizzazione voglia infiltrarsi in un'azienda concorrente, oppure condurre una campagna di disinformazione più raffinata, o ancora realizzare una truffa finanziaria di alto livello. Per fare tutto questo sarebbe molto utile disporre di un'identità falsa, ma **estremamente credibile**, con la quale ingannare interlocutori, costruire relazioni online, raccogliere informazioni sensibili. In passato creare un'identità convincente era difficile: servivano documenti falsi di qualità, foto di qualcuno che somigliasse a una persona esistente, e si rischiava di essere scoperti. Oggi, con l'IA, **chiunque può generare da zero un'identità digitale** completa, che **non esiste nella realtà** ma appare del tutto autentica. Come è possibile? Flora descrive il "kit" tecnologico a disposizione: anzitutto c'è la possibilità di generare un **volto umano fotorealistico** inesistente, grazie a reti neurali generative addestrate su milioni di volti reali (la tecnologia resa popolare dal sito *This Person Does Not Exist*, dove ogni aggiornamento di pagina mostra la foto di un volto di persona che in realtà non è mai nata). Questi volti sintetici sono indistinguibili da una normale fotografia. Poi c'è la creazione di una **voce**: con tecniche simili al voice cloning, o con modelli generativi text-to-speech avanzati, si può ottenere una voce digitale che legga testi con inflessioni naturali, anche qui senza che corrisponda ad alcuna persona reale. Infine, il pezzo forse più inquietante: si può addestrare un'IA (come un modello

linguistico) a scrivere testi – email, post sui social, articoli – **nel modo in cui li scriverebbe una certa persona fittizia**. Ad esempio, si può definire un profilo: Mario Rossi, esperto di cybersecurity di 45 anni, con una laurea in ingegneria, che ha lavorato in varie aziende tech, carattere polemico ma competente. L'IA potrà generare interventi sui forum con il tono tecnico e un po' ruvido di Mario Rossi, rispondere su LinkedIn a discussioni di settore sfoggiando conoscenze approfondite, e così via, costruendo pian piano la **reputazione online** di questa entità virtuale. Dopo qualche tempo, Mario Rossi avrà contatti professionali veri, inviti a convegni, magari accesso a informazioni privilegiate perché considerato un collega affidabile – peccato che dietro di lui ci sia il team di sabotatori iniziali. **L'impatto** di una tale operazione può essere devastante: un *fake* ben piazzato può rubare **segreti industriali** (ottenendoli da contatti ingenui che pensano di parlare con un esperto fidato), può lanciare campagne diffamatorie con maggiore credibilità (perché la voce viene da un "professionista del settore" riconosciuto online), può persino manipolare mercati o opinioni se conquista posizioni di influenza in community chiave.

Qual è l'insegnamento per i legali di questa terza storia? In primo luogo, mai come ora diventa essenziale fare **due diligence sulle identità digitali**. Un avvocato che, ad esempio, assiste un'azienda in una partnership, o che valuta il curriculum di un candidato a una posizione sensibile, o ancora che indaga su un profilo social che ha diffuso informazioni diffamatorie, deve considerare la concreta possibilità che dietro un'identità on-line apparentemente genuina **non ci sia nessuno**, se non un algoritmo. Ciò significa adottare nuovi approcci per la verifica: non basta vedere un profilo LinkedIn completo di foto e storia lavorativa – ora sappiamo che anche quelli possono essere artefatti. Magari si dovrà organizzare un incontro di persona o in videoconferenza (e anche lì, chissà, in futuro con deepfake video in tempo reale serviranno contromisure ulteriori), oppure chiedere referenze incrociate difficilmente simulabili. In breve, il concetto di **fiducia online** va ridiscusso. Inoltre, l'avvocato deve sapersi muovere nel caso in cui scopra di avere a che fare con identità false in contesti illeciti: potrebbe essere necessario lavorare con periti informatici per rintracciare la provenienza di certi account, o con le autorità per denunciare attività di spionaggio e manipolazione. Si entra qui in un terreno a cavallo tra diritto, investigazione digitale e cybersecurity. L'importante messaggio di Flora è che il **grado di complessità** delle frodi e degli schemi illeciti facilitati dall'IA sta aumentando, e di pari passo deve crescere l'attenzione e la competenza di chi tutela gli interessi delle vittime potenziali.

Giungiamo così alla quarta e ultima storia, che tocca un tema di **discriminazione algoritmica e asimmetrie di potere**. Flora ci fa spostare nei panni di un **consumatore qualunque**. Il nostro protagonista deve negoziare le condizioni di un servizio quotidiano – poniamo un piano tariffario telefonico, o le condizioni per un piccolo prestito al consumo – e, invece di parlare con un operatore umano, si trova di fronte a un **chatbot** intelligente, un assistente virtuale dell'azienda. Fin qui nulla di strano: molte società usano chatbot per l'assistenza clienti. Il punto è che questo chatbot, grazie all'IA, ha accesso a **una mole enorme di dati** sul cliente: conosce la sua cronologia di pagamenti,

il suo punteggio di credito, le offerte che ha sottoscritto in passato, persino dati demografici come l'età, la residenza, e chissà, magari ha inferito dal tono delle conversazioni precedenti qualcosa sulla sua personalità (se è propenso a lamentarsi, se tende ad accettare la prima offerta, ecc.). Il cliente chiede uno sconto o condizioni migliori; il chatbot, in apparenza "negoziando", in realtà sta applicando un algoritmo di **profilazione dinamica**: ha calcolato che il nostro cliente vive in un quartiere dove statisticamente le persone cambiano raramente operatore telefonico, e ha notato che in passato ha sempre accettato i rinnovi senza fare shopping di offerte concorrenti. Dunque, l'IA può decidere di **non concedere alcuno sconto**, anzi proporre un prezzo leggermente più alto rispetto a quello che offrirebbe a un altro cliente con un profilo diverso (magari un giovane di altro quartiere che tende a cambiare spesso operatore: a lui l'algoritmo offrirebbe lo sconto perché sa che altrimenti lo perderebbe). Ecco delineata quella che Flora chiama, con un'espressione colorita, la "**tassa sulla stupidità**" – dove "stupidità" è naturalmente dal punto di vista della macchina, che interpreta la fedeltà o la scarsa propensione a cercare alternative come una debolezza da sfruttare. Tecnicamente, è una **discriminazione di prezzo algoritmica**: ciascuno paga condizioni diverse in base a parametri nascosti. È importante notare che qui non parliamo di categorie protette (non è una discriminazione su razza, genere, ecc., cosa che sarebbe esplicitamente illegale), ma di un'**aggregazione di dati** che porta a trattare alcuni clienti meno favorevolmente di altri, sfruttando correlazioni statistiche. Se un essere umano facesse lo stesso – ad esempio scoprisse che un assicuratore fa pagare premi più alti solo perché il cliente abita in una certa zona "povera" – scattarebbero indignazione e forse cause legali per comportamento scorretto; con l'IA, tutto avviene dietro il velo dell'algoritmo, difficile da penetrare e dimostrare.

Questa storia mette in luce un fronte di impegno emergente per gli avvocati: la **contestazione delle decisioni automatizzate**. Sempre più spesso, i clienti (persone fisiche o imprese) si vedranno recapitare esiti, offerte, punteggi, rifiuti o condizioni contrattuali determinate da sistemi di IA. Il legale deve essere pronto a intervenire in loro difesa, chiedendo conto alle controparti di **come** tali decisioni sono state prese. Flora sottolinea che non basta la conoscenza astratta delle normative, pure importanti, come ad esempio l'art. 22 del GDPR che riconosce al cittadino il diritto a non essere sottoposto a decisioni interamente automatizzate senza adeguate garanzie, e il diritto di ottenere spiegazioni sul meccanismo usato. Serve anche la capacità pratica di **sfidare l'algoritmo** sul terreno tecnico: occorre capire quali domande porre per far emergere un eventuale trattamento discriminatorio o ingiusto. Ad esempio, l'avvocato potrebbe chiedere all'azienda: "Quali dati avete usato per profilare il mio cliente? Potete dimostrare che non includete fattori vietati come quelli razziali o di salute? Qual è il modello matematico alla base dell'assegnazione di questa tariffa? Vi rendete conto che clienti con caratteristiche simili ottengono risultati diversi?". Non è semplice, perché spesso le aziende invocano il segreto industriale sugli algoritmi, e perché dimostrare una disparità di trattamento richiede accesso a informazioni che il singolo cliente non ha. Ma è proprio qui che il giurista deve farsi strada, magari coinvolgendo **esperti di audit algoritmico**, e spingendo per maggiore **trasparenza**. Per esempio, in alcuni casi si potrà

ricorrere all'Autorità Garante per la protezione dei dati personali, denunciando la violazione dei principi di correttezza e proporzionalità nel trattamento automatizzato. In futuro, con l'entrata in vigore dell'AI Act, ci saranno probabilmente obblighi più stringenti per le aziende di documentare il funzionamento dei sistemi di IA, e questo potrà dare più armi ai legali per pretendere chiarezza. L'essenziale, come evidenzia Flora, è che l'avvocato **non dovrà più limitarsi alle sole argomentazioni giuridiche** tradizionali, ma dovrà entrare almeno un po' nella **logica dell'algoritmo**, per smontarne eventualmente i meccanismi quando ledono i diritti del cliente. Si prospetta quindi una figura di avvocato che sa dialogare con tecnici, che comprende termini come "modello di machine learning", "bias nel dataset", "discriminazione indiretta via proxy" e così via.

Dopo aver esposto queste quattro storie, Matteo Flora ne tira le fila con una riflessione comune. Tutti gli scenari descritti convergono su un punto: **l'IA sta cambiando la natura dei rischi** a cui sono esposti i clienti (individui, imprese, istituzioni). Minacce prima impensabili – diffamazione massiva automatizzata, frodi iper-realistiche, persone artificiali, trattamenti economici personalizzati in senso svantaggioso – sono oggi realtà o possibilità concrete. Di conseguenza, per un avvocato **non è più sufficiente la sola competenza giuridica classica**. Le conoscenze tradizionali del diritto devono essere integrate con un nuovo set di abilità e sensibilità. Flora ne evidenzia almeno tre. La prima è una solida **consapevolezza tecnologica**: ogni avvocato, pur non essendo un ingegnere, deve informarsi e capire a grandi linee come funzionano queste tecnologie di IA, quali cose possono fare e quali rischi comportano. Ciò include tenersi aggiornato sulle novità (per esempio sapere che esiste la clonazione vocale e come funziona, conoscere il concetto di deepfake, sapere che i chatbot possono profilare gli utenti, ecc.). Solo essendo consapevole l'avvocato può riconoscere in anticipo potenziali problemi e consigliare misure idonee. La seconda abilità è **l'analisi del rischio digitale**: in fase di consulenza legale, bisogna imparare a individuare le nuove **vulnerabilità** che un cliente può avere. Se si assiste una banca, bisogna chiederci: quali attacchi potrebbe subire? (Phishing evoluto, manipolazione di dati, ecc.) Se si assiste un personaggio pubblico: quali rischi corre? (Diffamazione online con deepfake, furto d'identità sui social media...) Questo tipo di *risk assessment* un tempo non rientrava nelle mansioni classiche dell'avvocato, oggi invece diventa parte integrante di una **tutela globale** degli interessi. Infine, la terza competenza è la **capacità di dialogare con la macchina**, per così dire. Significa sia saper **utilizzare gli strumenti di IA** (un po' come accennava Zanette, l'avvocato deve imparare a interagire col chatbot, a dargli i giusti prompt, a usarlo per ricerche, ecc., tutto con occhio critico), sia saper **interrogare e mettere in discussione** le decisioni degli algoritmi altrui, come visto poc'anzi. È un dialogo duplice: di collaborazione, quando l'IA è un nostro strumento, e di confronto-sfida, quando l'IA è dall'altra parte e dobbiamo verificarne l'operato.

In conclusione, Flora dipinge l'immagine dell'**avvocato del futuro** – un futuro molto vicino, praticamente il presente in fieri – come quella di un professionista che sa **proteggere il cliente non solo nel mondo fisico e giuridico tradizionale, ma anche nel mondo digitale**. Un mondo in cui le intelligenze artificiali sono sempre più potenti e

talvolta imprevedibili. Questo avvocato dovrà essere un po' **stratega della sicurezza**, un po' **esperto di tecnologia**, oltre che giurista, e dovrà saper collaborare in team multidisciplinari per affrontare problemi complessi. È una sfida impegnativa, ma anche entusiasmante: chi saprà raccoglierla sarà in grado di offrire un valore enorme ai propri clienti, facendo da ponte tra il diritto e la tecnologia e garantendo che, pur in mezzo alla rivoluzione digitale, i **diritti e gli interessi delle persone rimangano tutelati**.

Capitolo II

Intelligenza Artificiale e professione legale

1. Introduzione

La seconda lezione del **Corso di Perfezionamento in Intelligenza Artificiale per Avvocati** (16 aprile 2025) è stata concepita come un'immersione profonda e pragmatica nel mondo dell'IA applicata al settore legale. L'obiettivo formativo era duplice: da un lato fornire ai professionisti del diritto le basi **concettuali** per comprendere questa tecnologia, dall'altro trasmettere competenze **operative** per utilizzarla efficacemente nella pratica quotidiana. La sessione si è sviluppata in modo progressivo – partendo da un'analisi storico-tecnologica dell'IA per demistificarne il funzionamento, passando attraverso un'esplorazione dettagliata di tecniche avanzate di *prompting*, e culminando in dimostrazioni pratiche dal vivo.

Sin dalle premesse, gli interventi dei relatori hanno mirato a superare la visione semplicistica dell'IA come “**amico o nemico**” dell'avvocato. È stata invece presentata come uno **strumento** potente ma intrinsecamente imperfetto, il cui reale valore dipende interamente dalla capacità dell'utente di guidarlo, controllarlo e integrarlo nel proprio lavoro. Durante la lezione sono state illustrate le differenze tra i vari **modelli di IA**, le strategie per la scelta degli strumenti (soluzioni *open source* vs. piattaforme proprietarie) e le metodologie migliori per strutturare richieste complesse, redigere atti legali e analizzare documenti con l'ausilio dell'IA. In definitiva, la lezione ha inteso dotare l'avvocato di una solida «*cassetta degli attrezzi*» per interagire con l'IA in modo **consapevole, efficiente e sicuro**, preparandolo a un'evoluzione significativa della professione legale nell'era digitale.

2. Il bicchiere è sempre pieno: capire l'IA per sfruttarla al meglio - Dott. Gianluca Gilardi

Gianluca Gilardi, forte della sua duplice esperienza di avvocato e consulente tecnologico, ha aperto la lezione offrendo una visione fondamentale per inquadrare correttamente l'intelligenza artificiale, sfatandone miti diffusi e spiegandone i meccanismi di base. Il suo intervento, dal titolo evocativo “*Il bicchiere è sempre pieno*”, sottolinea un approccio ottimistico ma realistico: l'IA infatti presenta sempre qualche utilità (**bicchiere mezzo pieno**), ma comporta anche limiti e rischi da non ignorare (**bicchiere mezzo vuoto**). In questa prospettiva, Gilardi invita a comprendere a fondo la tecnologia per poterla sfruttare al meglio delle sue potenzialità, evitando sia facili entusiasmi che paure infondate.

Demistificare l'IA – storia e hype: Per prima cosa, Gilardi ha voluto *demistificare* l'intelligenza artificiale ricorrendo a un breve excursus storico. Ha ricordato che l'IA non è affatto una novità recente nata con ChatGPT: vanta una storia di oltre 70 anni, risalente almeno agli studi pionieristici di Alan Turing nel 1950. Nel delineare l'evoluzione del settore, il relatore ha richiamato il concetto di **Hype Cycle** formulato

dalla Gartner, spiegando che l'IA generativa oggi si trova al culmine di un'aspettativa quasi messianica. In altre parole, siamo nel pieno di un "picco delle aspettative esagerate", in cui il clamore mediatico e l'entusiasmo generale rischiano di far apparire l'IA come una soluzione miracolosa a qualsiasi problema. A differenza di altre tecnologie, tuttavia, questo picco nell'IA tende ad autoalimentarsi: l'evoluzione tecnica è talmente rapida che, invece di ridimensionarsi, le aspettative continuano a crescere. Questo rende difficile raggiungere in tempi brevi una valutazione matura e oggettiva delle **reali capacità** e dei **limiti** dell'IA. Gilardi ha inoltre rievocato il periodo passato noto come "**inverno dell'IA**" (negli anni '70-'80), durante il quale la ricerca nell'ambito artificiale subì una forte stagnazione a causa della carenza di potenza computazionale e di dati. Oggi quei due ingredienti fondamentali abbondano, e proprio dalla loro combinazione è scaturita la rivoluzione contemporanea dell'IA, dopo decenni di preparazione sottotraccia.

I due volti dell'IA – simbolica vs. apprendimento automatico: Un punto cruciale su cui Gilardi ha insistito è la distinzione tra due paradigmi fondamentali dell'intelligenza artificiale: l'**IA simbolica** e l'**IA basata sull'apprendimento automatico**. Questa distinzione è particolarmente importante per i giuristi, perché i due approcci presentano caratteristiche di trasparenza e affidabilità molto diverse. L'*IA simbolica* (o *rule-based*) rappresenta l'approccio "classico" dell'IA: un sistema assimilabile a una *white box*, in cui un esperto umano definisce a priori tutte le regole che il programma deve seguire (ad esempio: "*SE si verifica X, ALLORA esegui Y*"). Tutto il comportamento del sistema è quindi determinato da regole logiche esplicite scritte dall'uomo, il che lo rende trasparente e relativamente **spiegabile**. Gilardi fa notare come questo modello di IA simbolica sia concettualmente molto vicino al ragionamento **logico-deduttivo** che ogni avvocato applica nel proprio lavoro quando interpreta norme e crea schemi decisionali.

Al contrario, l'approccio moderno basato su **Machine Learning** (apprendimento automatico) si fonda su algoritmi *data-driven* e rappresenta una sorta di *black box*. In questo caso non vengono fornite alla macchina regole esplicite; le si dà invece in pasto un'enorme quantità di dati dai quali il sistema deve autonomamente estrarre modelli e schemi. Un esempio tipico è l'addestramento di un algoritmo con milioni di email già etichettate come *spam* o *non spam*, lasciando che attraverso reti neurali la macchina "impari" da sola le caratteristiche che differenziano lo spam dal resto. Questo approccio produce sistemi con grandi capacità di **generalizzazione** e adattamento, ma con un importante rovescio della medaglia: il processo decisionale interno non è direttamente ispezionabile né comprensibile dall'uomo in modo semplice. In sintesi, l'IA simbolica offre trasparenza ma richiede conoscenze esperte per impostare le regole, mentre l'IA basata su apprendimento automatico offre prestazioni elevate apprendendo dai dati, però rinunciando in gran parte alla trasparenza sul *come* arriva alle conclusioni. Per un avvocato, comprendere questa differenza è essenziale: ad esempio, un sistema *black-box* di apprendimento automatico potrà fornire risultati sorprendenti, ma bisognerà sempre tener conto che manca di spiegabilità intrinseca, con tutte le implicazioni che ciò comporta in termini di fiducia e responsabilità nell'uso professionale.

Limiti attuali e requisiti dell'IA moderna: Nel delineare lo stato dell'arte, Gilardi ha evidenziato alcuni limiti strutturali e pratici dei moderni sistemi di IA, nonché i requisiti necessari per il loro impiego. In primo luogo, vi sono **costi fissi enormi**: l'addestramento e l'esecuzione dei grandi modelli di deep learning richiedono investimenti economici colossali e infrastrutture computazionali imponenti. Questo è uno dei motivi per cui le soluzioni più avanzate provengono tipicamente da colossi tecnologici o consorzi di ricerca con risorse molto ingenti. In secondo luogo, c'è una **mancanza di conoscenza di dominio**: i modelli generalisti di IA (come quelli addestrati su tutto il web) *sanno un po' di tutto*, ma non possiedono conoscenze specialistiche approfondite in settori altamente verticali come ad esempio il **diritto italiano**. Gilardi li paragona metaforicamente a «*un buon praticante appena uscito dall'università: ha delle nozioni, ma nessuna esperienza pratica*». Cioè, un modello generativo può aver letto migliaia di testi giuridici e sentenze presenti in Internet, ma non avrà comunque l'intuito, il bagaglio esperienziale e la sensibilità specifica di un avvocato esperto in uno studio legale italiano. Questo limite di dominio fa sì che l'IA debba sempre essere affiancata e guidata dal professionista, soprattutto nei compiti più specialistici.

Un'altra riflessione importante riguarda la **variazione dell'efficacia dell'IA a seconda del compito legale**. Gilardi suggerisce una sorta di gerarchia di applicabilità dell'IA nel diritto, notando che ci sono ambiti in cui queste tecnologie brillano e altri in cui mostrano ancora la corda:

1. **Ricerca giuridica:** è il campo in cui l'IA risulta **più efficace**. Un modello può passare al setaccio una mole sterminata di fonti, normative e precedenti giurisprudenziali, trovando correlazioni o riferimenti che potrebbero sfuggire all'occhio umano. Per l'avvocato questo significa velocizzare le ricerche ed eventualmente scoprire argomenti o decisioni rilevanti difficilmente rintracciabili con metodi tradizionali.
2. **Pareri pro veritate e contrattualistica:** in attività come la redazione di **pareri legali** (specie su questioni teoriche di principio) o la stesura di **contratti** standard, l'IA offre già oggi un'**efficacia buona**. Può supportare nell'analisi di clausole tipiche, nel suggerire formulazioni alternative o nel verificare la coerenza di un testo con modelli predefiniti. In pratica, per compiti che hanno un'impronta fortemente testuale e ripetitiva, o dove esistono già formule largamente condivise, l'IA può fornire un aiuto concreto e affidabile, pur dovendo sempre essere verificata.
3. **Atti difensivi (contenzioso):** in questo ambito l'efficacia dell'IA è **ancora limitata**. La preparazione di memorie e atti processuali richiede non solo conoscenza giuridica, ma anche la capacità di elaborare strategie narrative e persuasive, cogliendo **sfumature** del caso specifico, toni, ironie o implicazioni **emotive** che emergono dalle vicende fattuali e dai rapporti umani. Si tratta di elementi extratestuali e di contesto che al momento una macchina non possiede affatto. La **strategia processuale** inoltre implica creatività e responsabilità che non possono essere delegate interamente a un algoritmo. Gilardi, pur riconoscendo alcuni passi avanti anche in quest'area, avverte che per ora affidarsi all'IA per la redazione autonoma di atti difensivi complessi sarebbe imprudente: il contributo

umano esperto rimane centrale e insostituibile nella preparazione della difesa in giudizio.

Verso soluzioni open source e locali: In chiusura del suo intervento, Gilardi ha guardato al prossimo futuro delineando quella che ritiene la direzione più sensata per integrare l'IA nella pratica legale in modo sostenibile e conforme alle normative. Per superare i seri problemi di **compliance** legati all'uso di servizi di IA basati su cloud esterni (problemi soprattutto di riservatezza dei dati e tutela della proprietà intellettuale), la strategia suggerita è puntare sui modelli **open source** eseguibili in locale. In altre parole, invece di inviare informazioni sensibili a servizi cloud di terzi, gli studi legali dovrebbero dotarsi di soluzioni di IA interne: modelli aperti, installati sui propri server **"in casa"**, così da mantenere il **pieno controllo** dei dati e dei processi. Fortunatamente, ha osservato Gilardi, questa strada sta diventando concretamente praticabile perché il costo dell'hardware necessario è drasticamente diminuito negli ultimi anni. Oggi con un investimento relativamente contenuto (nell'ordine di poche migliaia di euro, ad esempio **circa 2.700 €** per una macchina dedicata) anche uno studio legale di medie dimensioni può permettersi l'infrastruttura hardware per far girare localmente un modello di IA sufficientemente potente. Questo approccio **"fai da te"** all'IA offre maggiori garanzie di **sicurezza e riservatezza**, eliminando il rischio che dati legali delicati escano dai confini dello studio, e consente una personalizzazione degli strumenti sul proprio dominio applicativo. In sintesi, Gilardi ha trasmesso ai colleghi un messaggio di cauto ottimismo: l'IA può diventare un prezioso **alleato** dell'avvocato, a patto di conoscerne bene il funzionamento, riconoscerne limiti e ambiti di applicabilità, e adottare soluzioni tecniche che ne permettano un uso etico e sotto controllo.

2. Il prompting come asset strategico: tecniche di programmazione conversazionale - Ing. Francesco Marinuzzi

Il secondo intervento specialistico è stato curato dall'**Ing. Marinuzzi**, esperto informatico con una lunga esperienza nel campo dell'IA. La sua presentazione ha spostato il focus dall'inquadramento generale dell'IA alle tecniche concrete di interazione con i modelli linguistici, elevando il concetto di **prompt** (la richiesta che forniamo all'IA) da semplice domanda estemporanea a vero e proprio **asset strategico immateriale** per lo studio legale. Marinuzzi ha esordito sottolineando che in un mondo in cui i modelli di IA sono sempre più potenti e accessibili, la differenza competitiva tra un professionista e l'altro non la farà tanto l'aver accesso all'IA (destinato a diventare comune), quanto la **capacità di sfruttarla al meglio**. In questa chiave, saper progettare un prompt efficace diventa una forma di **programmazione** della macchina, al punto che egli parla di **"No-Code for Coding"**: la possibilità di *programmare* un modello tramite il linguaggio naturale, senza scrivere una sola riga di codice nel senso tradizionale.

Il prompt è il nuovo software: Marinuzzi paragona quindi il *prompt* a un piccolo programma su misura che **specializza** il modello generico di IA su un compito specifico. Mentre tradizionalmente un software veniva sviluppato con codice per risolvere una

certa categoria di problemi, oggi possiamo spesso ottenere un risultato analogo *istruendo* opportunamente un modello generico attraverso un prompt ben congegnato. L'insieme dei prompt elaborati all'interno di uno studio legale – ad esempio, prompt specifici per redigere un contratto di locazione, per analizzare una sentenza alla ricerca di determinati principi, o per preparare una bozza di atto giudiziario su un certo modello – può arrivare a costituire un patrimonio di conoscenza **unico e di grande valore** per lo studio stesso. Invece di sviluppare un software proprietario, lo studio costruisce una **libreria di prompt** ottimizzati, che riflettono il suo know-how e le sue best practice. Questi prompt raffinati, non disponibili altrove, diventano un asset intangibile che offre un vantaggio competitivo. Si tratta insomma di considerare il *prompt engineering* come una nuova frontiera della conoscenza interna: chi padroneggia l'arte di *mungere* (cioè spremere a fondo) i modelli di IA otterrà risultati migliori e più mirati rispetto a chi si limita a porre quesiti generici.

La società degli agenti e l'importanza del controllo: Spingendosi oltre, Marinuzzi ha descritto uno scenario emergente in cui ai modelli di IA di grande scala si affiancano **Agenti IA** specializzati. Grazie alla tecnica dei prompt, è possibile infatti plasmare versioni personalizzate di un modello (ad esempio istanze specializzate di un Large Language Model) affinché svolgano compiti specifici in autonomia. Si può immaginare di avere una “società” di agenti artificiali cooperanti: per esempio, un agente progettato per **anonimizzare** documenti potrebbe preparare i dati eliminando riferimenti sensibili; un secondo agente, specializzato in **analisi giuridica**, potrebbe esaminare quei documenti anonimizzati individuandone i punti chiave; un terzo agente ancora potrebbe occuparsi di **redigere una sintesi** o un report finale destinato al cliente. Ciascun agente opera su un compito delimitato e circoscritto, massimizzando l'efficacia dell'IA in quel segmento del flusso di lavoro. Tuttavia, man mano che si moltiplicano agenti e automazione, diventa cruciale implementare anche degli “**agenti di controllo**”: in pratica dei sistemi di monitoraggio che supervisionano l'operato degli agenti operativi, verificandone i risultati e intervenendo in caso di output dubbio o incoerente. Questi agenti di controllo fungono da garanti dell'affidabilità, riportando poi l'esito all'essere umano che rimane al centro del processo decisionale. L'idea di Marinuzzi è che, in uno studio legale avanzato, l'avvocato potrà orchestrare una **squadra di agenti IA** ognuno dedicato a un compito (dalla ricerca documentale, alla prima bozza, all'analisi di rischio, ecc.), ma dovrà anche dotarsi di meccanismi di controllo per assicurare che l'intero processo automatizzato resti accurato e conforme alle aspettative. Si delinea così un possibile futuro prossimo in cui parte del lavoro ripetitivo sarà svolto da agenti artificiali altamente specializzati, sotto la stretta supervisione di strumenti e protocolli di controllo di qualità, con l'avvocato nel ruolo di **direttore d'orchestra** di questa società di agenti digitali.

Architettura di un prompt legale avanzato – l'esempio pratico: A supporto teorico di queste idee, Marinuzzi ha presentato un esempio concreto di come costruire un **prompt complesso** per un'attività legale sofisticata (nello specifico, l'analisi di un contenzioso legale). Ha mostrato un'architettura di prompt suddivisa in diversi componenti, ciascuno

con un preciso scopo, quasi fosse la **scaletta** di un programma vero e proprio. I suoi elementi chiave erano:

- **Definizione del ruolo:** specificare fin dall'inizio *"chi"* è l'IA all'interno dello scenario. Ad esempio, il prompt può iniziare con un'indicazione del tipo *"Sei un avvocato di grande esperienza, specializzato nell'analisi e nella preparazione di confutazioni strategiche..."*. In questo modo si imposta il "tono" dell'IA, facendole adottare il punto di vista e il livello di competenza desiderato (in questo caso, quello di un legale esperto, preciso e analitico).
- **Definizione del compito:** descrivere esattamente **cosa** deve fare l'IA. Nell'esempio proposto, si istruisce il modello con frasi come *"Analizza i documenti caricati... Esamina la citazione in giudizio... Separa il querelante dall'imputato e identifica i rispettivi argomenti..."*. Questa sezione del prompt chiarisce il **perimetro dell'attività**, elencando in modo puntuale le operazioni richieste e le informazioni da estrarre o elaborare.
- **Richiesta di chiarimenti:** prevedere una fase in cui, se al modello **manca qualche informazione**, esso non procede in modo errato ma chiede all'utente ulteriori dati. Ad esempio: *"Se ti manca qualche informazione, chiedimi di fornirti chiarimenti aggiuntivi finché non ti è tutto chiaro."* Ciò trasforma l'interazione in un **dialogo iterativo** e riduce il rischio di errori dovuti a input incompleti o ambigui: l'IA viene incoraggiata a colmare le lacune ponendo domande, proprio come farebbe un collaboratore umano diligente.
- **Specifiche sull'output:** indicare con precisione come deve essere la **risposta** dell'IA. Marinuzzi ha suggerito di includere nel prompt direttive del tipo *"Fornisci sempre una giustificazione per ogni conclusione... Prevedi possibili controargomentazioni della controparte... Usa un linguaggio chiaro ma tecnico... Struttura l'analisi in punti elenco logici e sequenziali..."*. In questo modo si dà una **forma** all'output atteso, impostando standard qualitativi (completezza, chiarezza, struttura) e adattandolo al contesto d'uso (ad esempio evitare tecnicismi se il destinatario finale è un cliente non addetto ai lavori, oppure al contrario utilizzare terminologia giuridica appropriata se si tratta di un documento interno).
- **Meccanismo di feedback ricorsivo:** inserire infine una verifica di comprensione da parte del modello stesso. Nell'esempio, il prompt termina chiedendo all'IA: *"Ripeti con parole tue ciò che hai compreso, per confermare che tu abbia capito correttamente le mie richieste."* Questo passaggio forza l'IA a un primo livello di **auto-riflessione**, obbligandola a riformulare il problema. È come farle rileggere la traccia per assicurarsi che abbia interpretato bene ogni punto: un ulteriore livello di controllo che può prevenire fraintendimenti prima che il modello passi a generare la soluzione vera e propria.

Con questo esempio dettagliato, Marinuzzi ha dimostrato che la scrittura di un prompt avanzato assomiglia molto alla stesura di un **piano di lavoro**. Non è qualcosa di estemporaneo, ma un procedimento strutturato in passi successivi, ognuno con una funzione ben precisa. L'avvocato che progetta un prompt in questo modo, quasi fosse un piccolo algoritmo conversazionale, sta in effetti programmando il modello di IA

affinché segua un percorso logico definito, massimizzando le possibilità di ottenere un risultato utile, accurato e mirato alle esigenze specifiche del caso.

3. L'IA in azione: regole di prompting e sperimentazioni pratiche - Dott. Luca Andreola

La terza parte della lezione, curata in tandem dal **Dott. Luca Andreoli** e dall'**Avv. Martina Mauloni**, è stata interamente dedicata a dimostrazioni pratiche dal vivo (*live prompting*) per mostrare concretamente come applicare i principi discussi nella teoria. In questa sezione iniziale, Luca Andreoli ha condiviso le sue *regole d'oro* per una corretta interazione con i modelli di IA e ha contribuito a illustrare alcune sperimentazioni volte a testare i limiti e le capacità effettive di questi strumenti.

Le regole d'oro del prompting efficace: Luca Andreoli ha sintetizzato alcuni principi fondamentali che un avvocato dovrebbe seguire quando formula un prompt, al fine di ottenere dall'IA le risposte migliori possibili. Queste *"golden rules"* del prompting, come le ha chiamate, sono state presentate come segue:

- **Definire il ruolo:** dichiarare esplicitamente *"Chi"* l'IA deve impersonare durante l'interazione. Ad esempio iniziare il prompt con *"Agisci come un avvocato specializzato in diritto amministrativo..."* oppure *"Immagina di essere un assistente legale esperto di contrattualistica internazionale..."*. In tal modo si fornisce al modello un contesto identitario e un punto di vista da adottare, coerente con il tipo di output atteso.
- **Definire l'obiettivo:** chiarire *cosa* si vuole ottenere come risultato finale. Se l'avvocato necessita di una tabella riassuntiva, di un elenco puntato di rischi legali, di un riassunto testuale o magari di una bozza di **post per LinkedIn**, questo va specificato nel prompt (*"L'obiettivo è produrre una tabella comparativa degli obblighi contrattuali..."* oppure *"Genera un riassunto di una pagina in linguaggio semplice per il cliente"*). L'IA deve sapere sin dall'inizio quale forma dovrà avere la sua risposta.
- **Non essere vaghi:** fornire sempre **contesto e dettagli**, non limitarsi a poche parole chiave generiche. Un prompt come *"Riassumi questo contratto"* è troppo scarso; molto meglio sarebbe: *"Riassumi il contratto di locazione allegato evidenziandone gli obblighi principali per il conduttore e per il locatore, in linguaggio non tecnico, in non più di 10 righe."* Più il prompt è specifico e contestualizzato, più alta è la probabilità di ottenere risposte pertinenti e utilizzabili.
- **Fornire esempi (*few-shot learning*):** quando possibile, mostrare all'IA **esempi** del tipo di output desiderato. Andreoli suggerisce di allegare documenti di riferimento o di includere nel prompt un estratto di testo che funzioni da modello. Ad esempio, per far redigere una clausola contrattuale, si potrebbe prima fornire un esempio di clausola simile ben formulata, così che l'IA ne imiti stile e struttura. Questa tecnica, chiamata *few-shot*, orienta il modello attraverso l'imitazione di pattern esistenti.
- **Usare delimitatori chiari:** separare nettamente le varie parti del prompt (istruzioni, contesto, esempi, testo da analizzare) usando simboli o formattazioni

dedicati. Un metodo semplice è racchiudere tra virgolette triple ""..."" eventuali testi da analizzare o brani da non confondere con le istruzioni. In questo modo l'IA distingue ciò che è **input** (documenti, testo grezzo) da ciò che sono le **richieste** dell'utente, evitando fraintendimenti.

- **Evitare le negazioni nelle istruzioni:** formulare sempre le richieste in positivo. Ad esempio, invece di dire *“Non usare un tono troppo sintetico”*, è preferibile dire *“Usa un tono dettagliato e approfondito”*. Le istruzioni in forma affermativa riducono l'ambiguità, mentre le negazioni possono confondere il modello (che talvolta tende a ignorare il “non” e fare esattamente ciò che sarebbe da evitare!). Una comunicazione assertiva e diretta aiuta l'IA a comprendere meglio cosa vogliamo da essa.

Seguendo queste regole d'oro, Andreoli ha evidenziato come il **prompt engineering** diventi una vera e propria competenza strategica per l'avvocato: è la *chiave* per dialogare proficuamente con l'IA. Un prompt ben costruito può fare la differenza tra un output inutilizzabile e uno che invece fa risparmiare ore di lavoro. Al contrario, una domanda posta in modo approssimativo rischia di generare risposte generiche o addirittura fuorvianti (*allucinazioni*), obbligando poi il legale a lunghe revisioni manuali. Investire tempo ed esercizio nell'imparare a formulare i prompt è quindi un percorso altamente consigliato per chi voglia integrare con successo l'intelligenza artificiale nella propria attività.

Esperimento del “cane tigre” – i limiti di comprensione dell'IA: Per testare in diretta le capacità e i limiti dell'IA, Andreoli ha mostrato al sistema (e al pubblico) un'immagine curiosa: la fotografia di un cane sul cui corpo viene proiettata l'ombra di una ringhiera, creando l'illusione visiva di striature simili a quelle di una tigre. L'IA è stata interrogata su cosa vedesse nell'immagine. Il modello ha descritto con sicurezza **“un cane dipinto per assomigliare a una tigre”**, e addirittura ha avviato una dissertazione sulla possibile crudeltà di chi avrebbe colorato il povero animale. Questo esperimento, divertente e a tratti sconcertante, ha dimostrato in modo chiaro un punto essenziale: l'IA **non “comprende” realmente** il contenuto visivo come farebbe un umano, ma si limita ad associare pattern appresi. In questo caso, le strisce scure sul manto del cane sono state interpretate come pittura, probabilmente perché il modello ha richiamato dal suo addestramento immagini di cani effettivamente dipinti o storie di animali verniciati. L'IA ha mancato di cogliere l'inganno dell'ombra e quindi il contesto reale della scena. Questa limitazione evidenzia come i modelli attuali, pur sofisticati, siano essenzialmente **statistici**: non possiedono buon senso né capacità di interpretazione visiva al di là di ciò che hanno già “visto” in fase di training. Per l'avvocato, ciò si traduce in una lezione più generale: *mai sopravvalutare la comprensione dell'IA*, soprattutto quando si affrontano situazioni nuove, sfumate o contenute che escono dalla mera testualità. La supervisione critica dell'umano resta fondamentale proprio perché la macchina può prendere abbagli clamorosi se il problema esce anche di poco dallo schema a cui è abituata.

4. L'IA in azione: dimostrazioni pratiche e implicazioni operative - Avv. Martina Elisa Mauloni

A concludere la lezione è stata l'**Avv. Martina Mauloni**, la quale ha condotto una serie di dimostrazioni pratiche dal vivo volte a mostrare come l'IA possa essere effettivamente impiegata in attività tipiche dello studio legale. Dopo aver appreso principi e tecniche nelle sezioni precedenti, i partecipanti hanno potuto vedere in azione un modello di IA alle prese con compiti concreti di redazione e analisi di documenti giuridici. Mauloni ha preparato tre scenari dimostrativi, nei quali ha utilizzato **prompt avanzati** accompagnati da **documenti allegati**, simulando situazioni reali in cui un avvocato potrebbe trarre beneficio dall'assistenza dell'IA. Queste esercitazioni hanno permesso di evidenziare tanto le potenzialità quanto le modalità corrette di impiego pratico degli strumenti di IA generativa. In particolare, Mauloni ha presentato le seguenti tre demo:

- **Istanza di revoca di sequestro preventivo:** in questo esempio, l'avvocato ha fornito al modello come input sia il *provvedimento di sequestro* da contestare, sia un *modello formulario* di una tipica istanza di revoca. Il **prompt** chiedeva all'IA di redigere una nuova istanza di revoca, basandosi sul modello fornito per lo stile e la struttura, ma al contempo *confutando puntualmente* le motivazioni contenute nel provvedimento concreto. In risposta, l'IA ha prodotto una bozza di atto giuridico completa e ben strutturata: il testo generato riprendeva i **fatti** specifici tratti dal provvedimento (riconoscendo ad esempio i riferimenti temporali, le parti e gli oggetti del sequestro) e li confutava uno ad uno, mentre seguiva la **forma** e l'impostazione argomentativa suggerita dal formulario modello. Il risultato, pur non definitivo, ha mostrato come il sistema possa assistere l'avvocato nella preparazione di una prima bozza molto articolata, riducendo di molto i tempi di stesura.
- **Lettera di risarcimento danni:** nella seconda dimostrazione, Mauloni ha simulato la redazione di una lettera di richiesta danni in ambito civile. Al modello sono stati dati in pasto due elementi: la relazione di un incidente stradale (contenente i fatti, le responsabilità presunte, i danni subiti) e un modello di lettera di richiesta risarcimento. Il prompt in questo caso chiedeva di *estrarre* dalla relazione gli elementi salienti (dinamica dell'incidente, danni materiali e fisici, riferimenti a norme violate, ecc.) e di *compilare* una lettera formale di richiesta danni indirizzata alla controparte o alla sua assicurazione, utilizzando il tono e la struttura del fac-simile fornito. Anche in questa prova l'IA è riuscita a produrre un documento ben organizzato: ha riassunto la **dinamica dei fatti** dall'incarto dell'incidente, inserito i dettagli specifici (luogo, data, parti coinvolte, entità dei danni), e formulato una richiesta di risarcimento coerente, rispettando l'impostazione e il linguaggio tipici della lettera professionale allegata come esempio.
- **Analisi di un Modello 231:** l'ultima esercitazione ha riguardato l'area della *compliance aziendale*. È stato caricato nel sistema un **Modello Organizzativo 231** di una società (documento che descrive il sistema di controllo interno ex D.Lgs. 231/2001), insieme ad alcuni documenti di riferimento esterni rilevanti: le *Linee*

Guida di Confindustria in materia di modelli 231 e la normativa sul **whistleblowing** recentemente aggiornata. Si è quindi chiesto all'IA, attraverso un apposito prompt, di verificare se il modello 231 in esame contenesse tutti i requisiti minimi previsti da legge e linee guida, indicando eventuali mancanze. In questo scenario l'IA ha funzionato da assistente nell'**audit legale**: sfruttando i documenti caricati, ha confrontato i contenuti del modello aziendale con le best practice e vincoli normativi forniti, evidenziando punti di forza e carenze del modello. Ha ad esempio segnalato se mancava una sezione dedicata al whistleblowing, oppure se il sistema disciplinare non era descritto con sufficiente dettaglio, conformemente a quanto richiesto dalla legge.

In tutti e tre i casi, l'IA non ha dovuto "inventare" nulla di propria fantasia, bensì ha **arricchito** le sue capacità generative attingendo ai documenti forniti dall'utente. Si è trattato dunque di generare testi nuovi ma **ancorati** a una base conoscitiva documentale reale e pertinente. L'output dell'IA risulta molto più affidabile e aderente al contesto specifico, perché ogni affermazione può essere fatta risalire a una fonte concreta (il provvedimento giudiziario, la relazione tecnica, la normativa di riferimento). Per l'avvocato questa è una considerazione cruciale: anziché interrogare un modello generale che potrebbe commettere errori o allucinazioni, è preferibile fornirgli i **documenti del caso concreto** (leggi, sentenze, atti di controparte, modelli interni) e chiedere di lavorare *su di essi*. In questo modo l'IA diventa uno strumento di **elaborazione e sintesi** della conoscenza contenuta nei documenti forniti dal professionista, velocizzando enormemente analisi e stesura, ma senza perdere il **collegamento alle fonti** e quindi garantendo una maggiore sicurezza del risultato.

Un ulteriore passo in questa direzione consiste nella creazione di una propria architettura RAG (Retrieval-Augmented Generation), alimentata con normative, sentenze e documentazione interna rilevante per uno specifico settore (ad esempio, in ambito privacy). In tal modo, l'avvocato potrà interrogare il modello sapendo che le risposte saranno costruite a partire dai contenuti inclusi nella RAG stessa, ossia nella base documentale strutturata che egli ha predisposto.

A valle di queste esercitazioni, Mauloni – riprendendo anche i contributi dei colleghi – ha evidenziato alcuni **punti chiave** e considerazioni pratiche che ogni avvocato dovrebbe tenere a mente quando si accinge a usare l'IA nel proprio lavoro quotidiano. In particolare, la lezione si è conclusa rimarcando i seguenti aspetti fondamentali:

- **L'IA non "capisce", calcola probabilità:** un modello di IA generativa non possiede una vera comprensione semantica o intenzionale del testo; esso elabora input e produce output sulla base di un calcolo di probabilità. Questa natura probabilistica spiega sia la sua **potenza** (nel cogliere schemi e connessioni invisibili all'uomo su grandi moli di dati), sia il suo principale **limite**, ovvero la tendenza a generare anche contenuti apparentemente plausibili ma falsi o fuori contesto (*allucinazioni*) quando esce dal dominio su cui è stato addestrato o riceve istruzioni poco chiare.

- **Il prompting è una competenza strategica:** la capacità di formulare prompt chiari, dettagliati e contestualizzati diventa determinante per ottenere risultati di alta qualità e ridurre i rischi di allucinazioni. Non si tratta più solo di saper *cosa* chiedere, ma di *come* chiederlo: il professionista che investe nel migliorare questa abilità potrà sfruttare l'IA in modo nettamente più efficace rispetto a chi improvvisa richieste mal definite.
- **La conoscenza di dominio è insostituibile:** per quanto avanzata, l'IA resta paragonabile al più a un praticante brillante ma inesperto. Solo l'avvocato esperto possiede il *contesto giuridico* completo, sa quali sfumature cercare in un caso, quali riferimenti normativi sono davvero pertinenti e può valutare la *qualità* e la correttezza dell'output. L'IA non rimpiazza il giurista, ma ne estende le capacità; il **giudizio umano** rimane l'elemento decisivo per validare le analisi e prendere decisioni strategiche.
- **Ancorare l'IA ai propri dati:** come mostrato nelle demo, per risultati affidabili e pertinenti è spesso necessario **fornire all'IA dati e documenti specifici del caso su cui lavorare**. Agganciare il modello a una base di conoscenza mirata (contratti, atti, normative rilevanti) permette di contenere le divagazioni e focalizzare le risposte. In pratica, l'IA dà il meglio quando viene utilizzata come strumento di sintesi e collegamento **sui materiali che l'avvocato stesso seleziona**, più che come oracolo generico a cui porre domande nel vuoto.
- **Soluzioni open source e locali per sicurezza e controllo:** la scelta della tecnologia non è neutra. Per proteggere la riservatezza dei dati, è fortemente consigliabile adottare modelli **open source** da scaricare in locale, evitando per quanto possibile l'utilizzo di servizi di AI basati su cloud esterni quando si lavora su dati dei clienti. Ciò riduce i rischi legali legati a privacy e compliance.
- **Sperimentare è necessario:** l'IA è un campo in **rapidissima evoluzione**. Nuovi modelli, servizi e tecniche emergono di continuo. L'invito conclusivo è di non aver timore di sperimentare diverse soluzioni e approcci, partendo magari da compiti interni a basso rischio (ad es. bozze di documenti non critici, ricerche di giurisprudenza, ecc.) per poi estendere l'uso dell'IA man mano che si acquisisce confidenza. Solo attraverso la pratica costante l'avvocato potrà trovare la combinazione più adatta alle proprie esigenze e sviluppare un approccio davvero efficace e personale all'impiego dell'intelligenza artificiale.

In conclusione, la lezione del 16 aprile 2025 ha dipinto un quadro ricco e stimolante dell'integrazione tra **IA e professione legale**. Dalle solide basi teoriche fornite da Gilardi, passando per le strategie operative illustrate da Marinuzzi e Andreola, fino alle prove sul campo guidate da Mauloni, è emerso chiaramente come l'IA stia diventando uno **strumento complementare imprescindibile** per l'avvocato moderno. L'evoluzione tecnologica in atto non relega il giurista a un ruolo marginale, anzi ne ridefinisce le competenze: l'avvocato del futuro prossimo sarà colui che saprà orchestrare sapientemente conoscenza giuridica tradizionale e strumenti di intelligenza artificiale, sfruttando la velocità e la potenza di calcolo di questi ultimi senza mai rinunciare al proprio spirito critico, alla creatività e al rigore deontologico che la professione da sempre richiede. Questo equilibrio dinamico tra uomo e macchina rappresenta, in

definitiva, la chiave per **sfruttare il bicchiere (dell'IA) sempre pieno**, guidando l'innovazione nel solco della qualità e dell'etica professionale.

Capitolo III

Intelligenza Artificiale e Studi professionali

1. Introduzione

La sessione del **14 maggio 2025** del *Corso di Perfezionamento in Intelligenza Artificiale per Avvocati* si è sviluppata lungo due binari paralleli, strettamente interconnessi. Da un lato vi è stato un taglio pratico-operativo, focalizzato sull'uso concreto degli strumenti di **IA generativa** nella professione forense; dall'altro, una riflessione più ampia sull'impatto etico e sistemico che questa tecnologia sta avendo sul mondo accademico e sul sistema giustizia.

La **prima parte**, condotta dall'**Avv. Tommaso Ricci**, ha assunto i toni di un vero e proprio manuale d'uso: Ricci ha illustrato metodologie e tecniche avanzate di *prompting* (interazione efficace con modelli di IA generativa) per redigere pareri e atti, analizzare documenti e ottimizzare la ricerca legale, offrendo strategie concrete per integrare questi strumenti nella pratica quotidiana dell'avvocato.

La **seconda parte** ha ampliato lo sguardo oltre l'operatività immediata. Con il **Prof. Amedeo Santosuosso** si sono esplorate le criticità strutturali che ostacolano una sana adozione dell'IA nel sistema giudiziario (dalla governance alla carenza di competenze, fino al nodo della formazione). Successivamente, con i professori **Giovanni Ziccardi** e **Claudia Sandei**, si è discusso delle profonde trasformazioni in atto nell'università – dalla didattica alla ricerca – con un focus particolare sui principi etici e sui nuovi percorsi formativi resi necessari dall'avvento dell'IA.

Fornire ai professionisti sia competenze operative immediate sia una visione critica e consapevole delle sfide culturali e istituzionali poste dall'intelligenza artificiale all'ecosistema giuridico nel suo complesso è stato l'obiettivo dichiarato della lezione.

2. Manuale d'uso del Prompting - Avv. Tommaso Ricci

L'**intervento di Tommaso Ricci** è stato estremamente pratico e orientato all'operatività. Ricci, esperto di tecnologie applicate al diritto, ha offerto ai colleghi un vero e proprio vademecum sull'impiego dell'IA generativa nella professione forense, mostrando come gli strumenti di intelligenza artificiale possano concretamente migliorare il lavoro quotidiano dell'avvocato.

In apertura, Ricci ha contestualizzato il fenomeno, evidenziando come il mercato dell'AI legale sia in piena espansione. Ha citato un recente report Thomson Reuters, secondo cui **il 41% degli avvocati a livello globale** già utilizza strumenti di IA e **ben il 78%** li considera centrali per il futuro della professione. A riprova di quanto l'IA sia ormai parte integrante del presente (più che un semplice scenario futuro), ha portato l'esempio

provocatorio di uno “*studio legale*” nel Regno Unito recentemente autorizzato ad operare **senza avvocati umani**, occupandosi di recupero crediti di piccolo importo tramite sistemi di intelligenza artificiale. Il messaggio è chiaro: l’IA non è più solo futuribile, ma una realtà con cui l’avvocatura deve confrontarsi **imparando a utilizzarla al meglio**.

Per far comprendere l’atteggiamento corretto nell’adottare queste tecnologie, Ricci ha utilizzato la **metafora del trapano** elettrico: le persone non comprano un trapano per avere un trapano in sé, ma per ottenere “*un buco nel muro*”, cioè il risultato che il trapano produce e che consente loro di appendere meravigliosi quadri ai propri muri. Analogamente, il cliente di uno studio legale non desidera in quanto tali un atto o un parere, bensì **la soluzione efficace a un problema giuridico**. L’intelligenza artificiale va dunque vista come un mezzo per raggiungere più efficientemente quel risultato, non come un gadget da proporre di per sé al cliente. In altre parole, l’IA è uno strumento – potente – al servizio della capacità del legale di risolvere problemi.

Successivamente, Ricci ha illustrato una carrellata di **casi d’uso pratici** dell’IA nel lavoro forense, mostrando la versatilità di questi strumenti in diversi ambiti:

- **Ricerca giuridica potenziata:** a differenza delle tradizionali ricerche booleane per parole chiave, la ricerca condotta con IA generativa avviene in modo semantico, per concetti. Viene così drasticamente velocizzata l’analisi di grandi quantità di testi normativi o giurisprudenziali. Ricci ha citato in proposito un esempio in cui un modello generativo ha analizzato un regolamento europeo di cento pagine: in pochi istanti l’IA è stata in grado di estrarre dati statistici complessi (ad esempio le tipologie più frequenti di violazione previste da quel testo, oppure la distribuzione delle sanzioni per paese) e persino di produrre grafici riassuntivi – un lavoro che manualmente avrebbe richiesto ore di tempo.
- **Revisione di documenti:** anche l’analisi e il proofreading di contratti e atti può essere migliorato con l’IA. Un sistema avanzato è in grado di passare in rassegna le clausole di un contratto e, sulla base di regole predisposte dall’avvocato, *suggerire* formulazioni alternative, evidenziando con codici colore il livello di rischio associato a ciascuna clausola. Ciò risulta utile sia per velocizzare la revisione interna di contratti (individuando subito le clausole problematiche), sia per fornire ai clienti – ad esempio agli uffici acquisti di un’azienda – uno strumento preliminare per valutare eventuali criticità nei testi proposti dalla controparte.
- **Analisi per investigazioni (e-discovery):** l’intelligenza artificiale può setacciare **migliaia di email o documenti** alla ricerca di schemi ricorrenti, termini-chiave o scambi sospetti che un essere umano difficilmente noterebbe in tempi brevi. In ambito di litigation e compliance questa capacità di analisi massiva permette di individuare, ad esempio, flussi anomali di informazioni riservate (spie di una possibile fuga di dati) oppure correlazioni nascoste tra eventi e comunicazioni che potrebbero indicare condotte illecite. L’IA diventa quindi un alleato potente nelle attività di *e-discovery* e indagine interna, riducendo tempi e costi nel filtraggio di enormi quantità di dati.

- **Automazione di contratti e atti:** un ulteriore esempio illustrato è la generazione automatica di documenti legali standardizzati su larga scala. Predisponendo a monte un questionario strutturato per il cliente (ad esempio per raccogliere tutte le informazioni rilevanti ai fini di un contratto di locazione), l'IA può utilizzare le risposte fornite per pre-compilare **decine o centinaia di contratti personalizzati** in pochi secondi, tutti conformi al modello predisposto. Questo tipo di *document automation* consente agli studi legali di gestire volumi elevati di contrattualistica ripetitiva, liberando tempo per attività a maggior valore aggiunto.
- **Valutazione del rischio compliance:** infine, la lezione ha toccato l'uso di strumenti AI per la **verifica automatica della conformità** a normative complesse. Attraverso apposite *checklist* alimentate dall'IA – ad esempio per valutare l'aderenza di un progetto aziendale ai requisiti dell'AI Act (la proposta di Regolamento Europeo sull'Intelligenza Artificiale) o del GDPR in materia di protezione dati – è possibile ottenere un'analisi preliminare dei rischi legali. Il cliente risponde a una serie di domande mirate e il sistema genera un report che evidenzia eventuali criticità e livelli di rischio, che l'avvocato potrà poi approfondire. In tal modo l'IA funge da strumento di **compliance preventiva**, utile sia per i consulenti legali sia per le imprese stesse.

Ciascuna di queste applicazioni, ha aggiunto l'Avvocato Ricci, può essere utile a rendere più veloci, più precise o più affidabili parti delle attività operative svolte da un avvocato, ma **tali applicazioni non possono in nessun caso essere ritenute sostitutive dell'attività di analisi approfondita ed interpretazione sistematica che è propria dell'avvocato**, che ha la responsabilità di non fare cieco affidamento sugli output prodotti dai vari sistemi di AI, ma di rivederli con cognizione di cause e con consapevolezza di come funzionano questi strumenti e quali limiti hanno.

Dopo aver passato in rassegna *cosa* l'IA può fare, Ricci si è concentrato su *come* utilizzarla al meglio, introducendo alcune **tecniche avanzate di prompting**, ossia modalità efficaci per formulare le richieste ai sistemi di IA generativa ottenendo risultati ottimali. In particolare, ha presentato quattro tecniche essenziali, sottolineando che la qualità dell'output dipende in gran parte dalla qualità dell'input fornito:

- **Zero-Shot Prompting** – formulare una richiesta diretta all'IA *senza fornire esempi*. Affinché questo approccio produca buoni risultati, è fondamentale essere estremamente specifici nella domanda, definendo il ruolo che il modello deve assumere, il contesto, il formato dell'output atteso e ogni eventuale vincolo. *Ad esempio:* si otterrà una risposta più utile chiedendo «Sei un assistente legale esperto in contratti commerciali. Riassumi gli obblighi delle parti in questo contratto, in forma di elenco puntato, citando i riferimenti alle relative clausole» piuttosto che un generico «Leggi questo contratto e dimmi di cosa parla».
- **Few-Shot Prompting** – fornire all'IA uno o più **esempi** di output desiderato, così da «istruire» il modello sul formato o sullo stile richiesto. Mostrando esempi concreti, l'IA allineerà meglio la propria risposta alle aspettative dell'utente; questa tecnica risulta utilissima, ad esempio, per estrarre clausole contrattuali in

forma strutturata: basta fornire un paio di esempi di clausole già estratte e il modello comprenderà che deve restituire le informazioni nello stesso schema.

- **Chain-of-Thought** (*Catena di pensiero*) – guidare l’IA attraverso una **sequenza di passaggi logici**, anziché porre una singola domanda sintetica. Invece di chiedere genericamente «*Questa clausola è valida?*», si può istruire il modello a seguire un ragionamento in più fasi: «*Analizza questa clausola considerando: (1) la legge applicabile; (2) l’ambito geografico; (3) la sua durata; (4) fornisci infine una conclusione sulla sua validità, motivandola*». Inducendo l’IA a “pensare a voce alta” passo dopo passo, si ottiene un’analisi molto più accurata e coerente, riducendo il rischio di risposte affrettate o illogiche.
- **Metodo IRAC** – sfruttare il classico schema anglosassone di analisi legale *Issue, Rule, Application, Conclusion*. Ricci suggerisce di chiedere esplicitamente al modello di organizzare la risposta seguendo queste quattro fasi: individuare il **quesito** giuridico da risolvere (*Issue*), richiamare la **norma** o regola applicabile (*Rule*), **applicarla** al caso concreto in esame (*Application*) e formulare infine una **conclusione** motivata (*Conclusion*). In questo modo si induce l’IA a produrre un argomento rigoroso e ben strutturato, molto vicino a quello che un giurista svilupperebbe in una memo o in un parere tradizionale.

A chiudere la sua lezione, Ricci ha proposto anche una **strategia di implementazione** dell’IA in ambito legale, articolata in tre fasi operative. Queste rappresentano un percorso metodico che uno studio professionale può seguire per adottare con successo le nuove tecnologie:

1. **Analisi preliminare (interna ed esterna)** – Valutare da un lato la domanda di mercato e le mosse dei competitor (per capire quali servizi innovativi stanno già offrendo altri studi); dall’altro analizzare i **processi interni** dello studio per identificare quali attività risultano più ripetitive, laboriose o standardizzabili. Questa doppia analisi aiuta a individuare le aree in cui l’IA può avere maggiore impatto e a stabilire priorità d’intervento.
2. **Piano d’azione** – Definire un percorso concreto di adozione, rispondendo a domande chiave: meglio acquistare strumenti già disponibili sul mercato o svilupparne di propri su misura? Quali risorse (umane ed economiche) allocare al progetto? Quali obiettivi specifici perseguire (es. riduzione dei tempi di ricerca, maggiore efficienza nelle due diligence, ecc.) e in che tempistiche? In questa fase si tratteggia una vera roadmap, con **obiettivi misurabili** e scadenze, per passare dall’analisi teorica all’implementazione pratica.
3. **Valutazione dei risultati** – Una volta introdotto uno strumento di IA, è essenziale monitorarne e valutarne l’efficacia con criteri oggettivi. Ricci suggerisce di chiedersi, ad esempio: la soluzione adottata **riduce i costi operativi** o il carico di lavoro su attività ripetitive? **Migliora il servizio** al cliente in termini di qualità o rapidità? Può diventare essa stessa **fonte di profitto** (offrendo nuovi servizi a pagamento)? L’output generato è **accurato e aggiornato** rispetto all’evoluzione normativa? La tecnologia si **integra** bene nei flussi di lavoro esistenti? In base alle risposte, lo studio potrà decidere se estendere l’uso dello strumento, modificarlo o eventualmente sostituirlo.

Ricci ha illustrato che questi elementi fanno parte di una metodologia più ampia elaborata dall'Avvocato, la quale permette ai professionisti di valutare in modo strutturato l'opportunità di digitalizzare o automatizzare determinati processi, anche tramite l'adozione di soluzioni di intelligenza artificiale. L'approccio si focalizza sui processi che possono offrire il maggiore ritorno sull'investimento, tenendo conto delle esigenze specifiche di ciascuno studio.

In definitiva, l'intervento di Tommaso Ricci ha fornito ai partecipanti un ricco ventaglio di tecniche e casi concreti per avvalersi sin da subito dell'IA generativa. Il suo messaggio di fondo è improntato a un sano pragmatismo: l'avvocato non deve temere l'IA, ma **saperla utilizzare criticamente** come un nuovo attrezzo del mestiere, tenendo sempre ben presente che il fine ultimo è servire al meglio i bisogni del cliente.

3. Intelligenza artificiale e sistema giudiziario - Prof. Amedeo Santosuosso

Con il contributo del **Prof. Amedeo Santosuosso** la prospettiva si è spostata dal “micro” (il singolo studio legale) al “macro”. Santosuosso – magistrato emerito e pioniere in Italia negli studi su diritto e tecnologia – ha analizzato le **questioni sistemiche** che, a suo avviso, frenano una sana implementazione dell'IA nel sistema giustizia. Il suo intervento ha messo in luce tre criticità principali, riconducibili ad altrettanti nodi istituzionali e culturali.

(1) Governance – chi decide? Tradizionalmente, la dotazione di strumenti informatici per gli uffici giudiziari è gestita centralmente dal Ministero della Giustizia. Tuttavia, quando la tecnologia in questione è così pervasiva da poter influenzare anche il **contenuto** delle decisioni (come nel caso degli algoritmi di supporto al giudice), delegare ogni scelta al solo Ministero rischia di risultare inadeguato. Manca infatti, secondo Santosuosso, una **cabina di regia condivisa** che coinvolga in modo strutturato **tutte le componenti** del sistema – magistratura, avvocatura, personale amministrativo – nella definizione delle strategie di adozione dell'IA. Emblematico, in negativo, l'esempio di una *banca dati di merito* (ossia un database di sentenze di primo e secondo grado) sviluppata in passato in modo differenziato per i giudici e per gli avvocati: una scelta che tradiva una visione separata e non collaborativa tra attori che invece perseguono lo stesso obiettivo di giustizia. Per evitare simili distorsioni, è indispensabile immaginare **meccanismi di governance inclusivi**, in cui la magistratura e l'avvocatura partecipino attivamente – accanto al Ministero – alle decisioni su quali sistemi di IA adottare, con quali finalità e sotto quali controlli.

(2) Gap di competenze – il dialogo impari con i fornitori. Un secondo problema risiede nel **divario di competenze tecniche** tra chi propone la tecnologia e chi la dovrebbe adottare nel settore pubblico. Quando il Ministero (o in generale la PA) si interfaccia con i grandi fornitori di soluzioni di IA – spesso colossi multinazionali dell'ICT – parte da una situazione di debolezza non solo contrattuale ma soprattutto **conoscitiva**. All'interno delle istituzioni giudiziarie scarseggiano figure professionali con competenze

informatiche avanzate, in grado di valutare criticamente le proposte tecnologiche e dialogare **alla pari** con gli specialisti delle aziende fornitrici. Questa situazione comporta due rischi: da un lato, possono essere acquistati sistemi non ottimali o inadatti, perché manca la capacità di **valutarne tecnicamente** la qualità o anche solo di specificare correttamente le esigenze nei capitolati di appalto; dall'altro, di fronte a problemi tecnici complessi durante l'implementazione, l'ente pubblico potrebbe non avere personale in grado di comprenderli e risolverli, restando dipendente al 100% dal fornitore. Colmare questo gap di competenze è dunque fondamentale: significa investire in formazione tecnica per il personale pubblico, assumere consulenti esperti o creare strutture ad hoc che supportino la PA nelle scelte tecnologiche, così da riequilibrare il rapporto con i vendor privati e assicurarsi soluzioni davvero funzionali alle esigenze della giustizia.

(3) Formazione – un investimento mancato. Il terzo nodo cruciale individuato da Santosuosso è la **manca di adeguata formazione** del personale giudiziario sull'uso delle nuove tecnologie. Storicamente, nota il professore, l'approccio è stato spesso quello di **investire in hardware e software**, distribuendo nuovi strumenti nei tribunali, **senza investire in parallelo sulla formazione** di chi quegli strumenti avrebbe dovuto usarli. Il risultato prevedibile è un **sottoutilizzo** (quando non un uso scorretto) delle dotazioni informatiche: magistrati e cancellieri che ignorano funzionalità avanzate, o che addirittura evitano di usare il nuovo sistema perché non si sentono sicuri nell'adoperarlo. Santosuosso sottolinea che la formazione tecnologica non può più essere sporadica né lasciata alla buona volontà del singolo: deve diventare **continua, strutturale e istituzionalizzata**, esattamente come l'aggiornamento professionale su codici e procedure. E il problema, aggiunge, parte già dall'università: nei corsi di laurea in Giurisprudenza **l'informatica giuridica** è ancora troppo spesso un insegnamento opzionale e marginale, se non assente del tutto. Così i giovani laureati entrano in magistratura o nell'avvocatura privi di basi tecniche, alimentando un *"circolo vizioso della non conoscenza"* che poi prosegue lungo tutta la carriera. Rompere questo circolo vizioso richiede uno sforzo concertato: inserire seriamente competenze digitali nei percorsi formativi dei giuristi, prevedere programmi di training specifici quando nuovi strumenti vengono introdotti nei tribunali, e coltivare una cultura professionale che veda l'aggiornamento tecnologico come parte integrante dell'identità del giurista moderno.

L'analisi di Santosuosso mette in luce come le sfide poste dall'IA in ambito giudiziario non siano soltanto di natura tecnica, ma soprattutto **istituzionale e culturale**. Senza adeguate **strutture decisionali**, senza colmare i **divari di competenze** e senza un investimento serio in **formazione**, anche le migliori tecnologie rischiano di restare inefficaci o addirittura di produrre effetti indesiderati. Il messaggio è che l'innovazione nel diritto va governata con visione sistemica: solo così l'IA potrà davvero mantenere le promesse di efficientamento e miglior accesso alla giustizia, anziché generare nuove disparità o inefficienze.

4. I profili etici e deontologici dell'Intelligenza Artificiale nella professione legale - Prof. Giovanni Ziccardi

Il **Prof. Giovanni Ziccardi** ha affrontato la tematica dell'IA ponendo l'accento sugli **aspetti etici e di responsabilità**, portando la discussione nel contesto del mondo accademico e della formazione. In qualità di docente dell'Università Statale di Milano coinvolto attivamente nella definizione di policy sull'intelligenza artificiale, Ziccardi ha condiviso l'esperienza concreta dell'elaborazione di un *decalogo* – una lista di principi guida – per governare l'uso dell'IA in Ateneo. Da questo punto di partenza, il suo intervento ha toccato sia la necessità di aggiornare i principi etici nell'era dell'IA, sia l'importanza di stabilire limiti e regole chiare per l'utilizzo responsabile di queste tecnologie.

Come premessa, Ziccardi ha osservato che l'avvento dell'IA sta determinando **un ritorno prepotente dell'etica** al centro del dibattito. Se altre rivoluzioni tecnologiche del passato (dalla diffusione di Internet ai social media) hanno dato agli operatori e ai regolatori un certo tempo per adattarsi, nel caso dell'IA generativa la situazione è diversa: siamo di fronte a un'evoluzione **rapidissima e quotidiana**, che non lascia il tempo di “digerire” le novità né di valutarne con calma l'impatto. Questa **accelerazione senza precedenti** impone uno sforzo di riflessione immediato su rischi, valori e responsabilità. In altre parole, mai come ora è urgente dotarsi di una bussola etica per navigare il cambiamento: non possiamo aspettare di capire l'IA ex post, dobbiamo provare a guidarla ex ante con principi chiari.

In quest'ottica si inserisce il **Decalogo per governare l'IA** elaborato presso l'Università Statale di Milano, che Ziccardi propone come modello di riferimento anche per altri contesti. Molti dei suoi principi, infatti, risultano sovrapponibili a quelli contenuti sia nella Carta dell'Ordine degli Avvocati (relativa all'uso di strumenti tecnologici nella professione) sia nelle policy adottate da varie aziende per gestire l'AI al proprio interno. Tra i punti chiave del decalogo milanese emergono in particolare:

- **Alfabetizzazione all'IA (AI literacy)** – Formare l'intera comunità (studenti, docenti, personale) a un uso **consapevole e competente** degli strumenti di intelligenza artificiale. L'IA non dev'essere una scatola nera magica: tutti gli operatori devono acquisire le basi per comprenderne il funzionamento e utilizzarla in modo produttivo nei rispettivi ambiti.
- **Protezione dei dati e privacy** – Mantenere un'attenzione rigorosa alla tutela dei dati personali e sensibili. In concreto, ciò si traduce nell'evitare di inserire informazioni personali, dati aziendali riservati o altri contenuti confidenziali in piattaforme di IA (specie quelle pubblicamente accessibili), poiché tali dati potrebbero essere *riutilizzati* dal fornitore per addestrare ulteriormente i modelli o comunque finire fuori controllo.
- **Trasparenza e tracciabilità** – Dichiarare sempre **se e come** è stata utilizzata l'intelligenza artificiale in un certo lavoro o processo decisionale, e mantenere traccia delle fonti/dataset utilizzati dal modello. Ad esempio, in ambito

accademico, uno studente dovrebbe dichiarare se ha usato ChatGPT per generare una parte dell'elaborato, e un docente dovrebbe indicare se alcune dispense sono state elaborate con l'ausilio di strumenti AI. La trasparenza tutela sia chi utilizza l'IA (che evita accuse di plagio o scorrettezza) sia i destinatari del risultato (che possono valutarne meglio l'attendibilità).

- **Qualità dei dati** – Curare la qualità dell'input fornito ai sistemi di IA, nella consapevolezza che **dati di scarsa qualità producono output di scarsa qualità** (*"garbage in, garbage out"*). Questo principio invita a non dare per scontata l'attendibilità dei risultati dell'IA: se la base di conoscenza o le istruzioni date al modello sono errate, incomplete o distorte, anche la risposta lo sarà. In ambito legale, significa verificare le fonti normative/giurisprudenziali alla base di un parere generato dall'IA; in ambito universitario, significa addestrare eventuali modelli interni con dati controllati e aggiornati.
- **Sostenibilità ambientale** – Considerare il costo energetico e ambientale dell'IA. Addestrare e far funzionare modelli di machine learning avanzati richiede una quantità enorme di risorse computazionali e quindi di energia: è importante usarli in modo **sostenibile**, ad esempio prediligendo soluzioni meno dispendiose quando possibile, evitando richieste futili ai modelli e supportando la ricerca su algoritmi più efficienti. L'etica dell'IA include anche la responsabilità verso **l'impatto ecologico** della tecnologia.
- **Cybersecurity** – Proteggere i sistemi di IA (e i dati che elaborano) da attacchi informatici e tentativi di **"avvelenamento" dei dati**. Un modello di IA è tanto affidabile quanto sicuro è il suo ciclo di vita: occorre prevenire intrusioni che possano manipolare gli algoritmi o i dataset di training, compromettere la privacy degli utenti o sfruttare l'IA per veicolare attacchi (si pensi al rischio di output dolosamente manipolati). La sicurezza deve quindi essere parte integrante della progettazione e gestione di questi strumenti.
- **Supervisione umana** – Mantenere sempre un **controllo umano** sull'operato dell'IA, specialmente in contesti critici. L'essere umano deve restare il *"guardiano"* ultimo che verifica la qualità e la correttezza del risultato prodotto dalla macchina, intervenendo prima che errori o distorsioni possano causare danni. Questo principio riflette l'idea che l'IA debba aumentare, non sostituire, il giudizio umano: ad esempio, in uno studio legale, un avvocato dovrà sempre riesaminare e convalidare un contratto generato dall'IA prima di consegnarlo al cliente; in università, un docente deve vigilare sull'uso che gli studenti fanno delle AI nel preparare i compiti.
- **Responsabilità** – Chi utilizza uno strumento di IA è **responsabile** del suo output. Non ci si può nascondere dietro *"ha detto così il computer"*: se un professionista adotta un sistema AI per un parere o un'analisi, rimane pienamente responsabile di eventuali errori, omissioni o bias presenti nel risultato. Questo vale sia in sede accademica (ad esempio uno studente non eviterebbe la bocciatura dicendo che è stato ChatGPT a scrivere una risposta errata) sia – a maggior ragione – in sede professionale, dove l'avvocato risponde verso il cliente della consulenza prestata, a prescindere dal fatto che abbia usato o meno l'IA per produrla.
- **Governance strategica dell'IA** – Trattare l'implementazione dell'IA come un tema di **governo** e strategia, non relegarlo alle sole scelte tecniche dell'area IT. In un'organizzazione (che sia un'università, uno studio legale, un'azienda), le

decisioni su quali sistemi di AI adottare, per quali scopi e con quali politiche d'uso devono coinvolgere il livello dirigenziale e riflettere la visione etica e di missione dell'ente. Delegare tutto ai tecnici informatici sarebbe miope: l'uso dell'IA tocca aspetti deontologici, reputazionali e operativi che richiedono l'attenzione dei vertici e l'elaborazione di linee guida condivise.

La riflessione di Ziccardi, quindi, elenca una serie di principi che dovrebbero guidare chiunque – avvocato, docente o studente – nell'integrazione dell'IA nella propria attività. Ma l'analisi non si ferma a indicare *cosa fare*: un punto particolarmente significativo sollevato dal professore riguarda anche **cosa NON fare**. Ziccardi evidenzia la necessità di capire **quando è meglio non usare l'IA**. In qualsiasi ambito in cui l'elemento distintivo della prestazione risiede nell'intelletto, nell'esperienza, nell'empatia o nella creatività **tipicamente umana**, l'apporto dell'IA potrebbe rivelarsi non solo inutile ma persino dannoso, livellando verso il basso la qualità del risultato. Pensiamo all'argomentazione giuridica di alto livello, alla strategia processuale, alla negoziazione delicata di un accordo: sono frangenti in cui il valore aggiunto risiede nel pensiero critico, nella sensibilità e nell'originalità dell'operatore umano, doti che nessuna macchina può eguagliare. Avere la consapevolezza di questi **limiti intrinseci dell'IA** è parte integrante di un approccio etico e maturo: significa saper scegliere in quali casi affidarsi alla tecnologia e in quali invece fare un passo indietro e privilegiare il fattore umano.

In sintesi, l'intervento di Giovanni Ziccardi ha portato al centro della discussione il tema della **responsabilità condivisa** nell'uso dell'IA. Ha richiamato tanto gli accademici quanto gli avvocati a dotarsi di **principi etici chiari**, di politiche di utilizzo trasparenti e della capacità di riconoscere i confini oltre i quali la macchina deve lasciare il posto all'uomo. Solo così – suggerisce Ziccardi – l'adozione dell'IA potrà avvenire senza tradire i valori fondamentali della formazione e della professione legale.

5. L'intelligenza artificiale nel mondo accademico - Prof.ssa Claudia Sandei

A conclusione della giornata, la **Prof.ssa Claudia Sandei** ha analizzato come l'IA stia trasformando il mondo accademico e, di riflesso, la formazione dei futuri professionisti del diritto. In particolare, il suo intervento si è concentrato sull'impatto dell'intelligenza artificiale rispetto alle **tre missioni fondamentali dell'università**: la didattica, la ricerca e la cosiddetta *terza missione* (ossia l'insieme di attività con cui l'accademia interagisce con la società e il sistema produttivo, dal trasferimento tecnologico alla divulgazione).

Didattica. Sul fronte dell'**insegnamento**, Sandei evidenzia che l'innovazione tecnologica sta già modificando sia i contenuti sia i metodi della formazione giuridica. Da un lato stanno nascendo **nuovi percorsi di laurea** in cui diritto e informatica si fondono: un esempio emblematico è l'istituzione di una laurea triennale in *Diritto e Tecnologia*, pensata per formare **professionisti ibridi** con competenze trasversali (giuridiche e digitali) capaci di operare nell'era dell'AI. Queste iniziative rispondono alla necessità, sempre più avvertita, di figure che sappiano dialogare tanto con il codice legislativo

quanto con il codice informatico. Dall'altro lato, anche i **metodi didattici** tradizionali vengono messi alla prova: l'IA stessa può diventare strumento e al tempo stesso *oggetto* di insegnamento. La professoressa Sandei ha raccontato un esperimento condotto in aula con i suoi studenti: ha fornito loro un testo interamente generato da un'IA e ha chiesto di analizzarlo criticamente, individuando eventuali errori, inesattezze o assunzioni infondate. L'esercizio, a suo dire, si è rivelato molto istruttivo perché ha evidenziato un punto cruciale: **senza una solida conoscenza di base**, gli studenti faticano a riconoscere gli sbagli o le fallacie nei contenuti prodotti dall'IA. Alcuni allievi, di fronte a affermazioni apparentemente plausibili ma sostanzialmente errate, tendevano a dar loro credito perché provenienti da una "macchina". Questo dimostra che, paradossalmente, l'avvento dell'IA **rende ancor più centrale il ruolo del docente** nel fornire strumenti critici e conoscenze solide: solo con solide fondamenta gli studenti potranno utilizzare l'IA come supporto senza esserne sopraffatti. L'IA, insomma, può essere un utile **ausilio didattico** (si pensi alla possibilità di avere tutor virtuali, quiz automatici, riassunti di testi, ecc.), ma non sostituisce la figura del docente; piuttosto ne ridefinisce il ruolo, che diventa sempre più quello di guida nella navigazione tra informazioni generate automaticamente.

Un altro aspetto toccato da Sandei riguarda la definizione di **regole chiare con gli studenti sull'uso dell'IA**. Molti atenei – inclusa la sua università – hanno iniziato a inserire nei regolamenti didattici e nei **syllabus** dei corsi indicazioni esplicite su quali utilizzi degli strumenti di IA sono consentiti e quali no durante il percorso formativo. Ad esempio, potrebbe essere permesso usare ChatGPT per fare brainstorming o per migliorare la forma di un testo scritto, ma vietato inserirvi parti di elaborati d'esame o tesi di laurea. Stabilire questi *patti chiari* fin dall'inizio serve sia a **responsabilizzare** gli studenti (che sanno di dover dichiarare e limitare l'uso di AI nei loro lavori, evitando così forme di plagio o scorciatoie indebite), sia a **tranquillizzare** i docenti, che vedono riconosciuta la legittimità di porre limiti e controlli. La trasparenza sugli usi ammessi dell'IA in ambito accademico diventa quindi parte integrante di un'etica condivisa tra docente e discente, volta a promuovere un utilizzo dell'IA che sia *supporto* e non *sotterfugio*.

Ricerca. Passando all'ambito della **ricerca scientifica**, la professoressa Sandei individua due tendenze principali che l'avvento dell'IA sta accentuando nel panorama delle pubblicazioni giuridiche. Da una parte, nota una certa **frammentarietà** e specializzazione estrema: l'AI è un tema così caldo e in così rapida evoluzione che molti studiosi inseguono l'ultima novità, producendo una miriade di note a sentenza, commenti a nuovi regolamenti, analisi di casi particolari. Questa iper-produzione "mordi e fuggi" rischia però di far perdere di vista la **visione sistematica**: si commenta la singola norma dell'AI Act o la singola decisione su un algoritmo discriminatorio, ma si fatica a collocare questi tasselli in un quadro teorico più ampio, a interrogarsi su come l'IA stia ridefinendo concetti cardine del diritto (responsabilità, persona, decisione, ecc.). In parallelo a questa frammentazione, Sandei evidenzia un fenomeno opposto: la **perdita di settorialità** delle discipline giuridiche di fronte all'IA. Tradizionalmente, l'accademia è

organizzata in settori (diritto privato, diritto penale, diritto amministrativo, ecc.), ognuno col proprio oggetto e i propri confini. L'IA invece è per natura un tema **trasversale**: ad esempio, il *diritto dei dati* o il *diritto dell'intelligenza artificiale* comprendono aspetti di privacy (tradizionalmente diritto privato), aspetti di trasparenza algoritmica e procedimenti automatizzati (che toccano il diritto amministrativo), aspetti di tutela contro abusi o reati informatici (ambito penalistico), questioni di antitrust e concorrenza (diritto commerciale), e così via. Questa trasversalità mette in crisi le vecchie categorie e **spinge verso approcci interdisciplinari**. Giuristi appartenenti a diversi settori sono chiamati a collaborare, nascono convegni e gruppi di ricerca intersettoriali sull'IA, e persino le riviste scientifiche iniziano a dedicare numeri speciali al tema coinvolgendo autori di estrazione diversa. Secondo Sandei, questo *shock* epistemologico può rivelarsi salutare: obbliga la comunità scientifica a uscire dai propri recinti disciplinari e a elaborare **nuovi paradigmi condivisi**, adeguati a fenomeni (come l'AI) che non rispettano certo i confini accademici tradizionali.

Terza missione. Infine, Sandei ha esaminato l'impatto dell'IA sulla **terza missione** dell'università, cioè il complesso delle attività con cui gli atenei contribuiscono all'innovazione e al benessere della società. L'intelligenza artificiale, essendo una tecnologia di frontiera con forti implicazioni economiche e sociali, offre all'accademia l'opportunità di **rafforzare il proprio ruolo di motore dello sviluppo** sul territorio. Un esempio concreto citato è quello dei **Centri di Competenza** finanziati negli ultimi anni anche tramite il *Piano Nazionale di Ripresa e Resilienza (PNRR)*, nati per favorire il trasferimento tecnologico verso le imprese. In particolare, nel Nord-Est d'Italia opera il centro di competenza **SMACT**², che funge da nodo di collegamento tra università e aziende sui temi della digitalizzazione (il nome richiama infatti i termini Social, Mobile, Analytics, Cloud, IoT). Questi centri raccolgono i bisogni di innovazione provenienti dal tessuto imprenditoriale e formano **team multidisciplinari** di esperti universitari per sviluppare soluzioni applicative ad hoc – spesso vere e proprie soluzioni “chiavi in mano” – da trasferire alle aziende. Ebbene, molte delle richieste che giungono oggi riguardano l'adozione di sistemi di **intelligenza artificiale nei processi aziendali** (anche nel settore legale aziendale, basti pensare alla gestione dei contratti o alla compliance automatizzata). Attraverso strutture come SMACT, gli atenei contribuiscono attivamente a realizzare progetti di AI sul campo: ad esempio implementando algoritmi di analisi dati per un ente pubblico, oppure aiutando una startup legale a sviluppare un sistema di *legal analytics*. Questo modello collaborativo consente di colmare il divario tra teoria e pratica, facendo sì che le conoscenze sviluppate in ambito accademico trovino uno sbocco concreto e producano impatto reale. Sandei sottolinea come la partecipazione a tali progetti sia anche formativa per gli stessi accademici e studenti coinvolti, perché offre un ritorno di esperienza dal mondo reale verso l'università. In definitiva, l'IA sta **potenziando la “terza missione”**: l'università non solo studia l'IA in astratto, ma la applica e la *governa* in partnership con imprese e istituzioni, contribuendo a diffondere una cultura dell'innovazione responsabile nel tessuto sociale.

Prima di concludere, la professoressa Sandei ha voluto evidenziare un filo conduttore che unisce le sue osservazioni: l'adozione dell'IA in ambito giuridico (sia professionale che accademico) non pone solo sfide tecniche, ma **sfide culturali**. Richiede apertura al cambiamento ma anche spirito critico, richiede investimenti in nuove competenze ma anche capacità di mettere in discussione vecchi schemi. L'università, da questo punto di vista, ha un duplice ruolo cruciale: formare **giuristi nuovi** (in grado di capire e usare l'IA) e al contempo fungere da **coscienza critica** del cambiamento, interrogandosi sugli effetti di lungo periodo e orientando il dibattito etico e deontologico. Solo così – concludono idealmente tutti i relatori – l'intelligenza artificiale potrà essere integrata nel mondo del diritto senza snaturarlo, ma anzi esaltando ciò che resta insostituibilmente umano nella nostra idea di giustizia.

Capitolo IV

Soluzioni AI per Avvocati nel panorama Legal Tech Italiano

1. Introduzione

La quarta lezione del *Corso di Perfezionamento in Intelligenza Artificiale per Avvocati* (28 maggio 2025) ha offerto una panoramica pratica e diversificata del panorama **Legal Tech** italiano. In particolare, l'incontro – moderato dall'Avv. **Giuseppe Vaciago** – ha presentato una carrellata di soluzioni software che integrano l'**Intelligenza Artificiale** nella professione legale. Sin dalle premesse, è stato chiarito che la sessione non aveva un fine commerciale, bensì formativo: l'obiettivo era fornire ai partecipanti una comprensione concreta delle diverse applicazioni dell'IA in ambito legale. Si è dunque spaziato dalla ricerca giuridica potenziata dall'AI all'analisi documentale automatizzata, dalla redazione assistita di atti legali alla digital forensics, fino a soluzioni per la sicurezza delle comunicazioni e la gestione dei processi aziendali. Per ciascuna soluzione presentata sono stati evidenziati la **filosofia** di fondo, le **funzionalità principali**, le **implicazioni pratiche** per l'avvocato e, quando emersi, gli aspetti metodologici o strategici. Ne è emerso un quadro eterogeneo dell'offerta Legal Tech, utile per orientarsi in un mercato in rapida evoluzione e iniziare a valutare gli strumenti più adatti alle proprie esigenze professionali.

Di seguito, dopo una breve introduzione contestuale, verranno ricostruiti in forma narrativa e professionale i principali contenuti degli interventi dei vari relatori e delle aziende protagoniste della lezione: **Keplera**, (presentato da Marco Baroni), **Tiresia**, (presentato da Marta Elmi Camossi) e **Gaya-X** (presentato da Egidio Casati). Ciascun capitolo ripercorre fedelmente quanto emerso da ogni intervento, ponendo l'accento sulla filosofia dello strumento, sulle sue funzionalità e sulle ricadute pratiche per l'avvocato.

2. Keplera – La piattaforma “all-in-one” per la gestione documentale legale – Dott. Marco Barone

Il primo intervento ha visto protagonista **Keplera**, una giovane startup italiana che unisce competenze legali e tecnologie d'avanguardia. Nata inizialmente come *Alternative Legal Service Provider* tradizionale, Keplera ha poi compiuto una trasformazione in ottica *Legal Tech*, sviluppando un software proprietario equipaggiato con funzionalità di **IA generativa** per digitalizzare e rendere più efficiente la gestione dei documenti legali. Questo prodotto – presentato durante la lezione come una piattaforma “all-in-one” (in passato indicato anche col nome *LexHero*) – incarna la filosofia di Keplera: utilizzare l'IA per **accelerare i processi legali interni**, senza però rimpiazzare il professionista, bensì **colmandone le lacune** e automatizzando le attività ripetitive.

Durante la presentazione, il rappresentante di Keplera ha spiegato come la piattaforma sia originariamente nata per il **Contract Lifecycle Management (CLM)**, ossia per gestire l'intero ciclo di vita dei contratti, dalla redazione alla conservazione. In seguito, sono state integrate funzionalità di IA in modo pervasivo, con l'idea di "tappare le crepe" lasciate dai sistemi documentali tradizionali. In pratica, la piattaforma sfrutta modelli di linguaggio avanzati per assistere l'avvocato in ogni fase della gestione documentale.

Le **funzionalità principali** illustrate comprendono:

- **Editor modulare:** un editor di testi proprietario che scompone il documento giuridico in moduli flessibili (ad esempio per clausole, articoli, commi). Questo approccio modulare consente di riorganizzare e modificare con agilità le varie parti di un contratto, caratteristica molto apprezzata nella redazione di documenti complessi. L'avvocato, ad esempio, può riordinare clausole o inserirne di nuove senza compromettere la struttura complessiva del testo, grazie a questa gestione granulare.
- **Generazione automatica con IA:** la piattaforma integra funzionalità di **generative AI** sia per la creazione completa di documenti sia per la redazione di singole clausole. Ad esempio, è possibile fornire un *prompt* descrittivo e ottenere la bozza di un contratto completo e strutturato, con tutte le sezioni al posto giusto. In alternativa, mentre si lavora in editor, l'utente può chiedere all'IA di generare automaticamente una clausola specifica (ad esempio una clausola di riservatezza, o una clausola penale), che viene inserita nel punto desiderato del documento. Queste funzioni generative aiutano a velocizzare sensibilmente la stesura di contratti standard, lasciando però sempre al legale il controllo finale e la possibilità di modificare il testo generato.
- **Analisi intelligente dei documenti (*Document Analyzer*):** un altro tassello importante mostrato è lo strumento di analisi automatica dei contratti. Caricando un contratto già esistente sulla piattaforma, l'IA effettua il parsing del testo individuando le varie clausole in esso contenute, ne valuta il contenuto e assegna ad ognuna un **punteggio di rischio** (alto, medio o basso). Ad esempio, potrebbe segnalare come "alto rischio" una clausola di non concorrenza formulata in modo troppo vago, spiegandone i motivi (ad es. potenziale inapplicabilità o scarsa tutela del cliente) e proponendo un testo alternativo migliorativo. Questa funzionalità di *contract review* automatizzata risulta estremamente utile nella **due diligence contrattuale**: l'avvocato può rapidamente identificare le clausole critiche di un accordo e disporre di suggerimenti su come correggerle o rafforzarle.
- **Collaborazione in tempo reale:** infine, è stata evidenziata la capacità della piattaforma di supportare il lavoro collaborativo multi-utente sullo stesso documento, in tempo reale, in modo analogo a quanto avviene su Google Docs. Più professionisti possono quindi intervenire simultaneamente sulla bozza di un contratto, vedendo in diretta le modifiche reciproche. Questa caratteristica è pensata per gli studi legali strutturati, dove più colleghi lavorano insieme su atti complessi, o per gestire in modo interattivo le negoziazioni contrattuali con controparti esterne, mantenendo comunque traccia di tutte le revisioni in un ambiente sicuro.

Implicazioni pratiche: la soluzione di Keplera mira a **snellire il lavoro contrattuale quotidiano**. La filosofia “IA come assistente” è tangibile: l’avvocato rimane al centro del processo decisionale, ma può delegare all’algoritmo i compiti di drafting ripetitivo o di analisi preliminare. Ciò comporta un significativo risparmio di tempo nella redazione e revisione dei contratti, riducendo anche il margine di errore umano (ad esempio dimenticare clausole importanti o non accorgersi di un passaggio rischioso). Inoltre, grazie all’uso di fonti attendibili e modelli controllati, la piattaforma garantisce che i testi generati o suggeriti rispettino gli standard legali richiesti – un aspetto fondamentale per guadagnare la fiducia dei professionisti forensi più scettici verso l’AI. In sintesi, l’intervento su Keplera ha mostrato una tecnologia matura e **integrata a 360°** nel flusso di lavoro legale, capace di aumentare l’efficienza senza sacrificare il rigore e la sicurezza dei documenti trattati.

3. Tiresia – L’IA “super-collaboratore” per il Diritto Tributario – Dott.ssa Marta Elmi Camossi e Avv. Sebastiano Stufano

Il secondo intervento narrato è quello dedicato a **Tiresia Tax Consulting Copilot**, presentato dalla Dott.ssa **Marta Elmi Camossi** e dall’Avv. **Sebastiano Stufano** (tributarista). Tiresia è stato descritto come uno strumento di **Intelligenza Artificiale altamente specializzato nel settore fiscale**, un “Tax Consulting Copilot” progettato per affiancare l’avvocato tributarista nell’elaborazione di pareri e strategie. La filosofia alla base di Tiresia è chiara: non fornire risposte semplicistiche tipo “sì o no” a quesiti fiscali complessi, ma produrre **argomentazioni giuridiche strutturate**. In altre parole, il sistema agisce come un *super-collaboratore* dello studio legale, capace di analizzare normative e giurisprudenza pertinente e restituire un ragionamento articolato a supporto strategico del professionista.

Durante la lezione è emerso come Tiresia sia **verticale sul diritto tributario**, ambito caratterizzato da norme in continua evoluzione, giurisprudenza abbondante e prassi amministrative rilevanti (circolari, interpelli, ecc.). I relatori hanno sottolineato che, proprio perché la materia fiscale è complessa e “*fluida*”, un tool specializzato come questo può fare la differenza nel lavoro quotidiano del tributarista. Vediamo le sue funzionalità principali:

- **Prompt Assistant contestuale:** Nella chat principale, attraverso la funzione di **Redazione guidata del quesito**, il sistema presenta all’utente un percorso interattivo composto da alcune domande chiave (ad es. tipologia di soggetto, settore di attività, ambito territoriale, documentazione disponibile, inquadramento temporale, ecc.) progettate per raccogliere tutte le informazioni indispensabili a definire il perimetro del caso. In questo modo, si arricchisce il contesto del quesito in modo ordinato e strutturato, così da consentire a Tiresia di fornire la migliore risposta possibile, precisa e pertinente. Questo garantisce che ogni risposta sia personalizzata e aderente alla realtà concreta del cliente. Un’altra funzione, inoltre, è quella dell’**Assistente al Prompt** che permette o di avere un ulteriore ausilio nella strutturazione del quesito o di caricare e

analizzare documenti PDF (come accertamenti, sentenze, contratti, ecc.), estrapolando i punti chiave e aiutando il professionista a definire quesiti mirati. I dati contenuti nel documento caricato sono immediatamente anonimizzati, e lo stesso documento è trattenuto per il solo tempo necessario a generare la risposta al quesito posto. In poco tempo, Tiresia è in grado di **estrarre i riferimenti normativi e i temi controversi**, suggerendo quindi argomenti di difesa da approfondire. Quest'ultima funzione è stata dimostrata con un esempio pratico, mostrando come da un avviso di accertamento complesso Tiresia fosse in grado di estrarre i punti chiave controversi e proporre le possibili linee difensive da approfondire.

- **Risposte articolate in quattro parti:** ogni parere o risposta generata da Tiresia viene strutturato in maniera sistematica seguendo quattro pilastri, che riflettono il tipico **iter logico di un tributarista esperto**. Come spiegato dai relatori, quando si pone una domanda al sistema, la risposta di Tiresia comprenderà sempre:
 1. un **inquadramento normativo** della questione, nel quale viene definito e delineato il contesto entro cui il quesito richiesto si colloca rispetto alle leggi, ai regolamenti e alle altre disposizioni normative vigenti, arricchito con spunti su come tali norme debbano essere interpretate e applicate (es. il Testo Unico delle Imposte sui Redditi per una detrazione fiscale, o il DPR 600/1973 per le procedure di accertamento...);
 2. un **supporto giurisprudenziale** in cui sono elaborati e analizzati gli orientamenti interpretativi e citate sia le pronunce di legittimità (Corte di Cassazione) sia quelle di merito (Commissioni Tributarie Provinciali/Regionali) pertinenti al caso in esame;
 3. il richiamo alla **prassi amministrativa** rilevante, come Circolari, Risoluzioni e Risposte a Interpello dell'Agenzia delle Entrate;
 4. infine – ove opportuno – eventuali considerazioni conclusive o indicazioni operative.

Questa struttura quadrupla garantisce che nulla venga tralasciato: la **normativa** offre la base legale, la **giurisprudenza** fornisce l'interpretazione dei giudici (sia in astratto che in concreto), e la **prassi** chiarisce l'orientamento delle autorità fiscali. L'avvocato, ricevendo un output così organizzato, può rapidamente farsi un'idea completa del quadro e disporre di riferimenti precisi da approfondire o inserire in un ricorso tributario.

Oltretutto, dopo la risposta principale, si ha la possibilità di porre fino a 3 **domande di approfondimento** collegate alla prima, qualora permangano dubbi e/o si desideri indagare taluni profili specifici emersi dalla risposta principale.

Il parere generato può essere salvato in **PDF** e/o archiviato nel proprio **archivio personale**, completo della documentazione integrale rilevante citata nella risposta generata da Tiresia (fungendo da **Banca dati** sempre consultabile).

- **Gestione della complessità normativa e storica:** un aspetto su cui i relatori hanno posto l'accento è la capacità dello strumento di gestire la variabilità e l'evoluzione del diritto tributario. In sede di presentazione è stato spiegato che l'IA di Tiresia è in grado, ad esempio, di applicare in modo corretto norme diverse a seconda delle annualità coinvolte in una controversia. Se un quesito riguarda più anni d'imposta, il sistema sa che potrebbero essere cambiate le aliquote, le soglie o addirittura la norma applicabile, e tiene conto di queste differenze. Inoltre, Tiresia può ricostruire l'**evoluzione degli orientamenti giurisprudenziali** nel tempo: se su una certa detrazione fiscale la Cassazione aveva un orientamento fino al 2020 e poi l'ha modificato nel 2021, la risposta generata ne darà conto, segnalando il cambiamento (funzione di **Time machine**). Ciò è frutto di un continuo aggiornamento del dataset, che include **fonti ufficiali** come leggi, sentenze e prassi amministrative. In aggiunta, il sistema è in grado di affrontare anche **casi di fiscalità internazionale** complessi, integrando nelle risposte richiami a convenzioni contro le doppie imposizioni o normative UE se pertinenti.

Implicazioni per il professionista: Tiresia rappresenta per l'avvocato tributarista un **alleato prezioso**, in grado di ridurre drasticamente i tempi di ricerca e di elaborazione di un parere fiscale articolato, oltre a fornire un supporto strategico nella definizione delle proprie strategie difensive. Nella pratica quotidiana, ciò si traduce in risposte ai clienti più rapide e fondate. Ad esempio, invece di trascorrere ore a reperire manualmente sentenze sulle banche dati e a ricostruire la storia normativa di una certa agevolazione, l'avvocato può affidare questo compito preliminare all'IA, ottenendo in pochi minuti una bozza di ragionamento completa e un orientamento per le proprie strategie difensive. Ovviamente, il professionista dovrà poi verificare e affinare l'output – la “parola finale” resta umana – ma il grosso del lavoro di raccolta e *digestione* delle fonti viene automatizzato. I relatori hanno anche evidenziato un ulteriore vantaggio: **l'aggiornamento continuo** del sistema assicura che lo studio legale disponga sempre di informazioni allineate all'ultima giurisprudenza disponibile, cosa non scontata in un ambito dove settimanalmente escono nuove sentenze. In sintesi, Tiresia incarna la visione di un'IA **etica e di supporto**, che non sostituisce l'avvocato ma gli fornisce un Super-collaboratore (dotato di memoria sterminata, velocità di calcolo, capacità di correlare migliaia di documenti) addestrato per migliorare la qualità e l'efficacia del servizio al cliente.

4. Gayadeed AI Sphere – L'“AI Hub” sicuro per i dati sensibili – Dott. Egidio Casati

Un tema emerso durante la lezione è stato quello della **sicurezza dei dati e della sovranità digitale**, affrontato nell'intervento di **Egidio Casati**. Casati ha presentato la soluzione **Gayadeed AI Sphere**, pensata per aziende operanti in settori altamente regolamentati (come il fintech, la finanza decentralizzata, il settore Web3) e, per estensione, di grande interesse anche per gli studi legali che trattano dati sensibili. La domanda di fondo a cui questa soluzione cerca di rispondere è: *come sfruttare la potenza dei grandi modelli linguistici (LLM) garantendo al contempo che i dati aziendali riservati non escano mai da un ambiente privato e controllato?*

Contesto e problema: Casati ha esordito delineando la sfida che molte organizzazioni (imprese e studi professionali) affrontano oggi. Da un lato c'è l'attrattiva dei servizi di IA generativa offerti dai grandi provider (si pensi alle API di OpenAI, a Google Bard, ecc.), dall'altro c'è il timore – fondato – che utilizzare questi servizi comporti l'invio di informazioni riservate a server esterni, spesso extra-UE, con conseguenti rischi sul piano della privacy, del segreto professionale e della compliance normativa (GDPR, segreti industriali, ecc.). Nel settore legale, ad esempio, uno studio potrebbe voler interrogare un modello GPT su un corpus di contratti interni per trovare clausole ricorrenti o anomalie, ma non può certo caricare quei contratti su una piattaforma pubblica, pena violare la riservatezza dei clienti. **Come conciliare dunque l'uso di un LLM potente con la tutela assoluta dei dati?**

Soluzione – l'AI Hub Gayadeed AI Sphere: La risposta proposta da NymLab è la creazione di un **AI Hub privato** conforme dal punto di vista della sicurezza e della privacy e rispettoso della **sovranità dei dati**. In quest'ottica, l'AI Hub presentato funziona come un **gateway sicuro** tra l'organizzazione utente e i modelli di IA esterni. Più in dettaglio, il sistema prevede due componenti chiave:

- una **Knowledge Base privata** e
- un modulo di **interazione sicura con l'LLM esterno (no training no tracing no log)**.

Nel **cloud privato** (nel caso specifico basato su Oracle Cloud Infrastructure, configurato però in modo isolato e con credenziali d'accesso detenute unicamente dal cliente) l'azienda o lo studio legale carica tutti i propri documenti, contratti, database di conoscenza. Si tratta quindi di costruire internamente un patrimonio informativo strutturato – un po' come un archivio digitale dello studio – sul quale l'AI può effettuare **Retrieval-Augmented Generation (RAG)**, ossia recuperare spezzoni rilevanti per poi imbastire risposte più puntuali.

Quando un utente pone una domanda attraverso l'Hub, il sistema innanzitutto **arricchisce il prompt** attingendo dalla **knowledge base privata**: ad esempio, se un avvocato in azienda chiede *"Quali clausole di non concorrenza abbiamo utilizzato nei contratti con i distributori?"*, l'Hub effettua una ricerca nei contratti interni caricati e recupera alcuni estratti di clausole di non concorrenza presenti. Quegli estratti – opportunamente anonimizzati se contengono nomi o importi – vengono aggiunti al prompt originale sotto forma di contesto. **Solo a questo punto** il prompt arricchito viene inoltrato al modello di IA esterno (che può essere un servizio OpenAI utilizzato in modalità sicura, piuttosto che un modello open source hosted su un altro server, a seconda della scelta del cliente). Il LLM elabora la risposta, che potrebbe essere ad esempio un riassunto delle clausole tipiche trovate o un suggerimento di nuova clausola unificante i vari criteri. La risposta generata **torna all'AI Hub**, che eventualmente sostituisce nuovamente i token segnaposto con i dati originali (se erano stati anonimizzati) e la presenta all'utente.

Il risultato di questo flusso è che **nessun dato sensibile esce mai in chiaro dal perimetro privato**: il modello esterno “vede” solo il prompt già integrato, che contiene riferimenti alle informazioni interne in forma protetta o aggregata, ma non ha mai accesso diretto all’intero dataset del cliente. Casati ha paragonato l’AI Hub a una sorta di *cassetta di sicurezza con una fessura*: “da fuori il modello di AI può fornire risposte, ma non può sbirciare dentro la cassaforte dei dati”. Questa architettura garantisce che l’azienda o lo studio mantengano il completo controllo sulla propria conoscenza, conciliando quindi potenza dell’AI e sovranità del dato.

Flessibilità e compliance: un altro pregio del sistema Gayadeed AI sphere presentato è la sua **modularità**. L’Hub sicuro è stato pensato per potersi collegare a diversi servizi o modelli IA **intercambiabili**, comunque già esclusivamente disponibili in territorio EU e quindi effettivamente soggetti alla normativa vigente (gdpr, dora, etc). Questa portabilità, ha sottolineato Casati, è fondamentale in un’epoca in cui la regolamentazione dell’AI (si pensi al futuro **AI Act** europeo) potrebbe porre requisiti stringenti sull’uso di certi modelli o sull’esportazione di dati. Con l’AI Hub Gayadeed AI sphere, un dipartimento legale può già ora iniziare a usare l’AI in modo controllato, sapendo di poter poi “**agganciare**” l’Hub al servizio più conforme o conveniente mano a mano che il mercato e le norme evolvono.

Implicazioni per gli avvocati: l’entusiasmo attorno a questa soluzione deriva dal fatto che molti avvocati (specialmente d’impresa) sono frenati nell’utilizzare strumenti come ChatGPT proprio da considerazioni deontologiche e di riservatezza. L’AI Hub di NymLab Gayadeed AI Sphere offre una **risposta concreta a tali preoccupazioni**: consente di avere il meglio dei due mondi, ovvero sfruttare capacità avanzate di analisi del linguaggio naturale su un patrimonio informativo vastissimo, ma **senza dover mettere a repentaglio il segreto professionale** o la privacy dei clienti. Casati ha concluso il suo intervento evidenziando che, in prospettiva, soluzioni del genere potrebbero essere adottate anche dagli Ordini professionali o da reti di studi legali, creando dei veri e propri **circuiti chiusi di conoscenza legale condivisa** dove l’IA può operare in totale sicurezza. Un simile scenario aprirebbe la porta a una collaborazione tra studi basata su AI federate, mantenendo però ogni dato sotto il controllo di chi lo detiene. È una strada ancora in divenire, ma che ha suscitato molto interesse tra i partecipanti alla lezione, consapevoli che **la sicurezza non è un optional** quando si parla di innovazione digitale in campo legale.

5. Lexroom AI – Dott. Mauro Borselli

L’idea alla base di **Lexroom AI** è di fornire uno strumento di Intelligenza Artificiale generativa che **potenzi l’efficacia del professionista legale**, senza in alcun modo sostituirlo. A differenza dei chatbot generalisti - come ad esempio ChatGPT, Lexroom è progettato **specificamente per il diritto italiano**, basandosi su un modello addestrato con **fonti ufficiali e istituzionali locali**, per garantire risposte **affidabili, pertinenti e verificate**.

Questo approccio verticale **consente di velocizzare il lavoro dell'avvocato** - si stimano **oltre 5 ore risparmiate a settimana** tra ricerche e stesura di documenti - e di **minimizzare errori e imprecisioni**, grazie a un aggiornamento costante delle fonti e alla loro **citazione puntuale** in ogni risposta.

L'obiettivo non è fornire "risposte magiche" prive di fondamento, ma **affiancare il professionista legale** con uno strumento affidabile e veloce, in grado di generare **contenuti giuridici supportati da riferimenti normativi e giurisprudenziali verificabili**.

Tre funzionalità chiave per l'intero flusso di lavoro legale

1. Ricerca giuridica

L'utente può formulare un quesito in linguaggio naturale e ottenere una **risposta argomentata**, completa di **riferimenti normativi e giurisprudenziali pertinenti**. Lexroom AI effettua una **ricerca** su milioni di fonti ufficiali e genera **in pochi istanti una bozza di parere motivata**. Le fonti citate (es. articoli di legge, sentenze, provvedimenti di autorità di settore) sono **immediatamente consultabili e scaricabili**, in modo che il professionista possa verificarle e utilizzarle nel proprio lavoro.

Questo consente di **sostituire ore di ricerca tradizionale** su banche dati e codici con **una risposta esaustiva ottenuta in pochi secondi**.

2. Analisi documentale

Lexroom AI è in grado di **analizzare documenti caricati dall'utente** – come contratti, sentenze, atti – ed **estrarne le informazioni chiave** o produrre **una sintesi ragionata**.

Ad esempio, caricando una sentenza complessa o un regolamento articolato, la piattaforma è in grado di **fornire un riassunto dei punti salienti**.

Inoltre, può **confrontare più documenti tra loro**, evidenziandone **differenze, similitudini e correlazioni**.

Lexroom supporta documenti **in più lingue** e svolge queste operazioni di analisi, sintesi e confronto **in modo estremamente rapido**, permettendo all'avvocato di **focalizzarsi subito sugli aspetti rilevanti**.

In questo modo, agisce come **assistente di ricerca**, alleggerendo il professionista dall'onere di dover **leggere e comparare manualmente grandi quantità di testo**.

3. Redazione di atti

Una delle funzionalità più avanzate è la **generazione automatica di bozze di atti legali**. Partendo da **modelli predefiniti**, Lexroom può compilare un atto – ad esempio un **decreto ingiuntivo** – inserendo automaticamente nei campi pertinenti i **dati estratti dai documenti forniti**, come visure camerali, fatture o contratti.

In pochi secondi, il professionista ottiene **una bozza praticamente pronta**, che dovrà solo **rivedere e rifinire**.

Lexroom garantisce atti **formalmente corretti**, riducendo il rischio di errori di distrazione e **accelerando in modo significativo la fase di redazione**.

L'avvocato mantiene sempre **il controllo finale** sul contenuto, ma parte da una **base solida e coerente** generata dall'IA, risparmiando **tempo prezioso** sulle attività ripetitive di copia-incolla e formattazione.

Sicurezza e tutela della privacy: una priorità assoluta

Considerata la natura altamente sensibile dei dati trattati, **Lexroom AI adotta i più elevati standard di sicurezza e tutela della privacy**.

I documenti e i dati caricati non vengono **mai memorizzati in modo permanente** sui server: sono utilizzati **esclusivamente per il tempo necessario alla sessione** di ricerca o elaborazione, dopodiché vengono **eliminati**. Inoltre, Lexroom **non utilizza i dati degli utenti per addestrare i propri modelli di intelligenza artificiale**, prevenendo qualsiasi rischio di **diffusione non autorizzata di informazioni riservate**.

L'infrastruttura tecnica è **interamente ospitata nell'Unione Europea**, presso provider conformi agli **standard internazionali di sicurezza e compliance** (es. certificazioni ISO). Tutto ciò garantisce che l'utilizzo della piattaforma sia **pienamente conforme al GDPR**, permettendo allo studio legale di **mantenere il pieno controllo sui propri documenti**.

In sintesi, **Lexroom AI è un ambiente sicuro, affidabile e privacy-friendly**, pensato per mettere l'avvocato nella condizione di **lavorare meglio, più velocemente e con maggiore serenità**, sapendo che le informazioni dei propri clienti restano **sempre confidenziali**.

6. Kopjra– Dott. Tommaso Grotto

La Digital Forensics è l'insieme di tecniche e procedure volte a **raccogliere, conservare e analizzare** tracce digitali affinché possano valere come prove in un procedimento legale. Negli ultimi anni questa disciplina ha acquisito un ruolo sempre più cruciale, soprattutto a causa dell'avvento dei **deepfake** e, più in generale, delle sofisticate manipolazioni digitali di foto, video e documenti. Oggi è relativamente facile alterare un'immagine o generare un video/falso audio realistico tramite intelligenza artificiale, con il rischio di compromettere l'affidabilità delle prove elettroniche. Episodi recenti hanno mostrato truffe milionarie realizzate con audio fake o casi di disinformazione tramite video contraffatti. Di fronte a questa realtà, diventa fondamentale per investigatori e legali poter **verificare l'autenticità** e l'integrità di qualunque elemento digitale presentato in giudizio. La digital forensics fornisce dunque gli strumenti per certificare *chi, quando e come* ha generato o modificato un dato file, assicurando che la prova informatica non sia stata alterata e possa essere accettata dal giudice. In sintesi,

in un'era di *fake digitali*, l'informatica forense è il **baluardo** che assicura la **fiducia** nelle evidenze elettroniche.

L'Intelligenza Artificiale sta **rivoluzionando** molte fasi del processo forense, rendendo le analisi più rapide ed efficaci. In particolare, i contributi maggiori si vedono in due ambiti:

1. **Analisi dei dati:** le indagini digitali spesso comportano l'esame di quantità enormi di informazioni (si pensi alle e-mail di un'azienda, ai log di sistema, o a migliaia di immagini e video sequestrati). Strumenti di AI avanzata aiutano gli analisti a **setacciare automaticamente** questi dataset mastodontici, identificando pattern, correlazioni e indizi potenzialmente rilevanti che sfuggirebbero ad una ricerca manuale. Ad esempio, algoritmi di *computer vision* possono scorrere una collezione di fotografie alla ricerca di oggetti o volti specifici (es. individuare in mezzo a migliaia di foto quelle che mostrano un certo logo o una certa arma); allo stesso modo, motori OCR (riconoscimento ottico dei caratteri) potenziati da AI estraggono testo da documenti scansionati o da screenshot di chat, rendendolo ricercabile. Queste automazioni permettono di **classificare e filtrare** rapidamente i contenuti multimediali (ad es. isolando immagini contenenti materiale illecito), lasciando ai periti umani il compito di approfondire soltanto i dati veramente significativi. L'AI viene impiegata anche nell'e-discovery legale, dove deve passare al vaglio milioni di documenti alla ricerca di parole chiave o comunicazioni sospette: algoritmi di *text mining* e *natural language processing* possono identificare discussioni ricorrenti, collegare mittenti e destinatari, evidenziare anomalie linguistiche, accelerando indagini che manualmente richiederebbero mesi.
2. **Reportistica:** una volta raccolti ed analizzati i risultati tecnici, il perito forense deve stendere una **relazione** dettagliata da presentare in tribunale, spiegando metodi e conclusioni in modo chiaro. Anche qui l'IA offre un supporto prezioso: sistemi di **generazione del linguaggio naturale** (NLG) possono aiutare a redigere report strutturati, pre-compilando ad esempio le sezioni ripetitive (descrizione degli strumenti usati, parametri tecnici, ecc.) oppure fornendo bozze di sintesi dei risultati che l'esperto potrà rifinire. Inoltre, modelli linguistici avanzati possono tradurre automaticamente documenti tecnici o conversazioni (ad es. messaggi estratti in lingua straniera) per includerli nella relazione. Pur dovendo sempre essere verificato e validato dal professionista, l'AI velocizza la preparazione della **documentazione finale** e riduce il rischio di dimenticanze, generando schemi uniformi e completi. In sostanza, l'IA assiste il forense non solo nel *trovare* le evidenze, ma anche nel *presentarle* al meglio, rendendo più efficiente l'intero iter che va dall'analisi alla testimonianza in aula.

Per rispondere a queste esigenze, Kopjra ha sviluppato una piattaforma chiamata **Web Forensics**. Si tratta di un software fornito in modalità **SaaS** (Software-as-a-Service) che consente di effettuare **acquisizioni forensi di contenuti online** – ad esempio pagine web, profili social, conversazioni chat o email – preservandone il valore probatorio. La piattaforma crea infatti una *copia forense* del contenuto internet di interesse, operando in conformità con lo **standard ISO/IEC 27037:2012** (linee guida internazionali per l'identificazione, raccolta, acquisizione e conservazione delle evidenze digitali). In

pratica, Web Forensics permette a un operatore di navigare la risorsa web target all'interno di un ambiente controllato e monitorato: ogni azione svolta (click, scroll, inserimento credenziali, ecc.) viene registrata e ogni elemento visualizzato viene acquisito in forma immutabile. Al termine del processo, il software genera un **pacchetto di evidenza forense** completo di tutto il materiale raccolto (incluse la registrazione video dell'intera navigazione effettuata, gli screenshot delle pagine, i file scaricati, il traffico di rete catturato, le chiavi crittografiche utilizzate e vari log tecnici). Assieme a esso, viene prodotta una **relazione tecnica forense**: un documento descrittivo che dettaglia la metodologia seguita per l'acquisizione (con riferimenti temporali, strumenti utilizzati, hash calcolati sui file, ecc.), così da rendere **verificabile** e ripetibile il processo secondo i canoni richiesti dalla scienza forense. Questo doppio output – pacchetto di dati + relazione – costituisce la **prova digitale** pronta per essere depositata in giudizio. Grazie a Web Forensics, operazioni complesse come acquisire una pagina web dinamica (che potrebbe cambiare o sparire col tempo) diventano attività eseguibili anche da remoto, in modo **standardizzato e certificabile**, senza bisogno di sequestrare fisicamente un computer. Ciò risulta estremamente utile in casi di indagini su reati online (diffamazione a mezzo social, phishing, e-commerce illegali, ecc.), garantendo che il contenuto incriminato venga congelato in una forma idonea per la presentazione davanti a un giudice.

Un aspetto distintivo della soluzione di Kopjra è l'uso di algoritmi di intelligenza artificiale **all'interno** del processo di acquisizione forense online, attraverso quelli che l'azienda definisce *“adattatori IA”*. Si tratta di una tecnologia **brevettata** da Kopjra che opera dietro le quinte di Web Forensics: in sostanza, l'AI analizza **in tempo reale** la pagina web (o il servizio online) che si sta per acquisire, per comprenderne la struttura e adattare di conseguenza la procedura di copia forense. Ad esempio, l'**adattatore AI** è capace di riconoscere se la pagina richiede un login o qualche forma di autenticazione e, in tal caso, guida l'utente nelle operazioni da compiere (o eventualmente adatta il browser forense per gestire la sessione di login senza interrompere la validità della prova). Oppure, nel caso di applicazioni web complesse – come le **single-page application** in cui i contenuti si caricano dinamicamente tramite script – l'AI può individuare le chiamate necessarie e assicurare che tutti i dati vengano caricati e acquisiti correttamente, evitando che qualcosa sfugga alla registrazione. In pratica, questi moduli di intelligenza artificiale fungono da *“assistenti digitali”* durante l'acquisizione: **ottimizzano automaticamente** il processo in base alle caratteristiche del sito target, senza richiedere all'operatore di intervenire manualmente su impostazioni tecniche avanzate. Il risultato è una maggiore **robustezza** e completezza della prova raccolta, poiché l'AI riduce al minimo il rischio di errore umano (come dimenticare di scrollare una pagina, o non accorgersi di contenuti caricati in ritardo) e rende fruibile la procedura anche a utenti meno esperti. Questo connubio tra competenza forense e intelligenza artificiale rappresenta un notevole passo avanti: garantisce acquisizioni **affidabili** anche da fonti online complesse, mantenendo al contempo la semplicità d'uso. Non a caso Kopjra – forte anche di questa innovazione – **risulta tra i pochi soggetti certificati ISO 27037** e vanta brevetti sia in Italia che negli USA per le sue soluzioni

tecnologiche avanzate. In definitiva, l'intervento di Tommaso Grotto ha mostrato come l'IA non sia solo oggetto di indagine per i forensi (quando si tratta di scoprire deepfake), ma diventi essa stessa **strumento** al servizio della forensics, migliorando gli strumenti di investigazione digitale e aprendo la strada a nuove best practice nel trattamento delle prove elettroniche.

7. Conclusioni – Diversità di soluzioni e consigli operativi per gli avvocati

La lezione del 28 maggio 2025 si è conclusa con una serie di riflessioni generali che vale la pena richiamare, perché rappresentano i **punti chiave** emersi trasversalmente dai vari interventi e forniscono indicazioni pratiche per i professionisti legali che vogliano avvicinarsi all'IA:

- **Un mercato eterogeneo e in rapida evoluzione:** come dimostrato dagli esempi presentati, esiste un'ampia gamma di strumenti AI per avvocati, che spazia da piattaforme generaliste (applicabili un po' a tutte le branche del diritto) a soluzioni iper-specializzate (come Tiresia per il tributario). Non esiste *la* soluzione universale migliore in assoluto: la **scelta dipende dalle esigenze specifiche** dello studio e dalla natura delle attività svolte. Un piccolo studio boutique potrebbe trovare sufficiente uno strumento verticale per la propria nicchia, mentre una grande law firm dovrà probabilmente combinare più strumenti per coprire aree diverse.
- **La fonte dei dati fa la differenza:** uno degli elementi distintivi tra le varie soluzioni AI è la base di conoscenza su cui poggiano. È stato ribadito come alcune piattaforme cerchino risposte sul **web aperto**, rischiando di attingere a informazioni non verificate o obsolete (con tutti i problemi di **attendibilità** che ne conseguono, basti pensare a errori o *fake news*). Altre, invece, basano i loro output su **database proprietari di fonti ufficiali e verificate** (legislazione, case law, prassi consolidate). Per l'avvocato, essere consapevole di questa differenza è fondamentale: affidarsi a uno strumento con fonti controllate può dare maggiore tranquillità sul fatto che le analisi o i testi generati siano corretti e citabili in tribunale, mentre usare un generatore generalista richiede sempre un attento fact-checking a posteriori. In sintesi, *“spazzatura in ingresso, spazzatura in uscita”*: la qualità delle risposte dell'IA non può che riflettere la qualità delle fonti utilizzate.
- **La sicurezza e la riservatezza non sono optional:** praticamente ogni intervento ha toccato, da prospettive diverse, il tema della tutela dei dati. Che si tratti di **non conservare gli upload degli utenti** (come fa Lexroom AI), di **anonimizzare i prompt** (AidaGPT), di **tenere tutto in cloud privato** (Gaya-X Hub) o di **cifrare end-to-end le comunicazioni** (Fragmentalis), è chiaro che la gestione dei dati confidenziali rappresenta la sfida principale per l'adozione serena dell'AI negli studi legali. Le soluzioni affrontano la cosa in modo diversificato: alcune offrono semplici garanzie contrattuali o dichiarazioni di conformità GDPR; altre investono in tecniche sofisticate per eliminare o ridurre drasticamente il rischio di leak di informazioni. L'avvocato tecnologicamente informato dovrà sempre chiedersi, prima di adottare un tool: *“Dove finiranno i dati dei miei clienti? Chi*

potrebbe accedervi? È garantito il segreto?” e scegliere di conseguenza strumenti che offrano adeguate risposte a tali domande.

- **Make or Buy – una scelta strategica:** grazie anche al contributo di Aitho, si è compreso che nel campo dell’innovazione digitale ogni studio deve decidere se essere mero utilizzatore di prodotti altrui (cosa perfettamente legittima e spesso efficace) o se diventare co-creatore di soluzioni nuove su misura. Acquistare un prodotto pronto (*buy*) è più semplice, richiede meno competenze tecniche interne e dà risultati immediati, ma offre meno flessibilità ed esclusività (la stessa soluzione la possono avere anche i concorrenti). Sviluppare o personalizzare un prodotto proprio (*make*), invece, è più complesso e costoso sul breve termine, ma può portare a strumenti cuciti sulle proprie prassi e a un vantaggio competitivo distintivo. Non esiste una scelta giusta in assoluto: dipende dalla **visione dello studio**, dalle risorse disponibili e dall’orizzonte temporale con cui si pianifica l’innovazione.
- **L’importanza del testing diretto:** un consiglio molto pratico emerso è quello di **provare con mano** queste tecnologie. Vista la rapidità con cui evolvono e la varietà dell’offerta, l’unico modo per fare una scelta informata è testare direttamente gli strumenti in versione demo o pilot. Molte delle aziende presentate offrono trial gratuiti o progetti pilota. L’Osservatorio IA dell’Ordine di Milano stesso nasce con l’intento di facilitare questo processo, fornendo valutazioni di compliance e occasioni di incontro tra avvocati e fornitori. I relatori hanno incoraggiato i partecipanti a *“sporcarsi le mani”*: solo utilizzando sul campo un’AI per qualche settimana si può capire davvero se fa al caso proprio, se è intuitiva, se i risultati sono affidabili e se si integra bene nel proprio flusso di lavoro.
- **Integrare le soluzioni tra loro – il futuro del Legal Tech:** infine, una visione prospettica condivisa è stata quella della **integrazione**. Molte soluzioni nascono per risolvere problemi specifici in modo verticale. La vera sfida dei prossimi anni sarà farle **dialogare** tra loro, costruendo flussi di lavoro digitali coerenti ed efficienti. Un esempio citato: immaginare uno studio in cui un contratto venga redatto su una piattaforma come Keplera/Lexdo.it, poi analizzato automaticamente da un tool come quello di Keplera o da un AI custom, quindi inviato al cliente tramite Fragmentalis, e magari anche archiviato e monitorato in un AI Hub privato per future consultazioni. Il tutto integrato con l’assistente virtuale IO.AI che informa il team delle scadenze contrattuali o degli adempimenti successivi. Questo scenario di **ecosistema interconnesso** richiederà standard aperti, API e probabilmente nuovi attori che facciano da “collante” tra i vari servizi. Ma è lì che si direziona il mercato: non più singoli strumenti isolati, bensì **suite integrate** dove l’IA pervade l’intero ciclo di attività dell’avvocato, dalla prima chiamata del cliente fino all’archiviazione finale della pratica.

In conclusione, la lezione del 28 maggio ha fornito ai partecipanti non solo una panoramica ricca e concreta dell’attuale offerta di IA applicata al diritto, ma anche una serie di spunti critici su come approcciare questa rivoluzione tecnologica. Il tono generale è stato di cauto entusiasmo: **l’IA è pronta ad aiutare gli avvocati**, a patto che venga utilizzata con consapevolezza, etica e attenzione alla qualità dei dati e alla

sicurezza. L'avvocato del futuro prossimo sarà probabilmente chi saprà combinare il proprio **giudizio professionale** con gli strumenti innovativi disponibili, sfruttando l'IA come alleato per elevare il livello del servizio legale senza mai perdere di vista i principi deontologici e la centralità della figura umana nel sistema giustizia. In altri termini, come affermato da uno dei relatori in chiusura: *“Non abbiate timore che l'IA vi sostituisca; abbiate piuttosto timore di essere sostituiti da un collega che usa bene l'IA.”* Un messaggio che i partecipanti a questa intensa sessione formativa porteranno certamente con sé nella propria pratica quotidiana.

Capitolo V

Intelligenza Artificiale e Deontologia Forense: profili etici e responsabilità

1. Introduzione

La lezione del 4 giugno 2025 – quinta nel *Corso di Perfezionamento in Intelligenza Artificiale per Avvocati* – segna una svolta tematica: l'attenzione si sposta dagli strumenti pratici ai delicati profili **etici e deontologici** legati all'uso dell'Intelligenza Artificiale nella professione legale. Esperti di tecnologia, avvocati e rappresentanti istituzionali si confrontano sulle complesse implicazioni che l'IA comporta per l'avvocatura, con l'obiettivo di sviluppare nei partecipanti una solida coscienza critica. In altri termini, la sessione mira a fornire gli strumenti per navigare le “zone grigie” etiche, gestire in modo trasparente il rapporto con il cliente e integrare l'IA nella pratica quotidiana senza abdicare ai principi fondamentali della professione forense.

Nel corso della giornata si succedono quattro interventi, che seguono un filo logico coerente. **Alessandro Vitale**, esperto di IA e imprenditore, apre i lavori offrendo una panoramica sul funzionamento del *machine learning* e sui rischi intrinseci di questa tecnologia – dai bias alle *allucinazioni* informative, fino alle sfide poste dagli obiettivi assegnati agli algoritmi. **L'Avv. Giorgio Treglia**, coordinatore della Commissione Deontologia dell'Ordine degli Avvocati, affronta poi il tema delle **responsabilità deontologiche** dell'avvocato nell'uso dell'IA, delineando i doveri di competenza, supervisione e trasparenza alla luce dei principi etici della professione. A seguire, **l'Avv. Valentina Gariboldi**, Legal Innovation Manager presso lo studio internazionale Linklaters, cala la discussione nella pratica quotidiana dello studio legale, evidenziando i rischi “invisibili” che un uso disinvolto dell'IA può comportare e suggerendo strategie concrete di **prevenzione e governance interna**. Infine, **l'Avv. Elio Guarnaccia** offre uno sguardo sul **contesto normativo** in evoluzione: illustra un disegno di legge italiano che intende regolamentare l'impiego dell'IA nelle professioni legali e descrive un progetto pilota condotto presso il Tribunale di Catania, volto a integrare l'IA nell'amministrazione della giustizia.

Il filo conduttore che emerge da tutti gli interventi è chiaro: per sfruttare appieno i benefici dell'Intelligenza Artificiale nel settore legale è necessaria **consapevolezza critica** e rispetto dei principi deontologici. Nelle sezioni seguenti ripercorriamo ciascun intervento in ordine cronologico, sviluppando una narrazione accessibile ma fedele agli argomenti e ai casi presentati.

2. Demistificare l'IA: funzionamento, rischi ed etica – Dott. Alessandro Vitale

Alessandro Vitale, esperto di IA e fondatore di una start-up del settore, inaugura la lezione con un intervento introduttivo volto a *demistificare* la tecnologia e a metterne in luce i rischi etici. In modo semplice e intuitivo, Vitale spiega come funziona il **machine**

learning, il motore dell'IA moderna, per poi esplorare i principali pericoli insiti in questi sistemi e come possiamo mitigarli.

In apertura, Vitale ricorre a esempi concreti per far comprendere la logica interna dell'IA. Mostra anzitutto che ciò che noi percepiamo come immagini o testi, un algoritmo lo vede in realtà come pura **matematica**: ad esempio, una cifra scritta a mano non viene “vista” come un *nove*, ma come una griglia di pixel tradotta in 784 numeri. Una rete neurale è, in sostanza, una complessa funzione matematica con migliaia di parametri; “apprendere” significa aggiustare tali parametri in base agli esempi forniti, così che dato un certo input (es. i pixel di un'immagine) la macchina produca l'output corretto (es. riconoscere il numero). Vitale dimostra anche quanto l'IA dipenda dalla **qualità dei dati**: in un esperimento lampante, addestra un modello a distinguere immagini di auto e camion, ma gli fornisce solo foto di *auto bianche* e *camion rossi*. Il risultato? L'algoritmo impara a distinguere i colori, non i veicoli, fallendo l'obiettivo reale. Questo esempio mostra come l'IA possa rivelarsi molto meno “intelligente” di quanto si creda: il suo comportamento è totalmente determinato dai dati con cui viene addestrata, e dataset incompleti o sbilanciati possono portarla fuori strada.

Dopo aver chiarito il funzionamento di base, Vitale identifica **tre livelli di rischio etico** nell'uso dell'Intelligenza Artificiale:

- **Bias (pregiudizi nei dati)** – È il problema più noto e, sotto certi aspetti, il più semplice da individuare. Se i dati di addestramento sono incompleti o contengono pregiudizi, anche il modello riprodurrà quegli errori. Vitale cita esempi divenuti famosi: un algoritmo di riconoscimento immagini di Google che classificava persone di colore come “gorilla”, oppure sistemi che associano in modo stereotipato determinati generi a specifiche professioni. Nella **IA generativa** il tema diventa ancor più intricato: il tentativo di un recente modello di Google, chiamato *Gemini* (progettato per essere “iper-inclusivo”), ha prodotto immagini storicamente inaccurate – ad esempio raffigurando soldati nazisti con tratti somatici del tutto improbabili – nel maldestro tentativo di evitare stereotipi razziali. Per gestire i bias la soluzione tecnica più diffusa è implementare appositi *guardrail*: filtri e vincoli che bloccano le richieste o le risposte inopportune, cercando di prevenire output offensivi o distorti.
- **Obiettivi assegnati agli algoritmi** – Anche definire **cosa** chiediamo a un'IA può generare conseguenze inattese. Vitale porta l'esempio curioso di due algoritmi autonomi usati per fissare i prezzi di libri usati su Amazon: programmati per competere nel tenere il prezzo *appena più alto* l'uno dell'altro, finirono per far lievitare il prezzo di un libro fuori catalogo fino a 23 milioni di dollari! Un altro caso citato riguarda l'algoritmo di YouTube che, avendo come obiettivo massimizzare il *tempo di visualizzazione* degli utenti, tende a suggerire contenuti via via più estremi per catturare l'attenzione e rendere il comportamento dello spettatore più prevedibile. Questi esempi mostrano come un obiettivo mal calibrato – pur in assenza di cattive intenzioni – possa condurre a risultati eticamente o socialmente indesiderabili. La sfida, sottolinea Vitale, sta nel

- definire obiettivi che siano responsabili e positivi per la società**, evitando metriche troppo semplicistiche che possano indurre effetti collaterali.
- **Casi d'uso discutibili** – Infine, anche un algoritmo tecnicamente perfetto e privo di bias può essere impiegato in **ambiti eticamente controversi**. Vitale racconta di una società che ha scelto di *non* adottare un software di IA in grado di analizzare le espressioni facciali durante i colloqui di lavoro. Il motivo non era un difetto tecnico, bensì una valutazione etica: i dirigenti ritenevano sbagliato iniziare un rapporto di lavoro sottoponendo i candidati a uno *screening* automatizzato delle loro emozioni. In questo caso la tecnologia funzionava, ma il suo utilizzo appariva lesivo della dignità e della fiducia alla base del rapporto azienda-dipendente. L'aneddoto evidenzia dunque un terzo livello di riflessione etica: oltre ai bias e agli obiettivi, bisogna chiedersi **per quali scopi stiamo usando l'IA**, e se tali scopi siano accettabili dal punto di vista morale e sociale.

Affrontati i rischi, il discorso si sposta su **come mitigare** questi pericoli. Vitale sottolinea che nessun sistema di IA sarà mai infallibile: l'errore, in misura prima o poi, è inevitabile. Per questo è fondamentale progettare una **governance** intorno all'IA, fatta di linee guida chiare, test rigorosi e procedure di controllo. Ogni nuovo algoritmo andrebbe sottoposto a verifiche per individuare bias nascosti e comportamenti anomali *prima* di un impiego esteso. Inoltre, nel disegnare processi che integrano l'AI, bisogna prevedere sempre la possibilità di *intervento umano* a valle: in altre parole, creare canali di **appello e revisione** in caso di errore della macchina. Vitale cita a tal proposito un caso emblematico in negativo: un algoritmo utilizzato in ambito di welfare pubblico che, sulla base dei suoi calcoli, revocava automaticamente alcuni sussidi a persone in difficoltà, **senza prevedere alcuna procedura di ricorso o verifica umana**. Un tale approccio "cieco" è da scongiurare: la tecnologia deve essere inserita in processi che consentano all'utente o al cittadino di chiedere una rettifica e all'operatore umano di correggere eventuali ingiustizie causate dalla macchina.

In chiusura, Vitale volge lo sguardo al **futuro prossimo dell'IA**, parlando di una nuova generazione di sistemi sempre più autonomi. Mostra esempi di *agenti intelligenti* capaci di compiere azioni sul web in modo indipendente – ad esempio scaricare documenti contabili, prenotare appuntamenti o compilare una checklist di compliance normativa – senza un intervento passo-passo dell'utente. Questi agenti IA rappresentano un'evoluzione affascinante, ma pongono anche una **sfida di supervisione**: più l'IA acquisisce autonomia operativa, più l'utente umano dovrà mantenere un livello elevato di consapevolezza e controllo. In altre parole, conclude Vitale, l'avvocato (così come qualsiasi professionista) non potrà mai "distrarsi": dovrà anzi sviluppare nuove competenze per monitorare efficacemente il lavoro svolto dall'IA, assicurandosi che resti entro i binari desiderati e intervenendo prontamente in caso di deriva. Questa necessità di supervisione umana attiva è un tema che apre direttamente al cuore della discussione deontologica successiva.

3. Responsabilità professionale e doveri deontologici nell'era dell'IA - Avv. Giorgio Treglia

Subito dopo l'introduzione tecnica di Vitale, l'**Avvocato Giorgio Treglia** prende la parola per esplorare il tema cruciale della **responsabilità professionale** nell'uso dell'Intelligenza Artificiale. Treglia, che coordina la Commissione Deontologia presso l'Ordine degli Avvocati di Milano, affronta la questione dal punto di vista etico e disciplinare: quali sono i doveri dell'avvocato quando utilizza strumenti di IA? Quali errori bisogna evitare e di chi è la colpa se qualcosa va storto?

Per chiarire immediatamente il principio base, Treglia esordisce con un caso concreto che ha fatto scalpore: presso il Tribunale di Firenze, un avvocato aveva depositato una memoria difensiva contenente riferimenti giurisprudenziali completamente **inventati** dall'IA che aveva utilizzato per redigerla. Scoperto l'errore, il legale si era giustificato affermando: *“È stato il praticante a preparare l'atto, non me ne sono accorto”*. Una difesa del genere – commenta Treglia – è **drammatica** sotto il profilo deontologico, perché mette in luce un duplice fallimento: da un lato la mancata formazione del praticante, dall'altro la mancata supervisione da parte del dominus (l'avvocato senior responsabile). L'episodio serve a introdurre un concetto per Treglia imprescindibile: **la responsabilità è sempre e solo dell'avvocato**. Se l'Intelligenza Artificiale commette un errore, non la si può additare come capro espiatorio: la colpa ricade comunque su chi l'ha impiegata senza le necessarie cautele. Delegare a un algoritmo le proprie responsabilità professionali non è mai ammesso. In sintesi, l'avvocato risponde in prima persona degli atti che firma, indipendentemente dagli strumenti – umani o artificiali – di cui si è avvalso per prepararli.

Strettamente legato a questo tema è il **dovere di competenza**. La deontologia forense impone all'avvocato di mantenere competenze adeguate non solo nel diritto sostanziale e processuale, ma anche negli strumenti che utilizza. Treglia esprime preoccupazione per l'atteggiamento di alcuni giovani professionisti che vedono nell'IA una scorciatoia capace di sollevarli dallo studio e dall'aggiornamento continuo. Questo approccio passivo è pericoloso: l'IA deve essere intesa come un *ausilio*, non come un sostituto della conoscenza. Il rischio, avverte Treglia, è un **processo di “de-skilling”**: affidarsi ciecamente alle risposte della macchina può far perdere, col tempo, la capacità di ragionare in maniera giuridica autonoma. La formazione classica, lo studio dei principi fondamentali del diritto e l'esercizio quotidiano dell'analisi critica rimangono pilastri insostituibili. Anche nell'era digitale, insomma, un avvocato competente è colui che sa utilizzare gli strumenti tecnologici **senza diventarne dipendente**, conservando la padronanza della materia e del metodo giuridico.

Un altro punto cardine del discorso di Treglia è la **centralità della decisione umana** e il conseguente dovere di *controllo delle fonti*. Richiamando la *Carta dei Principi* promulgata dall'Ordine milanese, Treglia ribadisce che l'avvocato deve rimanere al timone delle scelte, senza lasciarle in balia di un algoritmo. In termini pratici, ciò significa

verificare accuratamente **ogni informazione** fornita dall'IA, esattamente come si controllerebbe a fondo una citazione giurisprudenziale ricevuta dalla controparte. Treglia fa un efficace paragone: se leggendo una sentenza citata dal nostro avversario notiamo un *“non” di troppo* (o di meno) che stravolge il senso di un principio, sentiamo il dovere di andare a recuperare quella sentenza e controllarla con i nostri occhi. Allo stesso modo, qualsiasi output generato dall'IA va preso con beneficio d'inventario e sottoposto al vaglio critico del professionista. In questa ottica, l'avvocato deve essere sempre **colui che “mette il dato”** nell'IA (formulando le giuste richieste, fornendo i fatti e i documenti pertinenti) e poi colui che interpreta e *valuta criticamente* il risultato ottenuto. L'IA, per quanto avanzata, resta uno strumento nelle mani del giurista: spetta a quest'ultimo il compito di assicurarsi che lo strumento lavori correttamente e che il prodotto finale sia accurato e aderente allo scopo.

Nel prosieguo, Treglia affronta un tema molto pratico: **la formazione dei giovani avvocati** nell'uso consapevole dell'IA. In risposta a una domanda posta dal coordinatore del corso, l'Avv. Giuseppe Vaciago, egli delinea il corretto approccio che deve svilupparsi tra *dominus* e *praticante* (mentor e tirocinante). Da un lato, **il dominus** ha il dovere di formare il praticante, spiegandogli i principi generali del diritto *prima* di affidargli ricerche o redazioni assistite dall'IA. Il giovane va coinvolto non come mero esecutore di compiti ripetitivi, ma come risorsa proattiva: spesso i praticanti padroneggiano le nuove tecnologie meglio dei loro maestri, e questo patrimonio di competenze può essere valorizzato in un dialogo costruttivo. Dall'altro lato, **il praticante (o neo-avvocato)** deve avere l'umiltà di continuare a studiare i fondamenti giuridici e, al contempo, il coraggio di “pretendere” formazione e confronto dal proprio dominus. È importante che il giovane professionista utilizzi l'IA per *affinare* la ricerca e velocizzare i compiti, ma non per sostituire lo studio personale. Solo così il tirocinio rimane un percorso di crescita reale e non si trasforma in un semplice *data entry* dietro cui celare lacune formative.

In conclusione al suo intervento, Treglia tocca un ultimo aspetto: **il rischio delle fake news e il valore della trasparenza**. La *trasparenza* nell'uso degli strumenti di IA – sottolinea – non è solo un obbligo deontologico verso il cliente e la controparte, ma anche un accorgimento che può tutelare lo stesso avvocato. Dichiarare espressamente di aver utilizzato un certo software o sistema intelligente per raggiungere un risultato dimostra infatti **competenza tecnica e buona fede** da parte del legale. Ciò aiuta anche a circoscrivere la sua responsabilità in caso di errore della macchina. Ad esempio, se un avvocato informa il cliente che per una prima bozza di contratto ha impiegato un algoritmo di supporto, sta implicitamente comunicando di aver adottato uno strumento moderno, ma di essere consapevole dei suoi possibili errori (e infatti poi controllerà la bozza attentamente). In caso emerga un problema dovuto a un *bug* o a un output fantasioso dell'IA, il fatto di aver agito con trasparenza e perizia può attenuare un giudizio di negligenza. Al contrario, usare l'IA di nascosto – magari per redigere un parere che si rivela impreciso – espone a rischi sia legali che reputazionali. In sostanza, **onestà e trasparenza** sull'uso dell'Intelligenza Artificiale fanno parte integrante del dovere di

corretta informazione e possono trasformarsi in un'ulteriore garanzia a favore dell'avvocato diligente.

(Dopo aver stabilito che l'IA non libera l'avvocato dalle sue responsabilità – anzi, richiede ancora maggiore attenzione e competenza – la lezione passa ad esaminare quali rischi pratici possono annidarsi nell'uso quotidiano di queste tecnologie e come predisporre misure efficaci per prevenirli. È questo il tema al centro dell'intervento di Valentina Gariboldi.)

4. Rischi “invisibili” nell'uso quotidiano dell'IA e strategie di prevenzione - Avv. Valentina Gariboldi

Il terzo intervento vede protagonista **l'Avvocato Valentina Gariboldi**, Legal Innovation Manager presso lo studio legale internazionale Linklaters. Gariboldi sposta il fuoco dell'attenzione dalla teoria alla pratica quotidiana dello studio legale, evidenziando i rischi più **insidiosi e nascosti** legati all'impiego disinvoltato dell'IA da parte di professionisti e collaboratori, e proponendo strategie concrete per prevenirli. Il taglio del suo contributo è fortemente operativo: attraverso esempi e *scenari tipo*, la relatrice mostra come anche attività all'apparenza banali possano celare insidie etiche, deontologiche o di sicurezza.

Anzitutto Gariboldi invita a riflettere sui **“rischi invisibili”**, quelli meno ovvi. Mentre un errore grossolano nel contenuto di un atto legale è qualcosa che un avvocato esperto può individuare e correggere relativamente facilmente applicando la dovuta diligenza, esistono pericoli più subdoli che spesso vengono sottovalutati. I principali sono collegati alla **riservatezza dei dati** e alla **sicurezza informatica**. In un'epoca in cui molti servizi di IA sono disponibili tramite piattaforme online di terze parti, c'è il rischio concreto di condividere inavvertitamente informazioni confidenziali dello studio o dei clienti con sistemi esterni, magari non protetti a dovere. Inoltre, affidarsi a strumenti digitali richiede attenzione sul fronte della *cybersecurity*: l'uso improprio di software non autorizzati o la mancanza di consapevolezza sulle impostazioni di privacy possono esporre lo studio a violazioni di dati e responsabilità.

Per dare sostanza a questi concetti, Gariboldi costruisce uno **scenario pratico**, scherzosamente battezzato del *“caffè mattutino”*, in cui descrive due situazioni tipo realmente plausibili in qualsiasi studio legale.

1. **La traduzione “fai-da-te” del praticante:** un avvocato chiede al praticante, chiamiamolo Marco, di tradurre un documento legale urgente. Marco, come milioni di altre persone, ricorre alla via più rapida: copia e incolla il testo in un traduttore online gratuito (ad esempio DeepL o Google Translate). Così facendo, però, *senza saperlo divide il contenuto riservato del documento con un servizio terzo su Internet*. Le policy delle versioni gratuite di questi servizi non garantiscono la confidenzialità: i dati inseriti potrebbero essere conservati sui loro server o analizzati per migliorare gli algoritmi. In pratica, l'entusiasmo

tecnologico del giovane – unito a inesperienza – ha creato un **breach** nella segretezza professionale, violando potenzialmente il dovere di riservatezza verso il cliente. **Soluzione:** secondo Gariboldi, episodi del genere vanno prevenuti dotando lo studio di strumenti professionali approvati e sicuri (ad esempio versioni *Pro* di servizi come DeepL, che offrono garanzie contrattuali di privacy dei dati) e bloccando a livello informatico l'accesso ai traduttori gratuiti o ad altre applicazioni non autorizzate. Inoltre, è fondamentale educare i collaboratori: far capire a Marco e colleghi che *non* si devono usare piattaforme esterne non controllate per dati sensibili, nemmeno per comodità. Una *policy* interna chiara e mezzi tecnici adeguati possono evitare che la frenesia di risolvere un compito in pochi clic si traduca in un disastro per la sicurezza delle informazioni.

2. **La trascrizione automatica di una videochiamata:** in un secondo caso, durante una videoconferenza con la controparte, un avvocato attiva la funzione di trascrizione in tempo reale offerta da un assistente virtuale (ad esempio Microsoft *Copilot* integrato nella piattaforma di meeting online). L'idea è comoda: ottenere subito il verbatim della discussione. Tuttavia, si apre un **dilemma etico-deontologico**. Infatti l'art. 38 del Codice Deontologico Forense richiede il **consenso di tutti i presenti** per registrare una conversazione, anche tra soli colleghi. Questa regola vale anche per la “mera” trascrizione automatica? E inoltre, chi garantisce la **correttezza** dell'output generato dall'IA? Un trascrittore automatico potrebbe commettere errori o, peggio, conservare nei server del fornitore cloud il contenuto della riunione. **Soluzione:** Gariboldi raccomanda di agire con estrema trasparenza e precauzione. Prima di tutto, *mai* attivare registrazioni o trascrizioni senza aver informato chiaramente tutti i partecipanti e ottenuto il loro esplicito consenso. È opportuno poi familiarizzare con i *settings* delle piattaforme digitali che usiamo: ad esempio, verificare dove vengono salvati gli eventuali file audio/testo e chi vi ha accesso. Se persiste incertezza sul rispetto delle norme deontologiche, meglio rinunciare alla funzione automatica e prendere appunti alla vecchia maniera, oppure concordare formalmente con le altre parti modalità sicure di verbalizzazione. L'importante, sottolinea Gariboldi, è **non farsi cogliere impreparati**: la tecnologia offre scorciatoie, ma l'avvocato deve valutare caso per caso se percorrerle o meno, tenendo presente sia l'etica sia la conformità alle regole.

Dopo aver messo in guardia contro questi rischi specifici, Valentina Gariboldi passa ad illustrare come uno studio legale possa attrezzarsi in senso più **sistemico** per un'adozione responsabile dell'IA. A tal fine, presenta il *percorso in cinque step* che il suo studio ha adottato come **governance interna** dell'innovazione:

1. **Definire la governance** – Stabilire chi fa cosa in materia di IA. In pratica, va creato un gruppo interdisciplinare che guidi la transizione tecnologica: non solo avvocati, ma anche figure di IT e compliance, per assicurare una visione d'insieme su aspetti legali, tecnici e regolamentari. Questo team avrà il compito di valutare i nuovi strumenti di IA, dettare linee guida e monitorare il rispetto delle policy.
2. **Elaborare i principi guida** – Prima di lanciarsi nell'uso degli strumenti, è utile che lo studio adotti una sorta di *carta etica interna* sull'IA, ovvero una filosofia

condivisa. Ad esempio: l'IA deve essere concepita come un **acceleratore** dei processi, non come un sostituto del lavoro umano; la responsabilità finale su ogni attività rimane sempre in capo alle persone. Chiarire e mettere per iscritto questi principi aiuta tutti i membri dello studio a capire lo spirito con cui la tecnologia verrà impiegata.

3. **Passare dalla teoria alla pratica** – Dopo i principi, bisogna attuarli concretamente. Gariboldi sottolinea due azioni chiave: **formazione diffusa** e **sperimentazione controllata**. Da un lato, è stato avviato un programma di *alfabetizzazione di massa*: training obbligatorio per tutto il personale – avvocati, praticanti e staff – sulle potenzialità e i rischi dell'IA, in modo da creare consapevolezza diffusa (ognuno deve capire almeno le basi di come funziona uno strumento di AI e quali sono i *red flag* da tenere d'occhio). Dall'altro lato, ogni nuovo software o sistema viene sottoposto a **test approfonditi** in ambiente controllato prima dell'adozione su larga scala: si effettuano prove pilota, si valutano gli output, si cercano bug o bias, così da poter correggere rotta o scegliere strumenti alternativi se necessario, prima che l'intero studio ne diventi dipendente.
4. **Formalizzare policy e procedure** – Uno step cruciale è tradurre principi e lezioni apprese in **procedure scritte**. Occorre redigere vere e proprie policy aziendali su ciò che è consentito o vietato fare con l'IA nello studio. Ad esempio: quali tipi di dati si possono caricare su quali piattaforme (e quali assolutamente no), quali strumenti sono approvati per certe attività (es. traduzioni, ricerche, analisi di dati) e quali invece banditi, come gestire la verifica degli output generati dalle macchine, ecc. Inoltre, Gariboldi consiglia di *mappare* ogni strumento di IA utilizzato nello studio e per ciascuno definire chiaramente le regole d'uso, i responsabili e le misure di sicurezza da adottare. Questa documentazione deve essere facilmente accessibile a tutti in modo che non vi siano incertezze operative.
5. **Monitoraggio continuo** – Infine, la governance non può considerarsi fissa una volta per tutte: poiché la tecnologia evolve rapidamente, è necessario prevedere un meccanismo di **rivalutazione periodica** delle policy e dei processi. Nuovi rischi emergono, normative si aggiornano, compaiono sul mercato strumenti migliori: lo studio deve restare agile e pronto ad adeguare le proprie regole interne. Un esempio citato è l'aggiornamento costante dei filtri e delle impostazioni man mano che le piattaforme cloud introducono nuove funzionalità o cambi nei termini di servizio. In sostanza, l'adozione dell'IA va gestita come un **processo dinamico**, con verifiche e migliorie continue, piuttosto che come un progetto implementato e archiviato.

Chiudendo il suo intervento, Valentina Gariboldi offre una prospettiva incoraggiante: **etica e governance non sono un freno, ma un acceleratore** per un'adozione dell'IA che sia consapevole, sicura e – in ultima analisi – più proficua. In altre parole, investire tempo ed energie per mettere paletti etici e regole all'uso dell'Intelligenza Artificiale non significa ostacolare l'innovazione, bensì creare le condizioni per sfruttarla al meglio evitando incidenti di percorso. Un approccio ben regolato permette infatti ai professionisti di fidarsi maggiormente della tecnologia e usarla con serenità, traendone

tutti i vantaggi. La conclusione di Gariboldi è che **un'IA “governata” bene può accelerare lo studio legale verso il futuro**, mantenendo però i valori fondamentali intatti.

(Dopo aver esaminato soluzioni pratiche di governance interna, l'ultima parte della lezione allarga l'orizzonte al contesto esterno: la normativa in cantiere e le esperienze istituzionali sul campo. Questo è il focus dell'intervento di Elio Guarnaccia.)

5. Verso una regolamentazione dell'IA: il DDL professioni e la sperimentazione nel Tribunale di Catania - Avv. Elio Guarnaccia

A conclusione della giornata interviene l'**Avvocato Elio Guarnaccia**, del Foro di Catania, che porta l'attenzione sul quadro normativo e sulle iniziative pubbliche relative all'Intelligenza Artificiale nel settore legale. Il suo contributo si articola in due filoni: da un lato l'analisi di un **disegno di legge (DDL)** italiano dedicato all'uso dell'IA nelle professioni intellettuali, dall'altro la descrizione di una **sperimentazione pratica** condotta presso il Tribunale di Catania in collaborazione con un ente di ricerca, per integrare l'IA nel lavoro giudiziario.

Il disegno di legge italiano sull'IA nelle professioni forensi. Guarnaccia riferisce che a livello nazionale è in discussione un DDL che mira a disciplinare l'impiego dell'Intelligenza Artificiale da parte di avvocati (e altre professioni intellettuali). Due sono i punti salienti di questa proposta normativa. In primo luogo, si **limita** espressamente l'uso dell'IA ad attività “strumentali o di supporto”, ribadendo che deve comunque prevalere l'apporto intellettuale umano del professionista. Ciò significa, ad esempio, che un avvocato potrebbe utilizzare software di IA per ricerche giurisprudenziali, analisi di documenti o altre attività ausiliarie, ma non delegare alla macchina decisioni o prestazioni professionali vere e proprie (come scrivere pareri legali completi o interagire con i clienti senza controllo umano). Questa formulazione tuttavia pone un problema di interpretazione: non è banale definire con precisione che cosa rientri nell'ambito “strumentale” e cosa no, lasciando una potenziale area grigia. Il secondo punto chiave del DDL è l'introduzione di un **obbligo di trasparenza** verso il cliente: l'avvocato dovrà informare in modo chiaro e semplice il proprio assistito ogniqualvolta impieghi sistemi di IA nello svolgimento dell'incarico, spiegandone l'utilizzo. Anche questa previsione, nota Guarnaccia, ha l'obiettivo di preservare la fiducia nel rapporto professionale e garantire consapevolezza da parte del cliente.

Tali norme, seppur animate da intenti condivisibili, hanno già suscitato **critiche**. La Commissione Europea, per esempio, ha espresso perplessità temendo che limitazioni troppo rigide possano confliggere con l'approccio più flessibile previsto dal regolamento UE in materia di IA (il cosiddetto *AI Act*, che per i sistemi non ad alto rischio non impone restrizioni così stringenti). Anche alcuni operatori del settore tecnologico, come la stessa OpenAI, vedono nel DDL italiano un potenziale ostacolo all'innovazione e alla competitività dei professionisti italiani, i quali si troverebbero gravati da obblighi aggiuntivi non richiesti altrove. Guarnaccia evidenzia quindi il delicato equilibrio che il

legislatore dovrà ricercare: garantire un uso etico e responsabile dell'IA senza però imbrigliare eccessivamente le opportunità offerte dalla tecnologia.

Un ulteriore aspetto di grande rilievo toccato dal DDL è che la **formazione in materia di IA** diventa un vero e proprio obbligo istituzionale. In pratica, la legge delega il Governo a predisporre percorsi obbligatori di alfabetizzazione e aggiornamento sull'Intelligenza Artificiale destinati agli avvocati. Ciò significa che l'acquisizione di competenze digitali e la comprensione dei meccanismi di funzionamento dell'IA saranno considerate parte integrante della formazione continua obbligatoria dell'avvocato. Si tratta di una novità assoluta, che ricollega l'Italia a quanto previsto a livello europeo: l'articolo 4 dell'**AI Act** europeo (dedicato all'"AI literacy") incoraggia infatti gli Stati e gli ordini professionali a promuovere competenze diffuse sull'IA, e le FAQ della Commissione UE hanno chiarito che potranno essere previste sanzioni in caso di inadempimento. In altre parole, l'aggiornamento del legale sulle nuove tecnologie non sarà più lasciato alla buona volontà del singolo, ma diverrà un **dovere codificato**, la cui violazione potrebbe avere conseguenze disciplinari. Guarnaccia sottolinea come questo cambiamento di paradigma miri a colmare quel gap di conoscenze di cui si parlava in precedenza: non sarà più ammesso che un avvocato *ignori* il funzionamento di base di strumenti che magari usa quotidianamente, perché la legge stessa imporrà di istruirsi al riguardo.

Dopo la teoria legislativa, Guarnaccia porta un esempio concreto di **sperimentazione sul campo** che vede protagonista l'IA nel sistema giustizia: l'accordo di collaborazione avviato tra il consorzio universitario *CINI* (Consorzio Interuniversitario Nazionale per l'Informatica) e la sezione specializzata in materia di impresa del Tribunale di Catania. Questa collaborazione, stipulata ai sensi di una norma che permette alle pubbliche amministrazioni di cooperare con enti di ricerca, ha un oggetto ben definito: *non* la redazione automatica di sentenze o decisioni, bensì lo sviluppo di un **supporto tecnologico alla ricerca giurisprudenziale** per i giudici e i cancellieri. In pratica, i ricercatori del CINI impiegano modelli di *Natural Language Processing* (NLP) per analizzare un vasto corpus di sentenze passate (opportunamente anonimizzate) allo scopo di individuare più rapidamente precedenti utili, schemi decisionali ricorrenti, tempi medi di definizione delle controversie e altri elementi statistici. Questo sistema intelligente funge dunque da *assistente virtuale* per il magistrato, velocizzando l'attività di studio e inquadramento dei casi, specie quelli più complessi o con molta documentazione. L'iniziativa viene presentata come orientata a finalità di **interesse pubblico**: da un lato mira a perseguire i principi di buona amministrazione – rendere la giustizia più efficiente ed economica nei tempi e nelle risorse impiegate – dall'altro aspira a semplificare il lavoro quotidiano negli uffici giudiziari, riducendo attività ripetitive e liberando energie per gli aspetti decisionali di merito che competono al giudice. È importante notare che, grazie a questo progetto, il Tribunale di Catania non delega mai all'IA decisioni o valutazioni giuridiche, ma la utilizza come **strumento conoscitivo**, un po' come un motore di ricerca avanzato specializzato sul proprio archivio di precedenti.

In chiusura, Guarnaccia propone una riflessione sul **potenziale dell'IA come livellatore** delle disparità nel mondo legale. Progetti come quello catanese – osserva – qualora dovessero essere previsti anche nel mondo forense, se orientati a sviluppare soluzioni semplici, efficaci e magari basate anche su sistemi open-source o comunque accessibili, possono avere una ricaduta sociale significativa. I *prompt* e gli strumenti messi a punto in quella sede, per esempio, potrebbero essere condivisi e adottati anche dai piccoli studi legali o dagli uffici meno dotati di risorse, che altrimenti faticherebbero a permettersi costose piattaforme commerciali di legal tech. In questo modo l'Intelligenza Artificiale potrebbe contribuire a **ridurre il divario digitale** nella professione forense, fornendo anche ai giuristi di realtà più piccole la possibilità di non rimanere indietro rispetto alle grandi *law firm* internazionali. L'auspicio finale è quindi quello di un'IA democratizzata, il cui impiego responsabile e diffuso aiuti l'intero sistema giustizia a progredire in efficienza ed equità, senza creare nuove disparità ma anzi colmandone di esistenti.

Capitolo VI

L'Intelligenza Artificiale e l'implementazione negli Studi Legali

1. Introduzione

La sesta lezione del **Corso di Perfezionamento in Intelligenza Artificiale per Avvocati** (12 giugno 2025) ha esplorato a fondo strategie, strumenti ed esperienze pratiche legate all'adozione dell'IA negli studi legali. Attraverso gli interventi di professionisti provenienti da realtà diverse – dai grandi studi internazionali a primarie boutique legali nazionali – la sessione ha offerto una visione poliedrica delle sfide e delle opportunità che la tecnologia presenta al mondo forense. Sono state analizzate le dinamiche di mercato, le metodologie per la scelta e l'implementazione degli strumenti, i modelli di governance interna (come la creazione di team di “*prompt lawyers*”) e specifici casi d'uso che già stanno trasformando il modo di lavorare degli avvocati. Particolare attenzione è stata dedicata al confronto tra studi di diverse dimensioni, sfatando il mito che l'innovazione sia appannaggio esclusivo delle grandi law firm e mostrando come agilità e specializzazione possano costituire un vantaggio competitivo anche per realtà più piccole.

L'obiettivo formativo dichiarato della sessione era fornire ai partecipanti una sorta di *bussola* per orientarsi nel complesso panorama del legal tech, condividendo lezioni apprese, errori da evitare e un approccio metodologico per integrare l'IA in modo efficace e responsabile – mantenendo sempre al centro il valore della professione forense e la supervisione umana. Nei paragrafi che seguono ripercorriamo in tono divulgativo i contributi dei cinque relatori, seguendo l'ordine cronologico degli interventi: Avv. **Giuseppe Vaciago**, Avv. **Massimo Sterpi**, Avv. **Gianluca Gilardi**, Avv. **Andrea Tuninetti** e Avv. **Benedetto Lonato**. Ciascuna sezione sviluppa una narrazione coerente e fedele ai contenuti esposti, arricchita da esempi concreti e collegamenti tematici tra le diverse prospettive emerse durante la lezione.

2. Rimescolare le Carte: l'IA come opportunità per Studi di ogni dimensione - Avv. Giuseppe Vaciago

Aprendo la giornata, l'Avv. **Giuseppe Vaciago** ha offerto una riflessione strategica sul ruolo dell'IA come potenziale **livellatore** nel mercato legale, sfidando l'idea diffusa che solo i grandi studi possano beneficiare dell'innovazione tecnologica. In altre parole, Vaciago invita a *rimescolare le carte*: l'adozione dell'IA può dare anche ai piccoli e medi studi legali nuove opportunità per competere, a patto di affrontarla con la giusta strategia.

Uno dei primi concetti affrontati da Vaciago è il **falso mito del “first adopter”**. Spesso si pensa che chi adotta per primo una nuova tecnologia godrà di un vantaggio automatico; Vaciago mette in guardia da questa convinzione, portando una significativa esperienza

personale. Racconta di aver acquistato, nell'ottobre 2022, un costoso software di *Natural Language Processing* (NLP) basato su tecnologie di *machine learning* tradizionale, appena un mese **prima** che venisse lanciato ChatGPT. Il risultato? Quell'investimento ingente divenne quasi immediatamente obsoleto con l'arrivo di uno strumento più avanzato ed economico. La lezione, sottolinea l'avvocato, è chiara: adottare una tecnologia per primi comporta rischi significativi e non garantisce affatto un vantaggio competitivo duraturo. Meglio quindi valutare con attenzione **quando** e **cosa** adottare, invece di farsi trascinare dalla sola novità.

A questo aneddoto segue una vivida metafora proposta da Vaciago: il **“gommoni” contro il “transatlantico”**. L'immagine contrappone i piccoli studi legali alle grandi law firm, evidenziandone punti di forza e limiti rispetto all'innovazione. Il *transatlantico* – metafora del grande studio – gode di enormi vantaggi in termini di economie di scala, capacità d'investimento e basi di conoscenza collettiva; tuttavia, è lento nel cambiare rotta e deve affrontare costi elevati per implementare qualsiasi novità (basti pensare alla differenza tra dover acquistare 100 licenze software invece di 5). Il *gommoni* – il piccolo studio – è invece agile: può permettersi il “lusso di sbagliare” perché un eventuale errore ha costi contenuti, e può testare soluzioni innovative con maggiore rapidità e flessibilità. In altre parole, un piccolo studio può spesso arrivare prima a destinazione e offrire stimoli e idee innovative anche ai “transatlantici” più lenti. Questa metafora sfata il mito che solo i giganti possano innovare con successo, evidenziando come la **dimensione ridotta possa trasformarsi in un vantaggio** nell'adozione dell'IA (capacità di reagire rapidamente ai cambiamenti, minor inerzia e minor costo del cambiamento).

Un ulteriore tema toccato da Vaciago riguarda i **costi imprevedibili** nell'ambito dell'IA e le conseguenti scelte tecnologiche. Nel mercato in rapida evoluzione dell'intelligenza artificiale, i modelli di pricing possono cambiare drasticamente da un giorno all'altro. Vaciago cita, ad esempio, l'impennata dell'abbonamento premium di OpenAI (la piattaforma dietro ChatGPT) schizzato in breve tempo da 20\$ a 200\$ al mese. Un incremento di dieci volte nel costo non era prevedibile e può mettere in difficoltà chi aveva basato la propria strategia su costi inizialmente bassi. I piccoli studi, grazie alla loro maggiore flessibilità (ad esempio preferendo abbonamenti mensili anziché contratti pluriennali), sono spesso in posizione migliore per gestire queste oscillazioni. Alla luce di ciò, Vaciago delinea **due possibili strategie** per implementare strumenti di IA in uno studio legale, bilanciando investimenti e rischi futuri:

- **Soluzione “on-premise” (locale):** installare sui propri server modelli di IA open source (ad esempio *Llama* di Meta o il modello cinese *Qwen*), sostenendo un investimento iniziale di circa 10.000€ e tenendo conto dei costi energetici di esercizio. Questa scelta offre massimo controllo sui dati e sul funzionamento dei modelli, ma richiede competenze tecniche interne e risorse hardware dedicate.
- **Soluzione in cloud (SaaS):** utilizzare servizi e software legali “pronti all'uso” forniti sul cloud (Software as a Service). Questa opzione riduce l'onere infrastrutturale interno e consente di iniziare rapidamente, ma Vaciago consiglia

di non vincolarsi con contratti a lungo termine, dato il ritmo serrato con cui nuove soluzioni emergono sul mercato. In pratica, è preferibile mantenere la **libertà di cambiare fornitore** o tecnologia man mano che l'IA evolve, evitando di restare intrappolati in piattaforme che potrebbero diventare obsolete in pochi anni.

In conclusione, del suo intervento introduttivo, Vaciago formula un **principio guida** che funge da filo conduttore per l'intera lezione: *prima di investire in IA, bisogna capire esattamente di cosa si ha bisogno*. Non bisogna lasciarsi abbagliare dal modello più potente o dall'ultima novità, se questa non risponde a un'esigenza concreta del proprio lavoro. È fondamentale identificare con chiarezza il problema o il compito che si vuole risolvere con l'IA. Ad esempio, spiega Vaciago, **estrarre dati da un PDF per compilare un decreto ingiuntivo** è un'attività che richiede un livello di complessità (e un costo computazionale) enormemente inferiore rispetto, poniamo, alla **redazione di un complesso parere legale**. In casi del genere, non serve il modello di IA più avanzato sul mercato; soluzioni più semplici ed economiche possono svolgere egregiamente il lavoro. Per questo *il primo passo* consigliato è spesso quello di sperimentare con strumenti accessibili e relativamente semplici, allo scopo di capire sul campo come usarli e quali sono i propri reali bisogni. Solo *dopo* aver acquisito questa consapevolezza, si potrà procedere a investimenti più mirati e strutturati.

3. Navigare il Mercato Legal Tech: Scelta, Adozione e Rischi nell'Uso dell'IA - Avv. Massimo Sterpi

Proseguendo il filo logico, l'Avv. **Massimo Sterpi** ha offerto una panoramica molto concreta e pragmatica delle sfide che uno **studio legale strutturato** deve affrontare nel processo di adozione dell'IA. Il suo intervento, idealmente, prende le mosse dal discorso di Vaciago: dopo aver compreso le potenzialità dell'IA per studi di ogni dimensione, come ci si orienta tra le tante soluzioni disponibili sul mercato? Come scegliere gli strumenti giusti, come implementarli e quali rischi considerare? Sterpi, con la sua esperienza, cerca di fornire risposte a queste domande, guidando l'uditorio attraverso *tre* grandi temi: la scelta del prodotto sul mercato (*globale vs locale*), il processo interno di adozione (formazione, licenze) e i rischi da presidiare (privacy, veridicità, deontologia), per poi concludere con una riflessione sull'impatto generazionale.

Anzitutto, Sterpi analizza il **mercato del legal tech** evidenziando una dicotomia irrisolta tra piattaforme globali e locali, almeno allo stato attuale. Le piattaforme **internazionali** (ad esempio soluzioni celebri come *Harvey* o *CoCounsel*) sono strumenti potentissimi per l'elaborazione del linguaggio naturale – eccellono nel riassumere testi, generare bozze, fare analisi in lingua inglese – ma presentano un limite significativo: **non sono addestrate sul diritto italiano** e sulla nostra giurisprudenza. Di converso, le piattaforme **italiane** (come *Lexroom* o *Sapienza* di Giuffrè) hanno una conoscenza approfondita del diritto nazionale e risultano più efficaci nelle ricerche giuridiche in lingua italiana; tuttavia, sul piano delle funzionalità avanzate di generazione ed elaborazione del linguaggio sono ancora in una fase più embrionale rispetto alle controparti americane.

In pratica, le soluzioni “nostrane” comprendono meglio le nostre leggi ma sono meno mature tecnologicamente, mentre quelle d’oltreoceano sono sofisticate ma culturalmente “ignoranti” sul nostro ordinamento. La soluzione attuale, osserva Sterpi, è probabilmente una **combinazione ibrida** dei due tipi di strumenti: i grandi studi internazionali stanno infatti sperimentando approcci integrati, utilizzando ad esempio modelli globali per le capacità di *language processing* avanzato, affiancati da strumenti locali o banche dati nazionali per colmare il gap di conoscenza giuridica locale. Questo consente di sfruttare i punti di forza di entrambi, mitigandone i rispettivi limiti.

In secondo luogo, Sterpi affronta la **sfida interna della scelta e adozione dell’IA** in uno studio legale di tipo *general practice*, cioè con molteplici aree di attività (dal tributario al M&A, dal contenzioso alla consulenza). In una realtà multidisciplinare, infatti, le esigenze variano enormemente da un dipartimento all’altro, rendendo difficile individuare una **soluzione unica** che vada bene per tutti. Occorre quindi mappare i bisogni specifici e, probabilmente, adottare un insieme di strumenti diversi. Inoltre, Sterpi sottolinea che il processo di adozione non è affatto **a costo zero** o automatico: introdurre l’IA in studio richiede investimenti di tempo e risorse. Ad esempio, *il training è necessario*: l’idea che un’interfaccia in linguaggio naturale (come ChatGPT) elimini del tutto la necessità di formazione è un’illusione. Anche se l’IA si presenta come “facile da usare” perché capisce il linguaggio umano, gli avvocati devono comunque imparare a strutturare efficacemente le richieste (i *prompt*) e a interagire con la macchina per ottenere risultati utili, oltre che verificarne la correttezza (v. *infra*). Si tratta di una competenza nuova che va coltivata, richiedendo un investimento di tempo significativo per il personale dello studio. Un altro aspetto pratico è **il problema delle licenze software**: non è affatto detto – avverte Sterpi – che si debba acquistare una licenza d’uso per ogni singolo avvocato in studio. Molti strumenti, infatti, vengono usati sporadicamente o da piccoli gruppi. Sterpi cita un esempio illuminante, ovvero di un software che sembrava a molti irrinunciabile, ma poi il *93% delle licenze acquistate non è mai stato utilizzato* nell’arco di un anno. La lezione è che la strategia di adozione deve essere **mirata e graduale**: meglio assegnare le licenze inizialmente solo a gruppi pilota o per specifiche esigenze, ed eventualmente estenderle in seguito se l’utilizzo reale lo giustifica. Questo approccio prudente evita sprechi e consente di misurare il reale valore aggiunto degli strumenti prima di un *rollout* completo.

In terzo luogo, Sterpi passa in rassegna i **rischi connessi all’uso dell’IA**, lanciando un monito chiaro per la professione forense. L’entusiasmo per le nuove tecnologie, infatti, non deve far dimenticare che il loro impiego è *estremamente delicato*, specialmente quando in gioco ci sono dati sensibili o decisioni legali importanti. Tra i rischi principali, Sterpi evidenzia innanzitutto **quello per la confidenzialità e la privacy**: utilizzare versioni gratuite o non certificate di strumenti di IA può esporre lo studio a pericoli gravissimi, perché i dati caricati (atti giudiziari, contratti, informazioni riservate dei clienti) potrebbero essere utilizzati dai fornitori per addestrare i loro modelli o, peggio, potrebbero essere soggetti a falle di sicurezza e diffusi involontariamente. Le conseguenze, in caso di violazione, sarebbero disastrose sia sul piano **reputazionale**

(perdita di fiducia dei clienti), sia su quello **risarcitorio** (possibili azioni legali), sia infine su quello **deontologico** (violazione del segreto professionale e delle norme deontologiche sulla riservatezza). Un secondo rischio è definito la “**trappola del verosimile**”: gli strumenti di IA generativa producono testi che suonano sempre grammaticalmente corretti, coerenti e plausibili, con un tono sicuro – in pratica, sembrano elaborati da un essere umano competente. Questo può indurre facilmente in errore, perché *il fatto che un output sia ben formulato non significa affatto che sia vero!* Sterpi racconta a tal proposito di aver condotto un test con una piattaforma IA di fama mondiale, la quale **inventava di sana pianta** le fonti giurisprudenziali che citava a supporto delle sue risposte. In pratica, l’IA forniva una risposta ben scritta e citava sentenze e articoli di legge inesistenti o irrilevanti pur di corroborare la risposta, sfruttando la fiducia dell’utente meno attento. Questo episodio evidenzia un pericolo: l’avvocato potrebbe essere tentato di prendere per buono un *output* solo perché è ben scritto e *sembra* corretto, trascurando di verificarlo.

Proprio per scongiurare tali pericoli, Sterpi ribadisce che **l’obbligo di verifica umana è assoluto** e inderogabile. L’avvocato ha il dovere professionale (oltre che etico) di **verificare ogni singola informazione prodotta dall’IA**. Nulla può essere accettato acriticamente: ogni riferimento normativo, ogni citazione e ogni conclusione suggerita dalla macchina vanno passati al vaglio dal professionista. Questo principio, sottolinea Sterpi, vale per tutti ma è particolarmente cruciale per i collaboratori più giovani, che hanno meno esperienza e strumenti critici per individuare errori o incongruenze nei risultati generati. In altre parole, paradossalmente i *junior* potrebbero essere più esposti al rischio di fidarsi dell’IA senza accorgersi delle sue falsità, ed è dunque necessario sensibilizzarli e formarli adeguatamente. Un ultimo rischio menzionato riguarda la **violazione di diritti di terzi** attraverso l’uso dell’IA: i contenuti generati potrebbero infrangere copyright (ad esempio riproducendo testi protetti senza autorizzazione), oppure contenere affermazioni diffamatorie, o ancora violare altri diritti di proprietà intellettuale altrui. Alcuni fornitori di servizi AI – nota Sterpi – offrono contrattualmente una sorta di **manleva legale** (ad esempio Microsoft promette di tenere indenne l’utente professionale da azioni di terzi per violazione di copyright imputabile all’uso dei suoi strumenti). Tuttavia, anche queste garanzie hanno un perimetro limitato: possono magari coprire l’aspetto economico di una causa, ma non **riparano il danno reputazionale** nel caso in cui lo studio legale finisca sui giornali per un uso improprio dell’IA, né tantomeno eliminano le possibili conseguenze disciplinari in ambito deontologico. Insomma, *prevenire è meglio che curare*: occorre adottare misure di controllo interne (policy, formazione, doppie verifiche) per evitare di incappare in questi scogli.

Sterpi conclude il suo intervento con una riflessione sull’**impatto dell’IA sulle diverse generazioni di avvocati**. Questo tema si ricollega idealmente alla metafora di Vaciago e aggiunge un ulteriore tassello: l’IA, in studio, non influisce in modo uniforme su tutti i professionisti. In particolare, osserva Sterpi, l’IA rappresenta un “*grande competitore*” per i giovani avvocati. Ciò può suonare controintuitivo – si tende a pensare che i giovani

siano più avvantaggiati con le nuove tecnologie – ma in ambito legale l’IA automatizza proprio molte delle attività tipicamente delegate ai praticanti e ai neo-associati: la ricerca di precedenti, la traduzione di testi, la predisposizione di prime bozze di atti, ecc.. Mansioni un tempo formative per i junior oggi possono essere svolte in parte da un algoritmo, riducendo quelle che erano opportunità per “farsi le ossa” sul campo. Di converso, paradossalmente l’IA **favorisce i professionisti senior**: sono questi ultimi, infatti, a possedere la profonda conoscenza del contesto giuridico e l’esperienza necessarie per porre le *domande giuste* alla macchina e valutare criticamente la qualità delle risposte fornite. In altri termini, l’IA potenzia chi ha già **conoscenza di dominio** (perché sa guidare meglio la tecnologia e capisce subito se un *output* è sensato oppure no), mentre rischia di sostituire alcune attività “di apprendistato” dei più giovani. Questa considerazione apre una questione importante per il futuro della professione: come conciliare l’uso efficiente dell’IA con la formazione dei nuovi avvocati? La chiave, suggerisce Sterpi, potrebbe stare in un utilizzo *complementare* dell’IA, che liberi tempo ai giovani per dedicarsi ad attività a maggior valore aggiunto (strategie, approfondimenti) e consenta loro al contempo di acquisire capacità nell’uso degli strumenti tecnologici stessi – competenza destinata a diventare parte integrante del bagaglio professionale dell’avvocato moderno.

(Connessione tematica: l’intervento di Sterpi prende gli spunti lanciati da Vaciago – opportunità per piccoli e grandi, necessità di cautela e consapevolezza – e li traduce in linee guida pratiche per l’adozione dell’IA in studio. Ci introduce inoltre ai temi della scelta tecnica (piattaforme globali vs locali) e della formazione interna, che verranno ripresi più dettagliatamente dai relatori successivi, Gilardi e Lonato. Anche la riflessione finale sull’impatto generazionale avrà un’eco nell’intervento di Tuninetti.)

4. Make or Buy? Modelli, Costi e Strumenti per una Scelta Consapevole - Dott. Gianluca Gilardi

Dopo aver considerato strategie generali e criticità, la lezione si è addentrata maggiormente negli **aspetti tecnico-operativi** con l’intervento dell’Avv. **Gianluca Gilardi**. Con un taglio molto pragmatico, Gilardi ha affrontato il classico dilemma che ogni studio si trova di fronte quando decide di innovare i propri processi con l’IA: *meglio acquistare una soluzione pronta all’uso sul mercato, oppure svilupparne una “su misura” internamente?* (il dilemma noto come “*make or buy*”). Inoltre, il suo contributo ha fornito criteri per valutare in modo consapevole i vari **modelli di IA** disponibili e suggerimenti su come testarli per capire quale sia il più adatto alle proprie esigenze. È un intervento che, idealmente, si collega alle indicazioni di Sterpi sulla scelta degli strumenti, approfondendo la questione dal punto di vista tecnico e decisionale.

Il primo punto affrontato da Gilardi è proprio la “**grande decisione**” **make or buy**. In sostanza, quando uno studio vuole dotarsi di strumenti di intelligenza artificiale, ha due strade:

- **Buy (comprare):** ovvero **acquistare** una piattaforma già pronta (*“a scaffale”*). Il vantaggio di questa scelta sta nella semplicità e rapidità di implementazione: si compra un prodotto fatto e finito, spesso *user-friendly* e subito utilizzabile. Di contro, ci sono alcuni svantaggi: in primis il **vendor lock-in**, cioè il rischio di legarsi a doppio filo a un unico fornitore e alla sua tecnologia; inoltre, si ha poca o nessuna **capacità di controllo** sui parametri interni del modello (non si può sapere né modificare come funziona esattamente l'IA “sotto il cofano”).
- **Make (creare):** ossia **sviluppare internamente** la propria soluzione, costruendo un'interfaccia o applicazione personalizzata che si appoggia ai grandi modelli linguistici disponibili tramite API. Fino a poco tempo fa questa opzione era appannaggio solo di chi poteva investire cifre enormi in ricerca e sviluppo; oggi, nota Gilardi, è diventata *accessibile anche a piccoli studi* con investimenti relativamente modesti (dell'ordine di poche migliaia di euro). L'approccio “make” offre potenzialmente **enormi vantaggi** rispetto al comprare soluzioni preconfezionate, in particolare per quanto riguarda:
 - **Flessibilità** – chi sviluppa in proprio può scegliere caso per caso **il modello di IA più appropriato per ogni singolo task** da svolgere. Ad esempio, si può decidere di usare un modello molto potente per compiti complessi e uno più leggero per attività semplici, ottimizzando così risultati e costi.
 - **Controllo** – realizzando la *propria* applicazione, si possono impostare direttamente certi parametri del modello (*gli “iperparametri”*), come ad esempio la *“temperatura”* che regola il livello di creatività/randomicità delle risposte. Questo tipo di fine-tuning è impossibile con le piattaforme commerciali chiuse, mentre diventa fattibile se si ha accesso diretto al modello tramite API. Il maggiore controllo consente di adattare il comportamento dell'IA alle proprie esigenze specifiche (ad esempio, rendendo le risposte più rigorose e aderenti al testo oppure più creative a seconda dei casi).
 - **Costi variabili** – invece di pagare licenze fisse, nell'approccio “make” si **paga solo per l'effettivo utilizzo** del modello (tipicamente calcolato in base ai token di testo processati, cioè al volume di dati elaborato¹), eliminando così i costi fissi e pagando in proporzione al beneficio ottenuto. Questo modello *pay-per-use* può risultare molto efficiente economicamente, specie se l'uso dell'IA non è costante ma concentrato in certi periodi o per specifici progetti.

Gilardi passa poi a sottolineare che **non tutti i modelli di IA sono uguali**, nemmeno all'interno dell'offerta di un singolo provider. Prendendo ad esempio OpenAI (che offre API per diversi modelli, da GPT-3.5 a GPT-4, ecc.), evidenzia come i modelli differiscano **enormemente** per caratteristiche e prestazioni. Alcune differenze chiave da considerare sono:

- **Capacità di “ragionamento”:** certi modelli di ultima generazione sono in grado di svolgere forme di *reasoning*, ossia di inferire e dedurre con una parvenza di logica, mentre altri – specialmente i modelli meno complessi – si limitano a produrre testo statisticamente coerente senza vero “ragionamento”. Per compiti semplici, osserva Gilardi, **non serve un modello dotato di avanzate capacità di**

- reasoning** (che sarebbe più costoso); modelli più basilari possono bastare e avanzare.
- **Finestra di contesto:** ogni modello ha un limite sulla quantità di testo che può essere introdotta in un singolo prompt (la cosiddetta *context window*). Questo limite varia enormemente: alcuni modelli gestiscono poche migliaia di token (parole), altri centinaia di migliaia, e gli ultimissimi arrivano a milioni. Gilardi cita il caso di *Llama 4* (un modello open source sviluppato da Meta) che promette una finestra di contesto fino a **10 milioni di token**. Una finestra ampia consente, ad esempio, di dare in pasto al modello documenti molto corposi o tanti documenti insieme, facendogli “leggere” un intero fascicolo prima di porre domande. Modelli con finestra ridotta richiederanno invece di spezzare l’input o sintetizzarlo in parte.
 - **Costo di utilizzo:** i costi (specie nei modelli accessibili via API) sono spesso calcolati “a consumo” in base ai token elaborati, e **variano significativamente da modello a modello**. Un modello più potente e avanzato di solito costa di più per ogni 1.000 token processati rispetto a uno meno performante. Bisogna quindi tenere presente il *trade-off* tra costo e prestazione: utilizzare sempre e solo il modello migliore può risultare oneroso se lo si impiega anche per mansioni in cui basterebbe un modello economico.

Alla luce di questa varietà di offerta, Gilardi fornisce infine alcuni consigli pratici su **come scegliere e testare i modelli** di IA prima di adottarli. Primo, suggerisce di utilizzare **piattaforme di benchmarking** specializzate, come ad esempio *OpenRouter*, che aggregano centinaia di modelli diversi e periodicamente stilano classifiche dei migliori per ciascun tipo di compito o settore (incluso quello legale). Queste piattaforme consentono di avere un quadro aggiornato delle prestazioni relative dei modelli senza doverli provare uno a uno manualmente. Secondo, Gilardi consiglia di **preparare internamente un set standard di domande/test** da sottoporre a vari modelli, in modo da confrontarne i risultati in condizioni controllate. Ad esempio, lo studio potrebbe predisporre una lista di **10 quesiti tipo**, toccando attività differenti (una richiesta di redigere un contratto, una di tradurre una clausola, una di riassumere una sentenza, ecc.) e poi provare a farli svolgere a più modelli di IA, confrontando le performance. Questo metodo empirico permette di individuare in concreto quale modello si comporta meglio sui problemi che allo studio interessano davvero, fornendo un **parametro oggettivo di confronto** tra le alternative. In sintesi, il messaggio di Gilardi è che la scelta del modello (o dei modelli) di IA non va fatta *sulla carta* o basandosi solo sulla fama del produttore, ma testando con mano le soluzioni sul campo, perché ciò che è migliore in generale potrebbe non esserlo nello specifico contesto del nostro studio.

5. L'Esperienza di uno Studio Internazionale: Contesto, Casi d'Uso e Giurisprudenza - Avv. Andrea Tuninetti

L’Avv. **Andrea Tuninetti** ha portato in aula la prospettiva di un grande **studio legale internazionale**, condividendo sia l’esperienza pratica di adozione dell’IA all’interno del suo studio, sia il quadro di riferimento normativo e giurisprudenziale più ampio in cui queste tecnologie si inseriscono. In un certo senso, Tuninetti ha allargato lo sguardo:

dopo gli approfondimenti su *come* scegliere e implementare l'IA, ha contestualizzato *perché* questo tema sia diventato così centrale, descrivendo il “terremoto” regolatorio in atto e come lo studio ha iniziato a sfruttare l'IA in concreto, mantenendo però un occhio vigile sugli sviluppi del diritto. Il suo intervento si articola in quattro parti: il contesto normativo europeo (paragonato a uno tsunami), alcuni casi d'uso interni allo studio, un'osservazione sull'adozione dell'IA da parte di giovani e senior, e una panoramica di giurisprudenza italiana recente in materia di algoritmi.

Per prima cosa, Tuninetti descrive **il contesto normativo** in cui si colloca l'avvento dell'IA, paragonandolo a un vero e proprio “*tsunami*” di *regolamentazione digitale*. Sottolinea che l'ondata di nuove normative sull'intelligenza artificiale *non è isolata*, ma rientra in uno sforzo più ampio del legislatore europeo per regolamentare le tecnologie digitali a tutto tondo. I numeri sono impressionanti: **il numero di provvedimenti normativi in ambito digitale che entreranno in vigore nei prossimi due anni supera quello dei precedenti quarant'anni**. Questa pioggia di nuove regole – che spazia dalla protezione dei dati (GDPR) ai servizi digitali (DSA/DMA), fino all'AI Act in arrivo – genera inevitabilmente un senso di smarrimento nelle aziende. Molte imprese faticano a comprendere e a stare al passo con obblighi che si moltiplicano e si aggiornano continuamente. Ciò, però, apre anche **nuove frontiere per la consulenza legale**: il ruolo dell'avvocato in questo scenario non si limita più a verificare la *compliance* (cioè il rispetto formale delle norme), ma si estende a un livello più strategico, aiutando i clienti a **ridefinire contratti, processi e strategie aziendali** in funzione di questo mutato contesto. In altre parole, il giurista diventa una guida nel *governare l'innovazione*, assicurando che l'adozione dell'IA e di altre tecnologie avvenga in modo conforme ma anche efficiente. Tuninetti, con questa premessa, fa capire che parlare di IA per avvocati non significa solo occuparsi del gadget tecnologico, ma inserirlo in un più ampio scenario di trasformazione normativa e di mercato.

Passando dal *macro* al *micro*, Tuninetti ha poi illustrato alcuni **casi d'uso concreti dell'IA** all'interno del suo studio legale internazionale, evidenziando come l'adozione sia avvenuta per gradi. L'approccio scelto è **graduale**, iniziando da compiti a **basso rischio** e sempre mantenendo una supervisione umana attenta e consapevole (“*mindful*”). Tra gli esempi pratici citati vi è l'uso dell'IA per la **sintesi di documenti**: in studio si impiegano strumenti di IA per riassumere automaticamente sentenze o altri documenti lunghi, così da fornire ai professionisti un *executive summary* immediato e avere rapidamente sott'occhio i punti principali. Questo non solo fa risparmiare tempo, ma aiuta l'avvocato a orientarsi più velocemente in materiali complessi, su cui poi comunque condurrà un esame approfondito. Un altro esempio è l'**analisi di più documenti in parallelo**: caricando in un sistema di IA un insieme di contratti o testi, è possibile porre domande specifiche (come si farebbe in un esame di comprensione del testo) oppure chiedere al sistema di individuare la presenza di determinate clausole di interesse (ad esempio, “trova tutte le clausole di riservatezza in questo pacchetto di contratti”). In questo modo l'IA diventa un assistente che *setaccia* una grande mole di documenti per estrarre le informazioni rilevanti, compito che manualmente richiederebbe ore o giorni. È

fondamentale notare – e Tuninetti lo sottolinea – che in tutti questi casi **l’IA non sostituisce l’analisi dell’avvocato, ma la supporta nella fase iniziale**. Il professionista infatti utilizza l’output dell’IA come base di partenza, che va poi *validata* e approfondita. Ad esempio, un riassunto automatico di una sentenza servirà a farsi un’idea generale, ma prima di citarla in un atto il legale dovrà comunque leggerla integralmente; analogamente, se l’IA individua clausole sospette in una serie di contratti, sarà il legale a esaminarle e valutarne l’effetto. L’adozione in studio quindi è guidata da **prudenza**: si sfrutta l’IA per velocizzare i processi interni, ma sempre con la supervisione e il controllo critico umano.

Un aspetto interessante emerso dall’esperienza di Tuninetti riguarda il **diverso approccio tra generazioni** all’interno dello studio nell’abbracciare questi nuovi strumenti. Contrariamente a quanto si potrebbe pensare, inizialmente c’è stata una maggiore adozione dell’IA da parte dei **professionisti più senior** rispetto ai colleghi più giovani. La spiegazione fornita è che molti giovani avvocati, meno sicuri su una materia ancora nuova, hanno preferito studiarla a fondo prima di utilizzarla e soprattutto **non vogliono mettere a repentaglio la propria reputazione** commettendo passi falsi con strumenti che non padroneggiano. I senior invece, forti di esperienza e posizione, si sono sentiti più liberi di sperimentare e hanno colto subito il potenziale di delegare all’IA alcune attività ripetitive, mantenendo però il controllo sul risultato. In parte questo rispecchia quanto osservato anche da Sterpi: i senior dispongono della “cassetta degli attrezzi” per capire e indirizzare meglio l’IA, mentre i junior temono (comprensibilmente) che affidarsi a uno strumento di cui non conoscono bene i limiti possa esporli a errori e critiche. Col tempo, naturalmente, anche i più giovani inizieranno a impiegare con disinvoltura queste tecnologie, man mano che ne comprenderanno il funzionamento e soprattutto vedranno che lo **studio ne incentiva l’uso consapevole**. Ma l’annotazione di Tuninetti fa riflettere: l’innovazione non è solo una questione tecnica, è anche un fatto culturale all’interno degli studi legali, dove convinzioni, timori e abitudini possono influenzare la velocità del cambiamento.

Infine, Tuninetti ha fornito una breve ma significativa **panoramica sulla giurisprudenza italiana** più recente in materia di algoritmi e decisori automatici, evidenziando come i nostri tribunali abbiano già affrontato casi pionieristici su questi temi. In particolare, nota come i giudici italiani siano stati **all’avanguardia in Europa** nell’affrontare questioni relative all’uso di algoritmi, ancora prima dell’entrata in vigore dell’AI Act (la futura regolamentazione europea sull’intelligenza artificiale)¹, basandosi sui principi di diritto esistenti. Tuninetti cita due esempi emblematici.

Il primo è il **caso Deliveroo**. Nel 2020 il Tribunale di Bologna si è trovato a giudicare il sistema algoritmico di gestione dei rider (i fattorini che consegnano cibo per la piattaforma Deliveroo). L’algoritmo assegnava punteggi ai rider anche in base alle assenze dal lavoro, senza distinguere tra assenze giustificate (ad esempio malattia o adesione a uno sciopero) e assenze ingiustificate. Questo meccanismo, di fatto, penalizzava i lavoratori che esercitavano diritti fondamentali (come appunto il diritto di

sciopero o la malattia) equiparandoli a “assenteisti”. Il tribunale ha ritenuto **discriminatorio** tale sistema ed è intervenuto per bloccarlo, applicando principi di diritto del lavoro e di non discriminazione già presenti nel nostro ordinamento. La sentenza Deliveroo, dunque, pur in assenza di una normativa specifica sugli algoritmi decisionali, ha saputo utilizzare i principi generali per censurare un uso distorto della tecnologia che ledeva i diritti dei lavoratori.

Il secondo caso menzionato è il **caso MeValuate**, affrontato addirittura dalla Corte di Cassazione. *MeValuate* era un sito web che forniva un servizio di **ranking reputazionale** delle persone, presumibilmente attraverso un algoritmo che aggregava varie informazioni disponibili online. Questo sito è stato sanzionato dalle autorità (il Garante Privacy) per **manca di trasparenza** sul funzionamento dell'algoritmo e sul trattamento dei dati personali degli interessati. La vicenda è arrivata fino in Cassazione, la quale ha confermato le criticità relative al trattamento dei dati personali, e affermando che ogni individuo ha **il diritto di conoscere la logica** in base alla quale vengono trattati i propri dati personali, soprattutto se da ciò deriva un giudizio sul proprio conto. In pratica, la Corte Suprema ha detto: se un algoritmo ti assegna un punteggio o un'etichetta che può influenzare la tua vita (fosse anche solo la reputazione online), hai il diritto di sapere in base a quali criteri e dati quel punteggio è stato calcolato. Anche questa pronuncia è significativa perché anticipa concetti che saranno al centro dell'AI Act (come il diritto alla spiegazione delle decisioni automatizzate) usando però strumenti giuridici già disponibili (in questo caso le norme del GDPR sulla trasparenza e i diritti degli interessati).

Gli esempi Deliveroo e MeValuate dimostrano come il diritto del lavoro e della privacy possano già oggi fronteggiare alcune sfide poste dall'IA, ma evidenziano anche la frammentazione di approcci che l'AI Act cercherà di sistematizzare. Per gli avvocati è importante conoscere questi precedenti perché forniscono indicazioni su come i giudici potranno valutare future controversie riguardanti algoritmi e IA.

*(Connessione tematica: l'intervento di Tuninetti ha ampliato l'orizzonte della discussione, inserendo l'adozione dell'IA in uno **scenario regolatorio in fermento** e mostrando esempi concreti di utilizzo “controllato” dell'IA in un grande studio. Ciò collega i temi tecnici affrontati prima con la **realtà quotidiana della professione legale** e con la cornice giuridica esterna che la circonda. Inoltre, la notazione sull'adozione differenziale tra generazioni riprende il filo lanciato da Sterpi, mentre i casi giurisprudenziali prefigurano alcune questioni di principio che rimandano all'importanza di mantenere il fattore umano e i diritti fondamentali al centro, proprio come evidenziato anche da Vaciago e dagli altri relatori.)*

6. La Filosofia Augmented vs. Automation: un Percorso Strutturato per l'Adozione dell'IA - Avv. Benedetto Lonato

A chiudere la ricca sessione del 12 giugno è stato l'intervento dell'Avv. **Benedetto Lonato**, che ha illustrato il percorso metodologico seguito dal **suo studio legale** per abbracciare l'innovazione dell'intelligenza artificiale in modo organico e strutturato. La particolarità di questo contributo è che Lonato ha condiviso una sorta di *case study*, mostrando non solo i principi guida adottati (la "filosofia" dello studio rispetto all'IA), ma anche le scelte pratiche fatte: come hanno organizzato il team interno, quali strumenti specifici hanno introdotto e con quali risultati operativi. In un certo senso, è la concretizzazione di molti concetti discussi dagli altri relatori – dall'importanza di un approccio strategico e cauto (Vaciago, Sterpi) alla necessità di scegliere gli strumenti giusti (Sterpi, Gilardi) e di mantenere l'avvocato al centro del processo (Tuninetti).

Il cuore della strategia di Lonato risiede nella filosofia di base adottata dallo studio, sintetizzata nel motto **"augmented vs. automation"**. Invece di cercare soluzioni che automatizzino completamente e **sostituiscano** intere fasi del lavoro legale, lo studio ha deciso di implementare strumenti che **aumentino le capacità del professionista** – in inglese si parla di *augmented lawyer*, l'avvocato potenziato. Ciò significa, in pratica, prediligere tecnologie che affianchino l'avvocato migliorandone l'efficienza e la qualità del lavoro, piuttosto che rimpiazzarlo in toto in qualche attività. L'obiettivo dichiarato è *migliorare la qualità* sia del risultato (un prodotto legale migliore) sia del tempo speso (eliminando compiti ripetitivi o a basso valore aggiunto), **mantenendo però l'avvocato al centro del processo** decisionale e di controllo. Questa presa di posizione è importante anche culturalmente: allevia il timore che l'IA "tolga il lavoro" e la presenta invece come uno strumento per **liberare tempo** da dedicare a ciò che solo l'uomo sa fare (strategie, creatività giuridica, rapporto col cliente). La filosofia "augmented" permea dunque tutto l'approccio all'innovazione di questo studio.

Un secondo passo fondamentale nel percorso di Lonato è stato definire un efficace **processo di selezione e implementazione** interna delle soluzioni di IA. Inizialmente, racconta, avevano provato un approccio diffuso: far testare le nuove tecnologie a team diversi, distribuiti per dipartimento, sperando forse che l'innovazione emergesse "dal basso" spontaneamente. Questo metodo però *si è rivelato fallimentare*, probabilmente perché senza un coordinamento centrale molti test si disperdevano, mancava condivisione di conoscenze, e alcuni reparti rimanevano indietro. La svolta vincente è stata allora quella di creare un **team dedicato e centralizzato di "Prompt Lawyers"**. Si tratta di un gruppo di **17-18 professionisti** interni allo studio, appartenenti a diverse practice ma con una passione o inclinazione per la tecnologia, ai quali è stato assegnato ufficialmente il compito di dedicare una parte del loro tempo (circa il 15% dell'orario) a testare strumenti di IA, selezionare quelli più promettenti e diventare essi stessi utenti esperti che poi fungano da tutor/evangelist per il resto dello studio. Questo team funziona un po' da *"testa d'ariete"* per l'intero studio: assorbe l'onere della sperimentazione, accumula competenza sugli strumenti e quindi riduce la curva di

apprendimento per tutti gli altri, facendo da riferimento interno. L'idea dei *prompt lawyers* – termine ormai entrato nel lessico del legal tech – è proprio quella di avere in casa specialisti capaci di dialogare con i modelli di IA (scrivere *prompt* efficaci, interpretare i risultati, escogitare usi innovativi) e di trasferire queste abilità agli altri avvocati man mano che gli strumenti verranno adottati più diffusamente. Questa mossa organizzativa riprende quanto suggerito anche dall'introduzione del corso: creare *team dedicati* all'innovazione è un modello di governance interna che può fare la differenza rispetto a un approccio dispersivo. E l'esperienza di Lonato conferma che **un nucleo interno focalizzato** ha accelerato e reso più coerente l'adozione dell'IA nello studio.

Passando agli strumenti concreti, Lonato illustra poi la **scelta degli strumenti tecnologici** operata: il suo studio ha adottato un vero *ecosistema integrato* di soluzioni, composto in totale da **sei strumenti principali**, combinando applicazioni “orizzontali” (utili trasversalmente a tutti) e soluzioni “verticali” specialistiche. Ecco, in sintesi, le categorie di tool implementati (con qualche esempio per ciascuna):

- **Piattaforma interna su cloud privato (Microsoft Azure):** un ambiente di lavoro proprietario, basato sull'infrastruttura cloud Azure già utilizzata dallo studio, in cui far girare modelli di IA e conservare i dati in modo sicuro. Questa scelta garantisce **massima sicurezza e controllo** sui dati trattati, poiché tutto rimane nel cloud dello studio, rispettando i requisiti di riservatezza e privacy.
- **Strumenti di traduzione automatica (es. DeepL Pro):** applicazioni dedicate alla **traduzione professionale** dei testi. La traduzione giuridica è stata uno dei primi ambiti in cui l'impatto dell'IA è stato massiccio: utilizzare un servizio come DeepL (nella versione Pro, che assicura riservatezza dei dati tradotti) permette di ottenere traduzioni di alta qualità in pochi secondi, laddove prima erano necessarie ore di lavoro umano o costose deleghe a traduttori esterni. Non a caso, nota Lonato, questo è uno degli strumenti la cui adozione è stata più **generalizzata** tra i professionisti – chiunque in studio può trarne beneficio immediato ed evidente.
- **Assistenti virtuali per Q&A legale (es. Lexroom):** si tratta di **consulenti legali IA** in grado di rispondere a quesiti giuridici comuni, fornendo riferimenti normativi o indicazioni di massima. In pratica sono chatbot addestrati sul corpus delle leggi e (in parte) delle sentenze italiane, che possono aiutare l'avvocato nelle ricerche preliminari. Nel caso dello studio di Lonato, l'uso iniziale di questi strumenti è stato **limitato ai professionisti senior** – cioè a chi ha l'esperienza per valutare criticamente le risposte dell'IA e contestualizzarle – mentre è stato **precluso ai praticanti**. Questa decisione è perfettamente in linea con la filosofia “augmented”: i praticanti devono prima formarsi sulle basi tradizionali e non rischiare di appoggiarsi acriticamente a una scorciatoia tecnologica; i senior invece, ben consapevoli dei limiti, possono usare l'IA come *sparring partner* intellettuale, senza compromettere la formazione dei junior.
- **Assistenti legali avanzati (es. Harvey, Lidia):** sono definiti la *vera rivoluzione*. Si tratta di strumenti più sofisticati, capaci di **analizzare migliaia di documenti in parallelo** e di generare bozze e spunti di lavoro. Ad esempio, *Harvey* è un assistente basato su GPT-4 progettato per studi legali internazionali, mentre

Lidia è una piattaforma italiana emergente; entrambi sono in grado di ingerire fino a 10.000 documenti contemporaneamente (non più uno dopo l'altro, ma appunto in parallelo) e poi rispondere a domande complesse o produrre sintesi integrate. Questo abilita casi d'uso prima impensabili, come condurre analisi trasversali su interi *dataset* di contratti, o redigere una prima bozza di due diligence report mettendo insieme informazioni sparse in centinaia di file. È in questo campo che si intravede come l'IA possa far fare un vero *salto di qualità* all'efficienza dello studio, fermo restando che l'output va sempre revisionato dall'avvocato.

Dopo aver elencato la “cassetta degli attrezzi” costruita nello studio, Lonato ha condiviso alcuni **casi d'uso operativi** per mostrare in concreto i frutti di questo ecosistema tecnologico. Un primo esempio riguarda la fase di **closing** nelle operazioni di M&A (*mergers & acquisitions*): tipicamente, dopo la negoziazione di un complesso contratto di acquisizione, bisogna preparare una *closing checklist*, ossia un elenco dettagliato e cronologico di tutte le attività ed adempimenti da completare affinché l'accordo vada a buon fine (dai documenti da firmare, alle comunicazioni da inviare, ecc.). Utilizzando l'IA, lo studio fa analizzare automaticamente all'assistente legale l'intero set di contratti e documenti di una transazione, e ottiene come output un **riassunto organizzato e cronologico di tutte le attività necessarie per l'esecuzione dell'accordo** – in pratica, una *bozza* di closing checklist già pronta da verificare. Ciò fa risparmiare ore di lavoro manuale nel setacciare ogni clausola alla ricerca degli obblighi post-firma: l'avvocato può concentrarsi nel controllare e rifinire la checklist generata invece di crearla da zero. Un secondo caso d'uso citato è la **due diligence su documenti pubblici**. Immaginiamo di dover esaminare centinaia di prospetti informativi finanziari per individuare determinati *trend* o rischi ricorrenti (ad esempio, l'eventuale menzione di crisi macroeconomiche, clausole particolari di uscita, etc.). Invece di leggere uno ad uno tutti i documenti, l'IA permette di analizzarli in massa, individuando pattern e segnalandoci le informazioni più rilevanti che ricorrono frequentemente. L'avvocato potrà poi approfondire solo i passaggi segnalati, ottenendo una panoramica più ampia in molto meno tempo. Questo esempio dimostra come l'IA possa potenziare la capacità di *knowledge management* all'interno di attività tipicamente gravose e “muscolari” come le due diligence, facendo emergere conoscenza dagli archivi documentali con maggiore rapidità.

A conclusione del suo intervento, Lonato evidenzia **una lacuna attuale** nell'offerta di IA per il settore legale: il *grande assente* è la **dottrina giuridica**, ovvero tutta la produzione intellettuale di libri e articoli scritti da giuristi ed accademici. Per ragioni di copyright, infatti, la maggior parte delle soluzioni di IA (soprattutto quelle commerciali) non hanno accesso ai contenuti dottrinali pubblicati dagli editori giuridici tradizionali. Ciò significa che, se un avvocato pone all'IA una questione che richiederebbe l'elaborazione di teorie o orientamenti tratti da libri e riviste scientifiche di diritto, oggi l'IA non è in grado di fornirli perché il suo *training* esclude quelle fonti coperte da diritti d'autore. Questo è un limite significativo, perché la **dottrina** rappresenta da sempre un pilastro nell'interpretazione del diritto e nella formazione del giurista. Lonato indica quindi **il prossimo fronte di sviluppo**: l'integrazione degli strumenti di IA con le grandi banche

dati dottrinali degli editori storici. In pratica, si auspica (e si prevede) che in futuro si possano stipulare accordi o sviluppare soluzioni tali per cui i modelli di IA possano consultare in modo controllato anche i testi dottrinali ufficiali, ampliando enormemente la profondità delle analisi possibili. Questo, ovviamente, dovrà essere fatto nel rispetto del copyright e magari generando nuove forme di licenza d'uso specifiche; ma è un'evoluzione naturale se si vuole che l'IA diventi un vero assistente a 360 gradi per il giurista. Si chiude così il cerchio: la tecnologia ha fatto passi da gigante, gli studi legali più avanzati la stanno già incastonando nei loro processi, e il mondo delle pubblicazioni giuridiche dovrà a sua volta evolvere per *dialogare* con queste nuove intelligenze, in un ecosistema integrato di sapere legale umano e artificiale.

7. Conclusioni – Districarsi in un futuro sempre più tecnologicamente ampio

La lezione del 12 giugno 2025 ha dipinto un quadro sfaccettato e altamente istruttivo del rapporto tra intelligenza artificiale e professione legale. Seguendo il filo cronologico degli interventi, siamo passati da una **visione strategica generale** – in cui l'IA emerge come opportunità di crescita e di riequilibrio competitivo per studi di ogni dimensione (Vaciago) – attraverso una **valutazione pragmatica** delle sfide di mercato, scelte operative e rischi da governare (Sterpi), quindi per un **approfondimento tecnico** su come decidere tra soluzioni preconfezionate o su misura e su come valutare i diversi modelli di IA (Gilardi), fino ad arrivare a **esperienze sul campo** sia di respiro internazionale (Tuninetti) sia di implementazione strutturata in uno studio italiano (Lonato). Ciò che emerge in filigrana da tutti questi contributi è una linea comune: **l'IA non è una bacchetta magica né un nemico da combattere**, ma uno strumento potente da integrare con intelligenza e senso critico nella pratica legale. Gli avvocati devono conoscerne le potenzialità e i limiti, sfruttarla per automatizzare compiti ripetitivi e accelerare il lavoro, ma sempre **mantenendo il controllo umano** sul processo e sul risultato finale.

In questo percorso di apprendimento collettivo, la sessione ha realmente fornito ai partecipanti quella *bussola* auspicata negli obiettivi formativi. Le “lezioni apprese” dai vari relatori possono infatti fungere da linee guida per qualsiasi studio legale che si affacci al mondo dell'IA: evitare l'ansia da *first adopter* e capire prima il proprio bisogno (Vaciago); studiare il mercato e non trascurare gli aspetti organizzativi interni, come la formazione e la gestione dei rischi (Sterpi); prendere decisioni informate sulle tecnologie da adottare, testandole e mantenendo flessibilità (Gilardi); inserire l'innovazione in un contesto di regole e con un approccio graduale e consapevole (Tuninetti); infine, adottare un metodo strutturato, con la giusta mentalità (“augmented” non “automated”) e investendo sulle persone oltre che sugli strumenti (Lonato).

Per un pubblico colto ma non specialista, il messaggio chiave è che l'intelligenza artificiale **può davvero trasformare in meglio il lavoro dell'avvocato**, rendendolo più efficiente e aprendo nuove possibilità, ma richiede di essere accompagnata da una trasformazione culturale e organizzativa. Non è solo questione di comprare il software

più avanzato: è questione di **ripensare processi**, formare professionisti alle *legal tech skills*, sperimentare con curiosità ma anche con cautela, e tenersi aggiornati su un panorama normativo in rapido mutamento. Solo così l'IA diventerà un *alleato* quotidiano nelle nostre scrivanie legali, e non un oggetto misterioso o – peggio – una fonte di problemi.

In definitiva, questa lezione ha mostrato come grandi studi e piccoli studi, senior e junior, tecnici e giuristi, possano collaborare per **sfruttare l'IA responsabilmente**. L'avvocato, figura chiave della tutela di diritti e interessi, può e deve restare **protagonista** anche nell'era dell'intelligenza artificiale: con umiltà nell'imparare i nuovi strumenti, ma anche con la fierezza di chi sa che il *valore umano* – il giudizio, l'etica, la creatività giuridica – resta insostituibile e anzi diventa ancor più centrale quando viene *potenziato* da strumenti digitali avanzati. In questo senso, per usare una metafora cara a Vaciago, l'IA è destinata a “rimescolare le carte” nel mondo legale, ma saranno sempre gli avvocati a giocare la partita, se sapranno far tesoro della bussola fornita da esperienze come quelle condivise in questa giornata.

Capitolo VII

Innovazione tecnologica nello Studio Legale: persone, processi e AI in azione

1. Introduzione

Nella lezione del 19 giugno 2025 del *Corso di Perfezionamento in Intelligenza Artificiale per Avvocati*, l'attenzione si è spostata dalla teoria normativa alla realtà operativa della gestione tecnologica e dell'innovazione nello studio legale. Diversi professionisti – dal responsabile IT di un grande studio, a consulenti di legal tech, fino ad avvocati impegnati in iniziative innovative – hanno condiviso esperienze eterogenee. Attraverso i loro interventi, sono state esplorate le sfide pratiche, organizzative e psicologiche che accompagnano l'adozione dell'Intelligenza Artificiale (IA) nel settore legale. I temi chiave emersi riguardano i bias cognitivi che frenano il cambiamento, l'importanza di un approccio strutturato basato sul trionomio **Persone, Processi, Tecnologia**, e metodologie concrete per mappare le esigenze, testare soluzioni e governare l'implementazione tecnologica. L'obiettivo formativo della sessione è stato chiaro: fornire ai partecipanti non solo una panoramica degli strumenti disponibili, ma soprattutto un metodo di lavoro e un approccio mentale per affrontare la trasformazione digitale. In sostanza, la tecnologia va trasformata da fonte di ansia a leva strategica per migliorare l'efficienza operativa, la qualità del lavoro e la soddisfazione del cliente.

Di seguito, seguendo l'ordine cronologico degli interventi dei relatori presenti, sviluppiamo una narrazione dei contributi di ciascun autore. Ogni sezione è dedicata a un relatore e al tema da lui affrontato, arricchendo i contenuti con esempi pratici, note esplicative ove utili e connessioni tematiche tra gli argomenti emersi.

2. L'impatto dell'IA generativa in un grande studio legale - Dott. Alejandro Perez.

Profilo e contesto: Alejandro Perez, manager con lunga esperienza in contesti strutturati (da Giuffrè a Thomson Reuters, fino allo studio Chiomenti), ha aperto la giornata offrendo una visione pragmatica di come l'ondata dell'IA generativa abbia impattato la realtà dei grandi studi legali. Il suo intervento ha messo in luce sia le sfide affrontate che le lezioni apprese nell'introdurre tecnologie di IA in contesti complessi, partendo dallo stato pregresso degli strumenti tecnologici fino ad arrivare alle prospettive future.

Il contesto prima dell'IA generativa: la mappa degli strumenti legal tech

Per prima cosa, Perez ha descritto l'ecosistema tecnologico dei grandi studi legali **prima** dell'avvento dell'IA generativa (indicativamente prima del 2022). Si trattava di una complessa "mappa" di strumenti specializzati, ciascuno dedicato a un compito specifico.

Ad esempio, erano già in uso software di **Document Automation** per creare bozze di documenti tramite questionari guidati, strumenti di **Document Review** basati su IA “estrattiva” per analizzare migliaia di documenti (tipicamente nelle due diligence), sistemi di **Document Management (DMS)** per l’archiviazione centralizzata degli atti, **Practice Management Software** per la gestione di contatti, fascicoli e calendari, oltre ad altri strumenti per **e-discovery**, ricerche giuridiche, e così via. L’insieme di questi tool formava il backbone tecnologico dello studio. Tuttavia, la sfida principale all’epoca – e ancora oggi – era l’integrazione fra sistemi così eterogenei. Far dialogare piattaforme diverse (spesso fornite da vendor differenti) richiedeva investimenti e competenze, e talvolta portava a inefficienze dovute alla mancanza di interoperabilità completa.

La pressione di ChatGPT: un unico tool per dominarli tutti?

Nel 2023, con l’arrivo di ChatGPT e degli strumenti di IA generativa rivolti al grande pubblico, i responsabili legaltech hanno dovuto affrontare una pressione crescente: si è diffusa l’idea, alimentata dall’entusiasmo mediatico, che un solo strumento potesse sostituire l’intero ecosistema di software specialistici costruito negli anni. Per chi, come Perez, si occupa di innovazione nei grandi studi legali, la prima sfida è stata riportare la discussione su un piano concreto, rafforzando la comunicazione interna e chiarendo che l’IA generativa può essere un alleato, ma non rimpiazza molte delle soluzioni verticali esistenti – come quelle per la ricerca giuridica, la traduzione o la gestione documentale – o quantomeno era necessario verificarlo attentamente. La gestione delle aspettative è diventata così centrale, per evitare sia disinvestimenti affrettati sia investimenti prematuri in tecnologie ancora immature per il contesto legale.

Il cambio di paradigma nelle competenze: il modello “sandwich”

La sfida forse più importante individuata da Perez è il **cambiamento nelle competenze** richieste all’avvocato nell’era dell’IA. Citando uno studio di Microsoft intitolato “Future of Work Report 2023”, egli ha illustrato il cosiddetto modello “**sandwich**” che descrive come si ridisegna il ruolo del professionista della conoscenza interagendo con l’IA. Il modello prevede tre fasi sovrapposte come in un panino:

- **Fase 1 – Input umano:** L’essere umano imposta il lavoro. In pratica, l’avvocato definisce il compito, fornisce il contesto necessario e scrive il prompt giusto per l’IA (ossia le istruzioni iniziali).
- **Fase 2 – Elaborazione automatica:** La macchina (il sistema di IA) esegue il lavoro operativo. Ad esempio, redige la prima bozza di un contratto o analizza un corpus di documenti generando un riepilogo.
- **Fase 3 – Verifica umana:** Infine, l’umano torna in gioco per verificare, valutare criticamente e adattare l’output prodotto dall’IA.

In questo schema a sandwich, le abilità richieste all’avvocato si spostano dalla *produzione diretta* del documento o dell’analisi, alla capacità di **gestire il processo**. Bisogna saper impostare bene il lavoro a monte (con domande precise, istruzioni

corrette, dati appropriati) e saper revisionare con occhio critico il risultato a valle. Perez ha messo in guardia i giuristi: usare l'IA non significa “fare di meno”, bensì fare *diversamente*. È necessario sviluppare competenze di *prompt design* (saper dialogare con l'IA in modo efficace) e di *critical review* (saper individuare errori o incongruenze nell'output generato). In sostanza, il professionista diventa un **manager** del lavoro svolto dall'IA, concentrandosi su strategia, controllo di qualità e giudizio finale, piuttosto che sulla stesura meccanica iniziale.

Il processo di adozione: dalla teoria alla pratica

Perez ha poi delineato un vero e proprio **processo operativo** per gestire l'ondata di nuove soluzioni tecnologiche all'interno dello studio, trasformando le idee in azioni concrete. Questo processo di adozione si articola in quattro fasi principali:

1. **Scouting**: si parte da un'analisi di mercato, ovvero dal monitorare e ricercare le varie soluzioni di IA e legal tech disponibili. In questa fase “esplorativa” si cerca di capire quali nuovi strumenti esistono, cosa offrono e se potrebbero adattarsi alle esigenze dello studio.
2. **Proof of Concept (PoC)**: dopo aver individuato una potenziale tecnologia interessante, è fondamentale testarla sul campo in scala ridotta. Realizzare una prova di concetto significa provare concretamente lo strumento su un caso d'uso reale dello studio, per valutarne l'efficacia. Perez ha enfatizzato l'importanza di passare dalla teoria alla pratica il prima possibile, per evitare la “*paralisi da analisi*”: se ci si perde in studi e comparazioni senza mai provare nulla, si rischia di non decidere mai.
3. **Negoziazione e sicurezza**: qualora il test dia esito positivo, si passa a una fase spesso lunga e delicata, cioè la verifica contrattuale e di sicurezza con il fornitore della tecnologia. In ambito legale, più che altrove, è cruciale analizzare i contratti di fornitura, le condizioni d'uso, le garanzie, nonché effettuare audit di sicurezza. Questa fase assicura che l'introduzione dello strumento sia conforme a standard legali, deontologici e di cybersecurity dello studio.
4. **Adozione vera e propria**: infine, se tutto va bene, si arriva all'adozione a regime della tecnologia nello studio. Paradossalmente, questa è la fase più complessa, perché comporta **formazione** del personale e, soprattutto, un cambio di mentalità. Spesso i vendor promettono strumenti “intuitivi” e facili da usare, ma Perez avverte che la realtà è diversa: per ottenere benefici occorre investire tempo per condividere casi d'uso concreti, e accompagnare i team finché non diventa parte integrante del loro modo di lavorare. Il successo dell'adozione, quindi, dipende più dalle **persone** (disponibilità ad apprendere e cambiare) e dai **processi** interni (adattati per includere la nuova tecnologia) che dalla tecnologia in sé.

Una tassonomia per valutare i tool: il "Vals Legal AI Report"

Un altro problema pratico emerso è come **valutare** in modo oggettivo gli strumenti di IA generativa disponibili sul mercato. Molti prodotti infatti promettono di “fare tutto”,

ed è difficile per un non-tecnico capire in cosa eccellono davvero e dove invece hanno punti deboli. Perez ha presentato a tal proposito il report “*Vals Legal AI*”, che propone una tassonomia di compiti specifici su cui misurare le performance di un sistema di IA legale. Invece di prendere per buone le affermazioni commerciali, *Vals* suggerisce di testare ogni strumento su alcune **funzionalità chiave** comuni:

- **Data Extraction:** ad esempio, la capacità di trovare una clausola specifica all’interno di un contratto complesso.
- **Document Q&A:** la capacità di rispondere a domande sui contenuti di un documento (esempio: “cosa prevede questa policy aziendale in caso di infortunio?”).
- **Summarization:** la capacità di riassumere testi lunghi, sintetizzando i punti principali con accuratezza.
- **Redlining:** la capacità di confrontare due versioni di un documento e evidenziare le differenze (utile nelle fasi di revisione contrattuale).
- **Transcription:** la capacità di trascrivere fedelmente un audio o un video (ad esempio una registrazione di una riunione o un dettato) in testo scritto.

Testando ogni tool su questi compiti specifici, lo studio legale può ottenere un **benchmark** oggettivo: confrontare le prestazioni della macchina con quelle di un operatore umano esperto su ciascun task. In questo modo diventa più chiaro se e dove l’IA apporta reale valore aggiunto. Ad esempio, se uno strumento eccelle in *data extraction* (trova clausole in pochi secondi che un umano impiegherebbe molto più tempo a individuare) ma magari è mediocre in *summarization*, lo si potrà adottare miratamente per il primo scopo, senza aspettarsi che svolga bene *anche* l’altro. Questa tassonomia aiuta dunque a **sgonfiare la retorica** del “tool miracoloso” e a capire concretamente come integrare vari strumenti in base ai rispettivi punti di forza.

Lezioni apprese e rischi persistenti

Concludendo il suo intervento, Perez ha condiviso alcune lezioni apprese e messo in guardia su rischi tuttora presenti nell’uso dell’IA generativa:

- **Le “allucinazioni” sono una caratteristica strutturale:** il fenomeno delle *hallucinations* in ambito IA – ossia quando il modello genera risposte apparentemente plausibili ma in realtà del tutto inventate¹ – non è un difetto isolato di un singolo software, bensì una conseguenza del funzionamento stesso dei modelli generativi. In altre parole, **ci si deve aspettare** che ogni assistente IA ogni tanto “fantasii”, perché questi sistemi non funzionano come banche dati precise, ma come modelli probabilistici del linguaggio. Sapendo ciò, l’avvocato deve adottare un atteggiamento prudente: trattare l’output dell’IA sempre come una bozza da verificare, mai come oro colato.
- **Qualità delle trascrizioni:** un esempio sorprendente di allucinazione si è riscontrato perfino in compiti apparentemente semplici come la trascrizione di audio in testo. Perez ha riferito che, in test riportati da Associated Press, anche strumenti di trascrizione automatica come Whisper a volte introducono parole o

frasi errate che **non** erano affatto pronunciate nell'audio originale, specie quando l'audio non è perfetto. Questo richiama alla massima **vigilanza**: anche quando ci affidiamo all'IA per compiti di routine, serve un controllo umano attento, perché l'errore può annidarsi dove meno ce lo aspettiamo.

- **Il futuro è degli “agenti” IA:** guardando avanti, Perez ha suggerito che la prossima evoluzione sarà costituita dagli **agenti intelligenti**, cioè sistemi di IA in grado di compiere workflow complessi composti da più passaggi e azioni in autonomia, anziché limitarsi a rispondere ad un singolo prompt. Un agente IA, ad esempio, potrebbe raccogliere informazioni da diverse fonti, eseguire analisi intermedie, prendere decisioni condizionate e infine produrre un risultato articolato – il tutto simulando un piccolo *processo* di lavoro multi-step. Questo modello di AI operativa per compiti composti somiglia molto di più al modo in cui lavora un avvocato, che raramente fa un'unica azione isolata, ma svolge una serie di attività concatenate (ricerca, redazione, verifica, comunicazione col cliente, ecc.). Prepararsi agli agenti IA significa dunque farsi trovare pronti a collaborare con sistemi ancora più autonomi e *proattivi* nel prossimo futuro, cogliendone le opportunità ma restando consapevoli dei rischi di delegare fasi intere di lavoro a una macchina.

(*Connessione tematica*: Nell'intervento di Perez emergono già alcuni fili conduttori che ritroveremo in tutta la lezione: l'equilibrio fra entusiasmo tecnologico e senso critico, la centralità del fattore umano (competenza e controllo) nell'uso dell'IA, e l'importanza di un metodo strutturato – *persone, processi, tecnologia* – per innovare con successo. I relatori successivi approfondiranno proprio questi aspetti da angolazioni diverse.)

3. Superare le barriere al cambiamento: psicologia e strategie per l'innovazione - Ing. Mauro Baldoni

Profilo e cambio di focus: Mauro Baldoni, Chief Information Technology Officer di uno studio legale, ha proseguito spostando il focus dalla tecnologia **all'elemento umano**. Il suo intervento ha esplorato i motivi per cui, anche quando le soluzioni tecnologiche sono disponibili, *le persone* e le organizzazioni spesso oppongono resistenza al cambiamento. Baldoni ha analizzato i principali bias cognitivi e le dinamiche psicologiche che ostacolano l'innovazione negli studi legali, offrendo al contempo alcuni consigli pratici per superare queste barriere.

Quattro ostacoli psicologici al cambiamento

Anzitutto, Baldoni ha identificato **quattro “default” psicologici** tipici che generano resistenza nelle organizzazioni legali di fronte all'innovazione tecnologica. Questi sono stati descritti quasi come *impostazioni di fabbrica della mente*, atteggiamenti istintivi che spingono a mantenere lo status quo:

- **Default di inerzia:** la tendenza a rimanere nella propria comfort zone professionale. È quel riflesso mentale riassunto dalla frase “abbiamo sempre fatto così”. In pratica, molti avvocati e staff di studio, di fronte a un nuovo

strumento o processo, provano un senso di disagio e preferiscono continuare con metodi tradizionali, anche se meno efficienti, semplicemente perché sono abituati ad essi.

- **Bias di conferma:** la propensione a cercare (e vedere) solo le prove che confermano le proprie paure o convinzioni preesistenti. Ad esempio, se qualcuno è scettico sull'IA, tenderà a ricordare e sottolineare il *singolo caso* in cui una IA ha prodotto un errore grossolano (come un'allucinazione o un suggerimento sbagliato), ignorando invece i numerosi casi in cui l'IA ha fatto risparmiare tempo e migliorato la produttività. Questo bias impedisce una valutazione equilibrata di rischi e benefici, perché la mente filtra la realtà in modo da auto-convincersi che "aveva ragione a diffidare".
- **Default sociale:** la riluttanza a promuovere innovazioni *disruptive* per timore di intaccare la propria reputazione o disturbare gli equilibri consolidati nella comunità professionale. In uno studio legale tradizionale, proporre cambiamenti radicali (come rivoluzionare il flusso di lavoro con nuove tecnologie) può essere visto come *andar controcorrente*. Chi vorrebbe introdurre la novità spesso si autocensura, preoccupato di come reagiranno i colleghi: "Sembrerò forse quello che gioca con i gadget invece di fare sul serio?". Questo conformismo sociale di gruppo può bloccare sul nascere idee innovative, perché nessuno vuole essere il primo a rompere le consuetudini consolidate.
- **Default dell'ego:** la resistenza che un professionista molto esperto può avere ad ammettere che una macchina possa svolgere meglio di lui certi compiti. Più anni di esperienza ha un avvocato senior, più può risultargli *inconcepibile* che un algoritmo rediga un contratto più velocemente, o che trovi una clausola nascosta più efficacemente. È quasi una difesa dell'orgoglio professionale: si tende a sminuire le capacità dell'IA per non mettere in discussione la superiorità acquisita con anni di mestiere. Questo atteggiamento, però, rischia di accecare di fronte ai *fatti* (quando la tecnologia effettivamente performa meglio su un micro-task) e impedisce di sfruttarla al meglio come alleato.

La trappola delle "scelte sicure"

A fianco dei bias cognitivi individuali, Baldoni ha evidenziato un comportamento diffuso a livello gestionale: l'**avversione al rischio** che porta a decisioni "di comodo" ma non ottimali. Si tratta della *trappola delle scelte sicure*. L'esempio emblematico citato è il detto: "*Nessuno è mai stato licenziato per aver scelto SAP*". Questo motto (variante locale del più noto "*nessuno è mai stato licenziato per aver comprato IBM*") riflette un approccio mentale molto comune: di fronte a una decisione tecnologica, spesso i responsabili preferiscono affidarsi al **leader di mercato** o alla soluzione consigliata dal blasonato consulente di turno, non perché sia oggettivamente la migliore per le proprie esigenze, ma perché è quella più *difendibile* sul piano politico interno. In caso di fallimento, infatti, poter dire "abbiamo seguito una strada collaudata adottando il prodotto più famoso" sembra tutelare da critiche personali, distribuendo la responsabilità sulle spalle larghe del fornitore rinomato. Questo atteggiamento prudente, tuttavia, ha un costo: porta a ignorare le grandi opportunità offerte magari da start-up più piccole e innovative, che però nessuno ha il coraggio di proporre perché

meno “sicure” dal punto di vista della carriera. Baldoni mette in guardia: *giocare sempre sul sicuro* in campo di innovazione tecnologica può voler dire perdere terreno rispetto a chi invece sperimenta soluzioni nuove più adatte al proprio contesto.

Le dinamiche decisionali negli studi legali

Restando sul tema delle decisioni organizzative, Baldoni ha analizzato come solitamente vengono prese le decisioni in materia di adozione di nuove tecnologie negli studi legali, evidenziando due approcci comuni, entrambi problematici:

- **Delega al singolo:** spesso la responsabilità di scegliere (o rifiutare) un certo strumento tecnologico viene delegata a un singolo partner o socio anziano. Questa persona si ritrova sotto un “occhio di bue” – cioè esposta, in piena luce, di fronte al resto dell’organizzazione – e sente su di sé tutto il peso della decisione. Comprensibilmente, un singolo decisore isolato avrà molta difficoltà a prendersi rischi: se la scelta innovativa fallisce, *lui o lei* ne porterà la colpa diretta. Questo porta spesso il delegato a optare per la scelta più conservativa o addirittura a rimandare la decisione, generando stallo.
- **Decisione comunitaria:** altre volte, per *diluire* la responsabilità, si cerca di decidere in modo collegiale coinvolgendo più soci o un comitato interno. In teoria questo dovrebbe portare a decisioni più equilibrate, ma in pratica – osserva Baldoni – spesso produce soluzioni di compromesso al **minimo comune denominatore**. Si finisce cioè per scegliere lo strumento che “vada bene un po’ a tutti”, spesso un prodotto generalista (*one-size-fits-all*, ad esempio un assistente basato su IA tipo Copilot integrato nella suite Office) che però non risolve in modo ottimale nessun problema specifico dello studio. La scelta collegiale “annacqua” l’innovazione per non scontentare nessuno, con il risultato che l’impatto trasformativo è modesto.

Strategie per promuovere l’innovazione

Dopo aver diagnosticato bias e dinamiche frenanti, Baldoni ha proposto una serie di consigli pratici per *promuovere l’innovazione* tecnologica negli studi legali, aggirando gli ostacoli psicologici e organizzativi. Queste strategie possono essere viste come passaggi di una roadmap culturale da seguire:

- **Riconoscere i propri bias:** Il primo passo è sviluppare **consapevolezza**. Sia i singoli professionisti sia i gruppi decisionali devono onestamente riconoscere di poter cadere nelle trappole psicologiche descritte (inerzia, conferma, ego, ecc.). Solo ammettendo di non essere completamente razionali nelle proprie paure e resistenze, si può provare a correggere il tiro.
- **Partire con progetti pilota:** Per ridurre il rischio percepito, Baldoni suggerisce di iniziare l’innovazione in **piccola scala**. Ad esempio, invece di implementare subito un nuovo software in tutto lo studio, si potrebbe avviare un progetto pilota in un singolo dipartimento o su un singolo processo, con un team ristretto di persone motivate. Un test in un contesto “sicuro” permette di raccogliere risultati e le eventuali criticità senza esporre l’intera organizzazione. Se il

progetto pilota fallisce, il danno è limitato; se ha successo, i suoi risultati positivi faranno da traino per convincere anche i più scettici ad estendere l'adozione.

- **Abbracciare l'open innovation:** Un altro consiglio è aprirsi alla collaborazione con l'ecosistema esterno, in particolare con **start-up locali** specializzate nel legal tech. Queste giovani aziende spesso hanno idee innovative tagliate su misura per il contesto nazionale (lingua italiana, procedure locali) e possono diventare partner preziosi. Coinvolgere attori esterni tramite hackathon, partnership o programmi pilota significa importare *energia fresca* e creatività, oltre a condividere il rischio dell'innovazione.
- **Imparare dagli altri, ma non copiare:** Guardare a cosa fanno gli studi legali internazionali più avanzati è certamente utile per trarre ispirazione. Tuttavia, Baldoni avverte che **copiare pedissequamente** modelli organizzativi altrui raramente funziona. Ad esempio, il fatto che uno studio anglosassone abbia interi team dedicati al *knowledge management* (gestione della conoscenza interna) o all'innovazione non implica che tali strutture siano replicabili in uno studio italiano di dimensioni medie. Ogni realtà ha la sua cultura e i suoi vincoli: meglio prendere spunto dalle best practice altrui adattandole criticamente alla propria situazione, piuttosto che importare ricette preconfezionate.
- **Promuovere una cultura dell'errore:** Innovare implica inevitabilmente la possibilità di sbagliare. Se lo studio coltiva una cultura dove ogni errore viene stigmatizzato e punito, nessuno si azzarderà mai a tentare strade nuove. Baldoni sostiene quindi la necessità di creare un ambiente in cui l'errore sia visto come **opportunità di apprendimento**. A tal proposito ha condiviso una colorita metafora nautica usata con il suo team: *"Se dovete fare dei buchi nella nave per testare, fateli sopra la linea di galleggiamento"*. In altre parole, sperimentate pure (bucate la nave) ma in modo intelligente, in alto, dove l'acqua non entrerà: ovvero predisponendo test controllati che non mettano a repentaglio la sopravvivenza dell'intero studio. E se qualche falla comunque si apre, niente caccia al colpevole, ma tutti a imparare come rinforzare lo scafo.
- **Diventare "data-driven":** Infine, Baldoni incoraggia un approccio **basato sui dati** nelle decisioni. Quando si provano nuove soluzioni, bisogna misurare i risultati in modo scientifico (tempi risparmiati, riduzione di errori, feedback dei clienti, ecc.), anziché affidarsi alle sensazioni o alle impressioni aneddotiche dei singoli. Ad esempio, se un software di analisi contrattuale viene testato, è utile raccogliere metriche concrete: quanti documenti ha analizzato, in quanto tempo, con che tasso di accuratezza rispetto a un revisore umano. Decisioni informate da dati permettono di *dimostrare* il valore (o la mancanza di valore) di una tecnologia, superando discussioni basate su opinioni soggettive.

(*Connessione tematica:* L'intervento di Baldoni rafforza un messaggio complementare a quello di Perez: anche le migliori tecnologie falliranno se non vengono affrontate le dimensioni umane e organizzative. Bias cognitivi, paure e abitudini radicate possono frenare la trasformazione digitale più della complessità tecnica. Ecco perché torna il mantra **Persone–Processi–Tecnologia**: l'innovazione efficace richiede di lavorare prima sulle persone (mentalità, competenze, cultura) e sui processi, per poi innestare gli strumenti tecnologici giusti. Il prossimo relatore, Di Pietro, approfondirà proprio l'idea di partire dai processi anziché dalla tecnologia in sé.)

4. Dal workflow allo strumento: un approccio pratico all'innovazione - Avv. Antonio Giuseppe Di Pietro

Profilo e idea chiave: Antonio Giuseppe Di Pietro, avvocato nell'ambito del diritto societario e dell'innovazione, ha offerto un punto di vista pragmatico, quasi **rovesciando l'approccio** comune. A suo avviso, spesso si commette l'errore di considerare l'innovazione come la ricerca della tecnologia più avanzata, per poi adattare il proprio modo di lavorare attorno ad essa. Di Pietro ha proposto l'approccio opposto: partire dall'analisi dei processi (i *workflow* quotidiani dello studio) e solo successivamente identificare lo strumento tecnologico più adatto in base alle esigenze specifiche. In questo modo vengono valorizzate le caratteristiche identitarie di ciascuno studio professionale, viene mitigata la difficoltà di apprendimento e inserimento di nuovi strumenti e viene assicurato l'efficientamento delle risorse.

"One tool, one workflow": prima il processo, poi la tecnologia

Secondo Di Pietro, l'errore più frequente è quello di **innamorarsi della tecnologia** del momento, imponendola sui processi esistenti o, all'estremo, forzando i collaboratori a usare uno specifico strumento solo perché lo studio ha sostenuto l'investimento in tal senso. L'approccio consigliato è invece l'inverso: analizzare con attenzione il proprio flusso di lavoro e, solo dopo averlo razionalizzato, scegliere lo strumento specifico utile a colmare ciascuna esigenza specifica. Questa filosofia viene riassunta dallo slogan **"one tool, one workflow"**, a significare che ogni strumento introdotto dovrebbe essere la risposta mirata a un'esigenza concreta per ciascun flusso organizzativo o gestionale, non una piattaforma generica calata dall'alto.

Come mettere in pratica questo approccio? Di Pietro ha illustrato alcune azioni chiave:

- **Mappare i flussi operativi:** La prima attività da svolgere è quella di descrivere ciascun processo interno nel dettaglio, anche graficamente in una mappa concettuale. Per esempio, il *workflow* relativo ad un'ingiunzione di pagamento avrà inizio dalla raccolta dei primi documenti presso il cliente, la loro organizzazione e nomenclatura, proseguendo con la verifica dei diversi presupposti speciali (come la provvisoria esecutività, la sussistenza di termini speciali di legge, ecc.), le modalità di redazione delle singole parti della costituzione in mora e del ricorso, la calendarizzazione delle scadenze, le istruzioni per il deposito telematico, e così via. Questa attività di mappatura e razionalizzazione è spesso già fondamentale per identificare ridondanze, passaggi superflui o inefficienze in genere, che nella routine quotidiana passano inosservate.
- **Identificare i punti dolenti:** Una volta disegnato il flusso, emergono chiaramente i punti dolenti: dove si verificano errori ricorrenti, quali informazioni sono necessarie in ciascuna fase o quali passaggi sono duplicati. Questi nodi sono le priorità su cui intervenire, così da introdurre soluzioni (anche software, ma non necessariamente) dove ci sono problemi.

- **Coinvolgere le persone giuste:** La mappatura e l'analisi non vanno fatte in astratto o in solitaria dal titolare, dal consulente o dall'Innovation Manager. Al contrario, è essenziale coinvolgere chi quotidianamente lavora sul processo che si vuole ottimizzare: avvocati operativi, praticanti, segretari, ecc. Queste persone hanno l'esperienza diretta e spesso sanno quali sono i veri problemi (e magari hanno già idee di soluzione). Inoltre, coinvolgendoli si ottiene un doppio beneficio: oltre a raccogliere informazioni utili, vengono *resi partecipi del cambiamento*, riducendo drasticamente la resistenza futura.
- **Solo dopo, scegliere lo strumento adatto:** Quando il workflow è messo nero su bianco e, sicuramente, è stato ottimizzato senza l'intervento di ulteriori investimenti, si passa a cercare la tecnologia. In questa fase vale la logica del *gap filling*: **uno strumento per ogni bisogno identificato**. Se, ad esempio, si scopre che un passaggio critico è l'estrazione di dati da una certa mole di documenti, ci si concentrerà su strumenti per l'estrazione di dati; se troppo tempo viene dedicato alla distribuzione dei compiti e alla pianificazione delle udienze, si valuterà un software di calendarizzazione intelligente, e così via. Non deve neanche trattarsi di strumenti disegnati per la nostra attività, l'importante è che lo strumento risponda a quel bisogno, invece di adottare piattaforme estremamente complesse sperando che "un po' risolvano tutto", per poi usarne solo alcune funzionalità.

In sintesi, "*one tool, one workflow*" significa che per ogni flusso di lavoro cruciale va trovata (o sviluppata) la soluzione *ritagliata su misura* per migliorarlo. Ciò porta a un'innovazione incrementale ma concreta, con benefici tangibili nel quotidiano, evitando il rischio di implementare strumenti sofisticati che restano inutilizzati perché non aderenti alle prassi reali. Tutto ciò, naturalmente, fino a quando il costo (di tempo e di denaro) per la gestione di una moltitudine di strumenti diventa a sua volta un'inefficienza e diventa necessario ricominciare l'analisi. Per questo motivo, l'innovazione è un ciclo e una linea retta.

Il ruolo delle figure non legali nell'innovazione

Per mettere in atto efficacemente quanto sopra, Di Pietro ha sottolineato l'utilità di coinvolgere **figure professionali non giuridiche**, in particolare esperti di organizzazione e progetto come i **Project Manager**. In un contesto di studio legale, un Project Manager dedicato (anche in modalità *fractional*, cioè a tempo parziale condiviso) può osservare dall'esterno i processi, coordinare la mappatura, facilitare i workshop con il personale e proporre migliorie. Queste figure hanno il vantaggio di portare prospettive diverse: non essendo giuristi, guardano ai flussi di lavoro con occhio laico, individuando incoerenze o sprechi che magari all'avvocato sfuggono perché "immerso" nella routine. Inoltre, liberano gli avvocati dal dover orchestrare in prima persona il cambiamento (cosa per cui spesso mancano tempo e competenze specifiche). Di Pietro ha fatto notare che investire in competenze manageriali e organizzative può avere un ritorno importante: un Project Manager può fare da "**ponte**" tra la tecnologia e le persone, assicurando che l'implementazione avvenga rispettando tanto le esigenze operative dello studio quanto gli obiettivi strategici.

Un esempio pratico di innovazione: le indagini interne

Per illustrare concretamente il suo approccio, Di Pietro ha raccontato un caso in cui ha applicato questi principi nell'ambito di un'*indagine interna* (ossia un'investigazione aziendale condotta dallo studio su richiesta di un cliente volta a verificare condotte illecite all'interno di una società). In questi contesti, ci si trova spesso a gestire **enormi volumi di dati** (email, documenti, log, ecc.), in un momento in cui l'utilizzo di IA generativa era indicata come la soluzione principale per ottimizzare il flusso di *e-discovery*. Ebbene, l'approccio adottato ha portato invece a non adottare soluzioni di IA generativa per il compito primario, ma tecniche di intelligenza artificiale non generativa con un mix di tecnologie più tradizionali ed efficaci per il caso concreto. In particolare, il *workflow* ridisegnato è stato il seguente:

1. **Identificazione del problema:** nel caso della specifica indagine da intraprendere, il problema era quello di identificare un gruppo ristretto di conversazioni tra centinaia di migliaia di email e documenti. Mettere nero su bianco tutti gli obiettivi ha permesso di elencare i criteri di selezione dei dati.
2. **Triage iniziale:** il primo passo non ha richiesto l'utilizzo di strumenti di IA avanzata, ma l'uso di strumenti classici di analisi forense per capire il volume e la natura dei dati. Software specialistici di lungo corso sono stati ottimali per indicizzare rapidamente tutti i documenti per parole chiave, metadati, interlocutori, ecc., fornendo una panoramica di quanti documenti (*artefatti*, in gergo) ci sono per ciascun criterio. Questo è servito a *delimitare il campo di analisi* alla fonte, ad esempio distinguendo tra 100.000 e-mail provenienti da un ufficio specifico e 2 milioni di file eterogenei.
3. **Classificazione assistita:** una volta eliminati i dati manifestamente irrilevanti, è stato introdotto il **Machine Learning** (non generativo, ma *supervisionato*). Sono stati addestrati modelli di intelligenza artificiale per la categorizzazione automatica di informazioni in insiemi significativi: ad esempio, raggruppandoli per mittente/destinatario, per reparto aziendale coinvolto, per argomento ricorrente, ecc. In questa fase l'IA c'è, ma non è quella generativa (come ChatGPT), ma una rete iterativa di algoritmi di classificazione e analisi testuale.
4. **Certificazione dei dati:** parallelamente, trattandosi di un flusso personalizzato, si è proposto di integrare la **Blockchain** per certificare la catena di custodia dei dati raccolti e processati. Nelle indagini è fondamentale garantire l'integrità delle prove digitali: registrare ogni accesso, spostamento e analisi dei file in modo immutabile serve a poter dimostrare, se necessario, che nulla è stato alterato. La blockchain pubblica o una rete *permissioned* consente di generare una "impronta temporale" sicura ad ogni file, a tutela della validità delle evidenze raccolte.

Seppur si tratti di una situazione particolare, stante che nella maggior parte dei casi l'IA generativa è uno strumento estremamente utile all'interno del flusso d'indagine, questo esempio dimostra due concetti fondamentali. Primo, che la soluzione efficace è spesso una **combinazione di tecnologie selezionate** in base alle esigenze specifiche, che non per forza include le ultime tendenze. In questo caso, l'IA generativa è stata

consapevolmente esclusa, poiché il problema da affrontare non era quello di riassumere o generare contenuti. Secondo, evidenzia l'importanza di **pensare al processo**: spezzettando il workflow in fasi (dall'identificazione del problema, al triage, alla classificazione, alla messa in sicurezza), si può decidere razionalmente dove applicare quale tecnologia.

Sguardo al futuro: la sfida della crittografia post-quantica

Nel chiudere il suo intervento, Di Pietro ha voluto allargare lo sguardo ai rischi e alle sfide tecnologiche all'orizzonte, citando in particolare l'avvento dei **computer quantistici**. Questa prospettiva, apparentemente lontana dal lavoro quotidiano dell'avvocato, in realtà pone una questione cruciale di sicurezza: i computer quantistici, quando maturi, saranno in grado di rompere gli attuali schemi di cifratura su cui si basano la riservatezza delle comunicazioni e l'integrità di moltissimi dati (la trasmissione di e-mail, le firme digitali, gli *hash* delle prove digitali, ecc.). In pratica, molte tecniche crittografiche oggi ritenute sicure diverrebbero inefficaci, esponendo potenzialmente tutti i dati sensibili ad alto rischio (compresi quelli creati fino a quel momento). Di Pietro ha sottolineato che questa *minaccia sul medio termine* richiede di pensare **fin da ora** a processi e misure di sicurezza *future-proof*, ovvero capaci di resistere anche all'era quantistica. Ciò significa iniziare a informarsi e investire su soluzioni di *crittografia post-quantum* (algoritmi di cifratura progettati per essere robusti anche contro attacchi effettuati con computer quantistici) e, più in generale, mantenere un approccio **proattivo** nella sicurezza informatica. Anche questo è innovare: non solo migliorare l'esistente, ma *anticipare* i problemi di domani per non farsi trovare impreparati.

(*Connessione tematica*: Il contributo di Di Pietro ribadisce, in chiave molto pratica, il messaggio che la tecnologia deve essere asservita al processo e non viceversa. L'innovazione davvero utile nasce dall'analisi del proprio lavoro (persone e processi) e dalla scelta mirata degli strumenti. Inoltre, emerge la necessità di guardare avanti con visione strategica, come nel caso della sicurezza post-quantistica. Il successivo intervento, di Marco Mendola, tornerà sul fattore umano e organizzativo, sposando un approccio quasi "terapeutico" all'introduzione del legal tech.)

5. Un approccio "terapeutico" al legal tech: Persone, Processi, Tecnologia – Dott. Marco Mendola

Profilo e visione umanistica: Marco Mendola, **Legal Technologist at TLT LLP**, vanta di esperienza professionale sia in Italia che nel Regno Unito, ha portato in aula una prospettiva originale e *umanistica* sull'innovazione. Egli parla di un approccio "terapeutico" al legal tech, indicando con questo termine la necessità di curare prima le *persone e le organizzazioni* nei loro timori e abitudini, prima ancora di prescrivere la "medicina tecnologica". Basandosi su alcuni aneddoti personali, Mendola ha sottolineato come la chiave del cambiamento non risieda tanto negli strumenti in sé, quanto nella capacità di ascolto e adattamento alle reali necessità.

Il problema non è la tecnologia, ma la mancanza di feedback

Mendola esordisce raccontando una sua **esperienza traumatica** in studi legali molto tradizionali. In quei contesti, la mentalità dominante era: *“Noi sappiamo cosa è meglio per il cliente, non serve chiedere”*. Tale atteggiamento – pur nato forse da fiducia nelle proprie capacità – si traduceva in una totale assenza di **feedback**: non si chiedeva mai al cliente un’opinione sul servizio, non si interrogavano i collaboratori sulle procedure interne, perché si dava per scontato di avere già tutte le risposte. Questo, a detta di Mendola, è un problema ben più grande della scelta di quale software usare: se manca l’ascolto, qualsiasi innovazione è destinata a fallire o a non essere compresa.

L’esperienza gli ha insegnato che prima di introdurre una nuova tecnologia bisogna creare una cultura in cui tutti i soggetti coinvolti – dai praticanti fino ai clienti – possano dare un riscontro, evidenziare esigenze e criticità. Solo *ascoltando* si può capire cosa davvero serve e quindi scegliere lo strumento giusto. In pratica, la tecnologia di per sé non risolve alcun problema se **a monte** non c’è stato un dialogo aperto sui problemi da risolvere. Al contrario, quando un’organizzazione è capace di recepire feedback e adattarsi, allora qualunque nuova soluzione (tecnologica o anche solo di prassi) verrà accolta e utilizzata molto più efficacemente. Ecco perché Mendola sostiene che l’innovazione ha un che di *terapeutico*: occorre prima guarire l’organizzazione dalla *sordità* verso i suoi membri e i suoi utenti, poi si potrà curare anche l’efficienza con gli strumenti digitali.

Il framework “Persone, Processi, Tecnologia”

Riprendendo un leitmotiv già introdotto nei precedenti interventi, Mendola ha ribadito che **“Persone, Processi, Tecnologia”** non deve essere uno slogan vuoto ma una guida concreta nell’implementare il cambiamento. Ha quindi declinato questo framework in termini molto pratici:

- **Persone:** Il punto di partenza di ogni innovazione riuscita è sempre l’ascolto attivo delle persone coinvolte. Un esempio portato da Mendola: invece di lanciarsi in un costoso progetto di digitalizzazione degli archivi perché “fa innovazione”, un managing partner farebbe meglio a chiedere a un praticante di studio o a una segretaria come trovano l’attuale gestione dei faldoni cartacei. Spesso, da una semplice chiacchierata con chi lavora sul campo emergono indicazioni preziosissime, più utili di mille analisi teoriche. Inoltre è fondamentale **circondarsi di persone giuste**, non solo giuristi ma anche figure con competenze diverse (tecnologi, analisti, project manager come visto prima). Queste prospettive diverse agiscono come anticorpi contro il pensiero unico e aiutano a progettare soluzioni più equilibrate.
- **Processi:** Prima ancora di pensare a quale nuova app o software adottare, bisogna guardare criticamente alle proprie **consuetudini di lavoro** e chiedersi *“perché facciamo così?”*. Spesso i processi interni vengono tramandati per inerzia e presentano passaggi inutilmente complicati o ridondanti. L’obiettivo è

semplificare il modo di lavorare eliminando il superfluo e ottimizzando ciò che resta, **prima** di provare ad automatizzarlo. C'è un detto nell'informatica: "un processo inefficiente automatizzato resta un processo inefficiente, solo più veloce". Mendola sottolinea proprio questo: automatizzare *dopo* aver migliorato il processo, non prima.

- **Tecnologia:** Solo a questo punto entra in scena la tecnologia. La regola proposta è "*meno compri, meglio è*". Ovvero, non bisogna collezionare software e piattaforme per il gusto di averli o perché "il competitor l'ha fatto". Ogni nuovo strumento va introdotto come **conseguenza** di uno studio serio: deve essere stato identificato un obiettivo preciso, valutate le risorse disponibili, analizzati i casi d'uso di altri (market analysis) e tratte lezioni dalle esperienze altrui. Solo allora si procede all'acquisto o sviluppo della tecnologia, con la consapevolezza esatta di cosa deve risolvere e come verrà adottata. In mancanza di questa disciplina, c'è il rischio di spendere budget in soluzioni alla moda che poi non attecchiscono perché non c'era un vero bisogno o non si è lavorato abbastanza con le persone e i processi per preparare il terreno.

In sintesi, *Persone, Processi, Tecnologia* è come uno sgabello a tre gambe: se anche solo una delle tre manca o è traballante, l'intero sforzo di innovazione crolla. Mendola fa eco così, in anticipo, a quello che sarà anche il messaggio conclusivo di Gilardi: la tecnologia deve essere l'ultimo tassello, sorretto dagli altri due.

Il ciclo virtuoso dell'innovazione

A coronamento del suo discorso, Mendola ha descritto l'innovazione come un **processo ciclico e continuo**, non un progetto con un inizio e una fine definitivi. Lo ha schematizzato in un ciclo virtuoso di cinque fasi ricorrenti:

Ascoltare → Semplificare → Automatizzare → Raccogliere Feedback → Ripetere.

In altre parole: prima si **ascoltano** le esigenze (di colleghi, collaboratori, clienti), poi si **semplifica** ciò che non va nei processi esistenti, quindi si **automatizza** con tecnologia ciò che è maturo per essere migliorato. Ma il lavoro non finisce qui: dopo l'implementazione, si **raccoglie di nuovo feedback** su come sta funzionando il cambiamento, e da queste nuove informazioni si riparte per ulteriori miglioramenti, iterativamente. L'obiettivo finale di questo ciclo non è mai fine a se stesso, bensì consentire di **vivere meglio la professione**: avere avvocati e staff meno appesantiti da compiti ripetitivi e frustrazioni, clienti più soddisfatti per la maggiore velocità e qualità del servizio, e in ultima analisi anche migliori risultati economici per lo studio (perché un cliente soddisfatto genera fidelizzazione e passaparola positivo). Questa visione dinamica dell'innovazione, quasi *organica*, implica che uno studio legale dovrebbe sempre considerarsi "in beta", cioè in costante evoluzione, adattandosi ai cambiamenti tecnologici e di mercato in modo graduale ma continuo.

(*Connessione tematica:* L'approccio di Mendola enfatizza, come gli altri interventi, che la tecnologia è efficace solo se inserita in un contesto umano preparato. La sua metafora

terapeutica e il ciclo di miglioramento continuo riecheggiano la necessità di visione a lungo termine e di cura costante dell'equilibrio **Persone-Processi-Tecnologia**. Questo concetto verrà splendidamente riassunto nelle conclusioni di Gianluca Gilardi, che ora andremo ad esaminare.)

6. Dall'approccio al metodo - Dott. Gianluca Gilardi.

Profilo e sintesi finale: Gianluca Gilardi, chiamato a concludere la giornata, ha avuto il compito di *tirare le fila* dei vari interventi, sintetizzando i temi emersi e fornendo una chiave di lettura unitaria. La sua chiusura ha evidenziato i concetti principali da portare a casa, delineando anche delle prospettive sull'immediato futuro per gli avvocati alle prese con l'IA. Gilardi ha strutturato il suo intervento finale attorno a cinque messaggi chiave:

Da “nice to have” a “must have”

Per prima cosa, Gilardi ha evidenziato come la **competenza in materia di IA** per un avvocato stia rapidamente passando da optional a indispensabile. Solo pochi anni fa, conoscere l'IA era un “nice to have” (cioè un extra interessante ma non essenziale); oggi è già diventato uno “should have” (sarebbe opportuno che un professionista ne sapesse qualcosa); e nell'arco di 2-3 anni diventerà verosimilmente un vero “must have” – un requisito imprescindibile. Questo significa che gli avvocati hanno ancora un piccolo margine temporale per sperimentare e mettersi in pari, ma non troppo: chi dovesse trascurare completamente queste competenze rischia di trovarsi presto tagliato fuori o fortemente svantaggiato nel mercato legale. La transizione da nice-to-have a must-have lancia un messaggio di **urgenza moderata**: non c'è bisogno di farsi prendere dal panico immediato, ma neppure di adattarsi, pensando che l'AI sia un interesse di nicchia. Diventerà parte integrante della cassetta degli attrezzi dell'avvocato e conviene attrezzarsi per tempo.

La *sunk cost fallacy* e il coraggio di abbandonare

Un altro concetto importante che Gilardi ha richiamato è la *fallacia del costo sommerso* (**sunk cost fallacy**), applicata ai progetti di innovazione. Questo bias cognitivo – complementare a quelli illustrati da Baldoni – consiste nella riluttanza ad abbandonare un progetto in cui si sono già investite molte risorse (tempo, denaro, formazione), anche quando sarebbe evidente che non funziona o non ne vale più la pena. Gilardi ha avvertito che nel testare nuove soluzioni tecnologiche bisogna avere la lucidità e il **coraggio di dire “ho provato, non ha funzionato, passo ad altro”**. Non c'è vergogna nel cambiare rotta: anzi, insistere su un errore per il solo fatto di averci già investito molto è il vero pericolo. In pratica, se uno studio adotta un certo software ma dopo sei mesi si accorge che crea più problemi di quanti ne risolve, dovrebbe saper fare marcia indietro senza troppi rimpianti, invece di continuare a investirci soldi e tempo sperando magicamente che le cose migliorino. L'innovazione efficace richiede anche **capacità di rinuncia**: ogni

tanto bisogna saper staccare la spina a un progetto fallimentare e reindirizzare le energie altrove, imparando la lezione per il futuro. Questo punto completa la serie di raccomandazioni di mindset per innovare: accanto all'entusiasmo e alla perseveranza, serve anche la *spietata onestà* nel riconoscere quando un'idea non ha funzionato.

Non esiste la bacchetta magica

Riprendendo un tema toccato specialmente da Perez e in parte da Di Pietro, Gilardi ha ribadito che **non esiste lo strumento magico** capace di risolvere tutti i problemi di uno studio legale. Nessun singolo software, per quanto potente, può sostituire in un colpo solo l'insieme di competenze, processi e strumenti necessari nella pratica quotidiana. Pertanto, la ricerca della soluzione unica e perfetta è destinata a deludere. Meglio invece adottare un approccio pragmatico: partire dal proprio bisogno specifico (ciò che *a me* serve davvero) e cercare la tecnologia che lo soddisfi, oppure in alternativa "giocare" e sperimentare vari tool per comprenderne le potenzialità in astratto, senza però illudersi che uno di essi eliminerà tutte le complessità. Gilardi suggerisce dunque di mantenere **aspettative realistiche**: l'IA e il software possono dare grandi vantaggi, ma vanno "a incastro" l'uno con l'altro e con il lavoro umano. In altre parole, si avrà magari un ottimo programma per la gestione documentale, un modello generativo per assistenza nella scrittura, un sistema esperto per la ricerca giurisprudenziale... ognuno eccellente nel suo compito, ma nessun *anello del potere* che comanda su tutti gli altri (parafrasando la citazione tolkieniana già evocata con ChatGPT). Così facendo, si evita la disillusione e si costruisce un arsenale tecnologico **diversificato** e appropriato.

Il vero rischio: mancanza di visione d'insieme

Tra i pericoli maggiori, ha continuato Gilardi, c'è quello di focalizzarsi troppo sul **singolo tool** senza comprendere i principi di funzionamento generali della tecnologia e senza una strategia complessiva. Questo atteggiamento espone al rischio di essere *travolti* dall'innovazione. In pratica, chi rincorre ogni nuovo gadget o applicazione senza avere una visione d'insieme finisce per reagire passivamente ai trend, integrando male le varie novità e magari stravolgendo di continuo il proprio modo di lavorare senza mai consolidare nulla. Gilardi invita invece a sviluppare una **visione sistemica**: capire come le varie tecnologie (IA, automazione, analisi dati, ecc.) si interrelano e quale filosofia comune le guida, così da poterle inserire armonicamente nei propri processi. Inoltre, avere strategia significa sapere *perché* si sta innovando: per migliorare quale aspetto (efficienza? qualità? competitività?) e con quale priorità. Senza questa bussola, si rischia di adottare strumenti a casaccio e di "perdere la rotta" cambiando direzione a ogni onda. In sintesi, la mancanza di visione d'insieme è l'antitesi di quel metodo strutturato su cui tutti i relatori hanno insistito. Per evitarla, Gilardi consiglia di ritagliarsi momenti per formarsi e riflettere sui principi dell'IA e del legal tech, magari facendo riferimento a frameworks consolidati (come, di nuovo, Persone-Processi-Tecnologia) per valutare ogni decisione innovativa nel quadro generale.

Sintesi finale: il tavolino a tre gambe

Giunto alla conclusione, Gilardi ha proposto una metafora semplice ma efficace per riassumere l'intera giornata: *l'innovazione è come un tavolino a tre gambe*. Le tre gambe sono **Persone, Processi e Tecnologia**. Se anche solo una di queste viene a mancare, il tavolino crolla. Questa immagine racchiude il messaggio chiave emerso da tutti gli interventi: l'adozione di nuove tecnologie (IA compresa) non può prescindere dall'analisi e dal coinvolgimento delle persone e dal ripensamento dei processi. Ogni studio legale, prima di correre ad installare l'ultimo software, dovrebbe fare un'**autoanalisi** onesta: capire quali sono i propri bisogni reali, dove sta andando il mercato, quali competenze interne ha e quali gli mancano, e soprattutto verificare di avere la volontà di mettere **persone e processi al centro**. Solo dopo questo lavoro preparatorio la scelta dello strumento (la terza gamba, la tecnologia) darà frutti. Diversamente, si rischia di investire in innovazione e ritrovarsi con niente in mano – come un tavolino monco che non sta in piedi. Gilardi, così, ci lascia con un approccio chiaro: affrontare la trasformazione digitale con metodo, equilibrio e visione integrata.

Capitolo VIII

AI Act e Professione Legale: rischi, trasparenza e strategie di adattamento

1. Introduzione

L’ottava lezione del **Corso di Perfezionamento in Intelligenza Artificiale per Avvocati** si addentra nel cuore della nuova regolamentazione europea sull’IA, con un’analisi approfondita del cosiddetto **AI Act** (il Regolamento europeo sull’intelligenza artificiale) e del suo impatto diretto e indiretto sulla professione legale. In questa sessione un panel di accademici e avvocati pratici esplora la normativa non come un mero onere di *compliance* burocratica, ma come un framework che ridefinisce rischi, responsabilità e strategie operative per gli studi legali. L’obiettivo è **demistificare i concetti chiave** dell’AI Act – dall’approccio basato sul rischio alla classificazione dei sistemi (con un focus critico sulla posizione degli strumenti usati dagli avvocati) – e chiarire gli **obblighi di trasparenza**. Vengono affrontati anche i “miti” legati all’IA, le sfide della sua governance all’interno di organizzazioni complesse come gli studi legali, e le interazioni con altre normative preesistenti (in particolare il GDPR sulla protezione dei dati). Il fine formativo dichiarato è fornire ai partecipanti una mappa concettuale e una *roadmap* pratica per navigare il nuovo panorama normativo, trasformando gli obblighi di legge in un’occasione per aumentare la consapevolezza, migliorare i processi interni e cogliere il vantaggio competitivo offerto da un uso informato e responsabile dell’intelligenza artificiale.

Di seguito ripercorriamo in ordine cronologico gli interventi dei relatori di quella giornata, ciascuno dei quali ha contribuito da una prospettiva specifica a questo quadro d’insieme. Ogni sezione è dedicata a un intervento, con titoli e sottotitoli che ne riflettono i temi principali, e arricchisce i punti trattati con esempi e collegamenti utili a una comprensione più chiara. Il tono è divulgativo e accessibile, pensato per un pubblico colto ma non necessariamente esperto di tecnologia.

2. Coordinate di Lettura dell’AI Act: Rischi, Trasparenza e Paradossi Normativi - Avv. Giuseppe Vaciago

Giuseppe Vaciago, avvocato esperto in diritto penale delle tecnologie, apre la giornata fornendo una cornice pratica per interpretare l’impatto dell’AI Act sulla professione legale. Nel suo intervento introduttivo affronta innanzitutto un **paradosso normativo** che emerge dalla bozza del regolamento, per poi concentrarsi su un tema di immediata rilevanza pratica: l’obbligo di trasparenza nell’uso dell’IA da parte degli avvocati.

Il paradosso dell’Alto Rischio per la Professione Legale

Vaciago evidenzia subito un’apparente contraddizione presente nell’**Allegato III** dell’AI Act, quello che elenca le applicazioni di IA considerate “ad alto rischio”. In base a tale

allegato, i software di intelligenza artificiale utilizzati **dai magistrati** (ad esempio nell'ambito decisionale giudiziario) o **per la mediazione** sono classificati come sistemi ad alto rischio. Al contrario, gli stessi tipi di software se utilizzati **dagli avvocati** – dunque dai professionisti forensi nel loro lavoro quotidiano, in qualità di utilizzatori o “*deployer*” del sistema – **non** rientrerebbero in questa categoria. Vaciago sottolinea come questa distinzione appaia paradossale: un giudice e un avvocato potrebbero infatti impiegare il **medesimo software** di IA (nel discorso viene fatto l'esempio di un ipotetico programma chiamato *Sapienza*) per analisi giuridiche o supporto decisionale, ma a livello normativo quel software sarebbe considerato rischioso solo nel primo caso.

Le conseguenze pratiche di tale impostazione, spiega Vaciago, sarebbero notevoli. **Allo stato attuale**, un fornitore di software destinato ai magistrati dovrebbe ottenere una **certificazione** ufficiale (una sorta di “bollino” di conformità) e seguire un rigoroso percorso di adeguamento alle prescrizioni del Regolamento, dato che il suo prodotto rientrerebbe tra quelli ad alto rischio. Lo stesso identico prodotto venduto agli avvocati, invece, non richiederebbe alcuna particolare certificazione né l'adempimento degli obblighi più stringenti previsti per gli usi “ad alto rischio”. In altre parole, il legislatore europeo – almeno nella bozza in discussione nel giugno 2025 – sembrerebbe considerare più pericoloso l'uso dell'IA in mano a un magistrato rispetto allo stesso uso fatto da un avvocato, nonostante l'output del software possa incidere su decisioni giudiziarie o strategie legali in modo simile. Vaciago indica di ritenere che questa **incongruenza verrà corretta**: è verosimile, secondo lui, che l'allegato in questione sarà aggiornato nelle fasi finali di approvazione dell'AI Act proprio per sanare il paradosso e includere anche gli utilizzi forensi tra quelli regolati più strettamente. Questo primo punto dell'intervento di Vaciago mette in luce le difficoltà di classificazione nell'approccio *risk-based* (basato sul rischio) adottato dall'AI Act, quando applicato a contesti specifici come la professione legale.

L'obbligo di Trasparenza: un Punto Fermo

A differenza delle incertezze ancora presenti sulla classificazione del rischio, Vaciago sottolinea che **l'obbligo di trasparenza** – ossia il dovere di dichiarare l'utilizzo di sistemi di IA – costituisce già **un punto fermo** e di applicazione immediata per gli avvocati. La professione forense, spiega, si trova attualmente in una sorta di “zona grigia”: molti avvocati **già oggi utilizzano** strumenti di intelligenza artificiale (si pensi a software per la ricerca giurisprudenziale, per l'analisi di contratti, o persino a assistenti virtuali per la gestione dello studio) **senza** però, in molti casi, **comunicarlo ai clienti**. Questa mancanza di trasparenza non è stata finora codificata da una norma vincolante specifica, ma il panorama sta cambiando rapidamente. Vaciago evidenzia come il dovere di informare il cliente dell'uso di IA trovi fondamento in più fonti concorrenti:

- **La Carta dei Principi dell'Ordine degli Avvocati di Milano (2024)**, un documento deontologico pionieristico, raccomanda espressamente un dovere di **trasparenza** nell'impiego di strumenti di IA in ambito forense. Questo significa

che, secondo tali linee guida, l'avvocato dovrebbe dichiarare al cliente quando si avvale dell'assistenza di un software intelligente nei servizi prestati.

- **Il disegno di legge italiano sull'IA** (il cosiddetto *DDL IA* in discussione in Parlamento) contiene anch'esso – se approvato nella sua forma attuale – un chiaro obbligo di trasparenza per i professionisti. In pratica, la legge nazionale richiederebbe a tutti i professionisti (avvocati inclusi) di avvisare clienti e utenti quando utilizzano sistemi di IA nello svolgimento dell'attività, rendendo questa prassi una prescrizione legale cogente.
- **Lo stesso AI Act europeo**, anche se non ancora in vigore, prevede un obbligo generale di trasparenza in capo agli utenti di sistemi di IA. In particolare, una volta adottato, il Regolamento UE imporrà di comunicare alle persone quando si interfacciano con un sistema di intelligenza artificiale (ad esempio nel caso di chatbot o sistemi decisionali automatizzati utilizzati nello studio legale). Dunque, indipendentemente dai dettagli su categorie di rischio, su questo aspetto l'AI Act stabilisce un principio netto: niente "IA occulta" nei servizi professionali.

Dato che **questo obbligo di trasparenza sembra ormai certo e imminente**, la domanda cruciale diventa operativa: **come** e **dove** applicare concretamente questo *disclaimer* nelle attività legali quotidiane? Vaciago solleva interrogativi pratici che molti professionisti condividono. È sufficiente, ad esempio, una menzione generale nel mandato professionale o nel contratto con il cliente in cui si dichiara che "lo studio potrebbe utilizzare strumenti di IA"? Oppure occorre inserire un'avvertenza esplicita **in ogni singola e-mail** o documento prodotto con l'ausilio dell'IA? Inoltre, se davvero si deve essere trasparenti, fino a che punto bisogna dettagliare? La normativa in arrivo potrebbe richiedere di specificare **anche quale tipo di IA** è stato utilizzato e **in che modo** lo si è impiegato nel caso concreto. Ciò ovviamente complica l'adempimento, perché significherebbe, ad esempio, dover spiegare al cliente se si è usato un algoritmo di *machine learning* per analizzare decine di contratti, oppure una rete neurale generativa per predisporre una bozza di atto, e così via. Al momento – riconosce Vaciago – **manca una risposta definitiva** su quale sia la modalità corretta di comunicare tutto ciò al cliente; è probabile che saranno le linee guida applicative (magari emanate dagli Ordini professionali o dalle autorità) a chiarirlo con il tempo. L'importante, intanto, è prendere coscienza che la **trasparenza verso il cliente** sull'uso dell'IA sta per diventare parte integrante della deontologia e presto anche un obbligo di legge: gli studi legali farebbero bene a prepararsi, definendo già ora una strategia per "dirlo in modo chiaro" ai propri assistiti.

(Transizione: Dopo questo quadro introduttivo sulle coordinate generali – rischio e trasparenza – il focus della lezione si sposta su un inquadramento più teorico-normativo. Il Prof. Alessandro Mantelero prende la parola per smontare alcuni luoghi comuni sull'IA e collocare l'AI Act nel contesto più ampio delle regole e responsabilità esistenti.)

3. Miti e Realtà dell'IA: Inquadramento dell'AI Act e della Responsabilità - Prof. Alessandro Mantelero

Alessandro Mantelero, docente di diritto e tra i massimi esperti di etica e regolazione dell'IA, offre nel suo intervento una visione d'insieme che sfata alcuni **miti comuni** sull'intelligenza artificiale e al tempo stesso contestualizza l'approccio normativo europeo. Il suo contributo aiuta a capire perché l'AI Act è stato concepito in un certo modo e come si inserisce nel quadro giuridico complessivo. Mantelero articola la sua trattazione in quattro punti principali: (1) la radice dell'approccio basato sul rischio, (2) due diffusi malintesi sull'IA da correggere, (3) il chiarimento che l'AI Act non esaurisce tutte le norme in materia di IA, e (4) una riflessione sul recente ritiro della proposta di direttiva europea sulla responsabilità civile da IA.

Le radici dell'approccio basato sul rischio

Per comprendere lo spirito dell'AI Act, spiega Mantelero, occorre partire da un dato di fatto: negli ultimi anni l'avvento dell'**IA generativa** (come i noti modelli linguistici alla ChatGPT) ha **rotto il paradigma** della cosiddetta "*Responsible Innovation*", cioè di un'innovazione tecnologica introdotta sul mercato solo dopo averne responsabilmente valutato e mitigato i rischi. In pratica, sono stati rilasciati al pubblico **prodotti basati su IA** con malfunzionamenti noti e rischi intrinseci significativi, senza aspettare di risolverli. Un esempio concreto è l'immissione in commercio di sistemi di IA che possono generare contenuti inesatti o discriminatori (si pensi ai modelli generativi che a volte producono informazioni false, le cosiddette "allucinazioni", o algoritmi di decisione automatica che mostrano bias). La velocità e l'aggressività con cui queste tecnologie sono state distribuite ha messo in crisi i meccanismi tradizionali di controllo e certificazione. **L'approccio risk-based** dell'AI Act – cioè la scelta di calibrare obblighi e divieti in base al livello di rischio del sistema – nasce proprio come **risposta** a questa situazione. In altre parole, di fronte a tecnologie IA "imperfette" e non del tutto controllabili, il legislatore europeo ha optato per un modello regolatorio graduale: tollerante verso le applicazioni a basso rischio, ma molto rigoroso con quelle ad alto rischio o addirittura proibitivo con quelle considerate inaccettabili (come i sistemi di social scoring). Questo spiega perché il regolamento classifica i sistemi in categorie e impone misure via via più stringenti man mano che aumenta il rischio che possano nuocere a persone o diritti.

Sfatare i miti sull'intelligenza artificiale

Un contributo prezioso di Mantelero è poi quello di riportare l'IA "con i piedi per terra", smontando due idee assai diffuse ma errate sul funzionamento di questi sistemi. **Mito 1:** "L'IA è intelligente come un essere umano". **Falso**, afferma il professore senza mezzi termini. L'intelligenza artificiale, per quanto sofisticata, **non ragiona come noi**. Può eseguire calcoli complessi e analizzare enormi quantità di dati, ma **commette errori di natura diversa** rispetto agli umani, spesso in modo seriale e sistematico. Le cosiddette "**allucinazioni**" ne sono un esempio: un modello di IA generativa può inventare fatti o

riferimenti del tutto inesistenti con assoluta sicurezza, fuorviando l'utente inesperto. Per questo, sottolinea Mantelero, la **supervisione umana** dei sistemi di IA non è solo un requisito normativo formale, ma una necessità pratica imprescindibile. La responsabilità finale degli output, infatti, **ricade sempre su chi utilizza l'IA**, che deve verificare e filtrare criticamente i risultati prodotti dalla macchina. *Un caso emblematico* a riguardo è balzato agli onori della cronaca nel 2023, quando un giudice negli Stati Uniti ha sanzionato due avvocati che avevano depositato una memoria contenente riferimenti giurisprudenziali del tutto inventati da ChatGPT. L'episodio illustra perfettamente come **affidarsi acriticamente** a un sistema automatizzato possa mettere nei guai: l'IA aveva generato citazioni di sentenze mai esistite, gli avvocati non le hanno verificate e il risultato è stato un atto fallace, con relative conseguenze disciplinari. In sintesi, l'IA può essere uno strumento utilissimo ma **non è infallibile** né veramente "intelligente" in senso umano: va considerata alla stregua di un assistente che lavora in modo differente dalla mente umana e pertanto **va sempre tenuta sotto controllo**.

Mito 2: "L'IA pensa e impara in modo creativo". Anche questo è **falso**. I moderni sistemi di IA – spiega Mantelero – sono in realtà **largamente retrospettivi**: si basano sull'analisi di grandi quantità di dati **passati** per produrre risultati. In altri termini, "imparano" da ciò che è già accaduto e tendono a replicare schemi preesistenti. Questo li rende intrinsecamente **conservativi**, limitando di molto la loro capacità di innovare veramente o di trovare soluzioni creative in senso proprio. Certo, un algoritmo può combinare informazioni pregresse in modi nuovi, ma resta vincolato al perimetro di ciò che ha già visto nei dati di addestramento. Per una professione come quella legale, che evolve anche attraverso interpretazioni originali e principi nuovi creati dai giuristi e dai giudici, questa **limitazione è cruciale**. Viene citato un esempio calzante: la creazione giurisprudenziale del concetto di "*danno biologico*" da parte dei giudici italiani (un tipo di danno alla persona risarcibile, non previsto espressamente dal legislatore ma riconosciuto dalla pratica giurisprudenziale). Un'IA addestrata solo su sentenze passate **non avrebbe mai potuto "inventare"** un istituto del genere, perché nei dati pregressi non ve ne sarebbe traccia. L'innovazione giuridica, dunque, richiede quella dose di creatività e "rottura" degli schemi che è propria dell'intelligenza umana ma estranea al modo di operare degli algoritmi odierni. Questo secondo mito sfatato serve a ricordare che, almeno per ora, l'IA **non può sostituire l'iniziativa creativa** e l'ingegno del giurista: può supportare, velocizzare, suggerire, ma difficilmente inventerà soluzioni giuridiche nuove. In definitiva, Mantelero ribadisce che tali miti – l'IA antropomorfa e l'IA creativa – devono essere superati, così da avere una visione lucida di cosa queste tecnologie possono e non possono fare, e regolamentarle di conseguenza.

L'AI Act e le altre normative vigenti

Un altro chiarimento offerto dal professor Mantelero riguarda un possibile equivoco sul ruolo dell'AI Act nel più ampio ecosistema normativo. **L'AI Act non è (e non sarà) "la" legge unica sull'IA**. Al contrario, ha un campo di applicazione specifico e mirato, e lascia espressamente impregiudicato tutto il resto dell'ordinamento. Significa che **tutte le**

altre normative esistenti – in particolare quelle sulla **responsabilità professionale e civile per danni** – **continuano ad applicarsi** ai casi che coinvolgono sistemi di IA. Ad esempio, se un avvocato utilizza un software di IA e da ciò deriva un danno al cliente, si guarderà non solo all'AI Act (quando pienamente operativo) ma anche alle regole ordinarie sulla responsabilità contrattuale o professionale. Ancora, resta in pieno vigore il GDPR per quanto riguarda la protezione dei dati personali trattati dall'IA, così come le leggi sulla proprietà intellettuale per gli output generati, e così via. L'AI Act dunque **non sostituisce né annulla** le leggi preesistenti, ma le integra focalizzandosi su specifici profili di rischio e trasparenza legati all'uso di IA. Per i giuristi e gli avvocati questo significa che l'adozione dell'IA va affrontata in un'ottica **multilivello**: bisognerà rispettare il Regolamento europeo, certo, ma anche **continuare a considerare tutte le normative applicabili** (prima fra tutte, appunto, le norme su privacy e responsabilità). Un approccio olistico rimane fondamentale per evitare vuoti di tutela.

Il ritiro della proposta di Direttiva sulla responsabilità civile da IA

In chiusura del suo intervento, Mantelero commenta una notizia recente nel panorama normativo: la Commissione Europea ha **ritirato la proposta di direttiva** sulla responsabilità civile per l'IA. Si trattava di uno strumento legislativo parallelo all'AI Act, che mirava a armonizzare a livello UE le regole per imputare i danni causati da sistemi di IA. Ebbene, secondo il professore si è trattato di **una scelta positiva**. Le ragioni sono diverse. Anzitutto, una direttiva avrebbe richiesto recepimenti nazionali da parte di ciascun Stato membro, con il rischio di **frammentazione normativa**: paesi diversi avrebbero potuto implementarla in modo non uniforme, vanificando in parte lo scopo di armonizzare. Inoltre – osserva Mantelero – la logica *risk-based* dell'AI Act dovrebbe già di per sé **ridurre drasticamente il numero di casi di danno**, perché imponendo misure preventive rigorose sui sistemi ad alto rischio ci si aspetta meno incidenti gravi. Ciò renderebbe meno urgente avere subito una normativa ad hoc sulla responsabilità civile. Infine, vi è da considerare lo **“stress normativo”** a cui sono sottoposti operatori e imprese nel digitale: GDPR, Digital Services Act, Digital Markets Act, ecc., e ora l'AI Act – aggiungere anche un'ulteriore direttiva complessa sarebbe stato prematuro e avrebbe rischiato di sovraccaricare il sistema. In altre parole, si è preferito evitare di mettere troppa carne al fuoco contemporaneamente. Ritirando la proposta, il legislatore europeo sembra voler prima vedere gli effetti pratici dell'AI Act e delle norme esistenti, lasciando eventualmente agli Stati la gestione con gli strumenti di responsabilità tradizionali (es. colpa, responsabilità oggettiva, ecc.) per i casi di danni da IA. Solo in un secondo momento, se necessario, si potrà ripensare a una normativa unitaria sulla responsabilità civile in questo campo. Questo approccio graduale e prudente è condiviso da Mantelero, che lo ritiene sensato per evitare confusione e sovrapposizioni legislative.

(Transizione: Dopo l'inquadramento teorico e normativo fornito da Mantelero, la lezione prosegue con l'intervento del Prof. Edoardo Raffiotta, che sposta l'attenzione sull'implementazione pratica dell'AI Act e sulle implicazioni organizzative per aziende e studi legali. Raffiotta – che tra l'altro siede nel Comitato strategico nazionale per l'IA –

offre una prospettiva orientata alla governance e al cambiamento culturale necessario per affrontare le sfide poste dall'IA.)

4. L'Implementazione dell'AI Act: Governance, Consapevolezza e Ruolo del Giurista - Prof. Edoardo Raffiotta

Edoardo Raffiotta, professore di diritto costituzionale ed esperto di diritto delle nuove tecnologie, concentra il suo intervento sul **processo di implementazione** dell'AI Act e su come questo possa diventare un **motore di cambiamento** positivo per gli studi legali e le imprese in generale. In qualità di membro del Comitato strategico per l'IA istituito dal governo italiano, Raffiotta condivide riflessioni pratiche su governance, consapevolezza tecnologica e ruolo del giurista nell'era dell'intelligenza artificiale. Il suo messaggio è ottimista ma anche sfidante: l'AI Act non va visto come un problema, bensì come parte della soluzione; la vera difficoltà sta piuttosto nell'adattarsi ed evolvere internamente per governare la tecnologia.

L'AI Act come soluzione, non come problema

Raffiotta esordisce smontando una critica frequente rivolta all'AI Act, ossia quella di essere l'ennesima **sovra-regolamentazione** che rischia di frenare l'innovazione. Al contrario, sostiene che le lamentele su un presunto eccesso di regole sono spesso **fuori bersaglio**. I problemi applicativi maggiori oggi, per chi opera con le tecnologie digitali, derivano piuttosto da normative **più datate** come il GDPR, concepite per un'era tecnologica diversa. Il Regolamento Privacy, pur fondamentale, è noto per la sua complessità e per oneri che hanno colto impreparate molte organizzazioni; anche normative di settore, come quelle sulla cybersecurity, pongono sfide notevoli. L'AI Act invece – afferma Raffiotta – **nasce già nell'era dell'IA** e, lungi dal creare problemi, sta in realtà offrendo **soluzioni**. La prima e più importante di queste soluzioni è la promozione della **consapevolezza** all'interno delle organizzazioni. In altre parole, l'AI Act costringerà (o forse “stimolerà” è il termine giusto) aziende e studi legali a prendere atto di quanta intelligenza artificiale stanno già usando, dove la stanno usando e con quali rischi. Questo **“reality check” tecnologico** è salutare: porterà molti, che finora magari hanno implementato strumenti di IA in modo informale e sporadico, a fare ordine e a introdurre procedure adeguate. Dunque, invece di vederlo come l'ennesimo ostacolo burocratico, l'AI Act andrebbe abbracciato come un'occasione per **maturare** e dotarsi di una **strategia** sull'uso dell'IA. Chi saprà muoversi per tempo ne trarrà un vantaggio competitivo, perché disporrà di processi più solidi e trasparenti, in linea con le aspettative di clienti e autorità.

La governance della tecnologia come parola d'ordine

Il cuore dell'adeguamento all'AI Act, secondo Raffiotta, può riassumersi in una parola chiave: **governance**. Il Regolamento spinge le organizzazioni a fare i compiti a casa in

termini di gestione interna della tecnologia. In pratica, le aziende e gli studi legali dovranno:

- **Mappare la tecnologia esistente:** prima di tutto, è necessario acquisire piena **consapevolezza di quanta tecnologia (IA compresa) si possiede** e per quali usi la si impiega. Può sembrare banale, ma non è raro che nelle realtà complesse nemmeno il management abbia un inventario chiaro di tutti i software, servizi cloud o sistemi automatizzati che vengono utilizzati quotidianamente. L'AI Act impone di fare questa ricognizione.
- **Inquadrare giuridicamente i sistemi:** una volta mappati gli strumenti, occorre **classificarli secondo la definizione di IA fornita dal Regolamento**. L'AI Act infatti definisce cosa si intende per "sistema di intelligenza artificiale" ai fini legali: non tutti i software avanzati rientrano necessariamente nella definizione. Bisognerà quindi capire quali strumenti interni sono effettivamente sistemi IA secondo la legge, e per ciascuno valutare se ricade in una delle categorie di rischio previste. Ad esempio, un semplice sistema di automazione di testi potrebbe non rientrare nella definizione di IA dell'Act, mentre un algoritmo di *machine learning* per l'analisi predittiva di cause legali sicuramente sì.
- **Definire ruoli e responsabilità:** l'adeguamento richiede anche di stabilire **chi, dentro l'organizzazione, fa cosa** in relazione all'IA. Servirà un approccio interdisciplinare: non basterà l'IT manager, né il solo ufficio legale; dovranno collaborare informatici, giuristi, responsabili della conformità, e probabilmente nasceranno figure nuove (come specialisti di etica dell'IA o responsabili AI per lo studio). Il Regolamento, ad esempio, richiede che chi utilizza sistemi ad alto rischio metta in atto misure di **sorveglianza umana**: andrà quindi individuato un responsabile di questa supervisione sui sistemi IA, con adeguata formazione.

In sintesi, **l'adeguamento alla normativa IA va oltre il semplice rispettare una lista di obblighi**: implica l'avvio di un processo organizzativo di comprensione, catalogazione e gestione della tecnologia. È un'attività di governance che, se affrontata con serietà, può migliorare il modo in cui lo studio legale funziona, rendendolo più consapevole e strutturato nell'uso degli strumenti digitali.

Il rischio più grande: l'immobilità

Raffiotta prosegue con una considerazione quasi provocatoria: **il rischio più grande, per i professionisti legali, è non fare nulla**. Richiamandosi alla Strategia Nazionale per l'IA (il documento programmatico italiano in materia), nota che il **primo pericolo identificato è proprio "stare fermi"**. In un contesto di innovazione rapida, l'immobilismo equivale a restare indietro. Gli avvocati non possono permettersi di ignorare l'IA o aspettare passivamente che tutto sia stabilito al 100% a livello normativo. Al contrario, **devono mantenere un atteggiamento di massima apertura** verso l'adozione delle nuove tecnologie. Ciò non significa ovviamente gettarsi acriticamente su ogni gadget hi-tech, bensì adottare l'IA con giudizio, utilizzando due strumenti come presupposti: **accountability** (responsabilizzazione) e **gestione del rischio**. Questi elementi devono essere visti come leve che permettono l'innovazione in sicurezza, non come freni. In altri

termini, l'avvocato dovrebbe dirsi: "Provo pure questo sistema di AI nello studio, ma contestualmente mi preoccupo di capire i rischi, adottare misure di controllo, e restare padrone del processo". Raffiotta insiste sul fatto che l'AI Act, paradossalmente, può aiutare in questo percorso perché fornisce proprio una **cornice di riferimento** per valutare e gestire i rischi tecnologici. Il vero nemico, sottolinea, è semmai la **mediocrità operativa** di chi non coglie le opportunità: rifiutare l'IA a priori per timore o inerzia potrebbe rivelarsi il più grande svantaggio competitivo nel medio-lungo periodo. In un mercato legale sempre più esigente, chi non utilizza strumenti avanzati rischia di offrire servizi meno efficienti e di qualità inferiore. Dunque, il messaggio è chiaro: meglio muoversi, magari con cautela, che restare immobili.

L'alfabetizzazione sull'IA come fondamento

Un punto chiave su cui Raffiotta converge – in linea con quanto accennato dai precedenti relatori – è la necessità di investire nell'**alfabetizzazione all'IA** (la cosiddetta *AI literacy*) all'interno della professione. Questo può essere visto come un vero e proprio **obbligo formativo emergente**. In pratica, gli avvocati (e il personale degli studi legali) devono acquisire una **cultura di base sull'IA**, comprenderne il funzionamento generale e i concetti essenziali. È fondamentale capire che l'intelligenza artificiale **non si riduce a ChatGPT** o a pochi strumenti famosi: dietro il termine c'è un intero ecosistema di tecniche e applicazioni diverse (dal *machine learning* nei motori di ricerca di sentenze, agli algoritmi di visione artificiale per analizzare documenti scansionati, ai sistemi esperti per il supporto decisionale, ecc.).

La **consapevolezza diffusa** sui meccanismi dell'IA porta due vantaggi. Primo, consente di **gestire meglio i rischi**: un professionista formato saprà, ad esempio, che non deve caricare dati sensibili su un servizio di traduzione automatica online senza garanzie, perché c'è il rischio di diffusione di informazioni riservate (tema della **condivisione dei dati**); oppure saprà che un eccesso di fiducia negli strumenti automatizzati può causare un **calo del pensiero critico** (se ci si abitua a non verificare più nulla "perché lo ha detto il computer"). Secondo, un avvocato alfabetizzato all'IA sarà in grado di **cogliere le opportunità** offerte da queste tecnologie in modo più efficace. Ad esempio, saprà identificare ambiti del proprio lavoro in cui l'IA può far risparmiare tempo (come la revisione di centinaia di pagine di discovery in un contenzioso) o migliorare l'analisi (come trovare correlazioni nei precedenti giurisprudenziali), mantenendo però il controllo sugli output. In ultima analisi, Raffiotta vede la formazione continua sull'IA come parte integrante della professionalità futura del giurista. Non si tratta di trasformare gli avvocati in ingegneri informatici, bensì di assicurarsi che abbiano la **"cassetta degli attrezzi" concettuale** per dialogare con i tecnologie, capire i limiti e i bias degli algoritmi e garantire un uso etico e efficace di questi strumenti. Solo così la categoria potrà mantenere **centralità e autonomia** nell'era digitale, senza subire passivamente le innovazioni ma governandole.

*(Transizione: A questo punto della lezione, dopo aver esplorato gli aspetti normativi e di governance, l'attenzione viene riportata sulla **pratica forense quotidiana**. L'Avv. Luciano Quarta interviene per calare i principi dell'AI Act nella realtà operativa dello studio legale, con esempi concreti di utilizzo dell'IA, obblighi pratici e uno sguardo alle competenze future necessarie agli avvocati.)*

5. L'AI Act nella Pratica Forense: Obblighi, Responsabilità e Scenari Futuri - Avv. Luciano Quarta

Luciano Quarta, avvocato con esperienza nell'implementazione di sistemi tecnologici negli studi legali, porta la discussione sul terreno concreto della **pratica forense**. Il suo intervento è ricco di esempi e fornisce una sorta di **roadmap pratica** per gli studi legali che vogliono adeguarsi all'AI Act senza perdere di vista l'operatività. Quarta inizia con una prospettiva storica sull'evoluzione tecnologica nella professione, passa poi a descrivere casi d'uso reali dell'IA in ambito legale evidenziando quando questi possano ricadere nell'"alto rischio", quindi elenca i principali obblighi pratici che gli studi legali (in qualità di *deployers* di IA) dovranno assolvere. Infine, riflette sulle **competenze ibride** che l'avvocato di domani dovrà possedere e delinea una possibile **tabella di marcia** per la compliance.

L'intelligenza artificiale come "salto quantico" nella professione

Quarta esordisce con un interessante parallelo storico: paragona l'avvento odierno dell'IA ad alcune **rivoluzioni tecnologiche** del passato recente che hanno rappresentato veri e propri "**salti quantici**" per la professione legale. Richiama ad esempio l'introduzione del **fax** negli studi, quella del **personal computer** e, negli anni '90, l'arrivo delle **banche dati giuridiche su CD-ROM**. Ognuna di queste innovazioni, vista retrospettivamente, ha cambiato radicalmente il modo di lavorare degli avvocati: il fax ha sveltito le comunicazioni rispetto alle lettere e ai tribunali; il PC ha rimpiazzato le macchine da scrivere e abilitato strumenti digitali; i database elettronici hanno permesso ricerche giurisprudenziali in ore anziché giorni. Allo stesso modo, **oggi l'IA sta diventando "nazionalpopolare"**, nota Quarta, nel senso che è divenuta accessibile a tutti e di utilizzo comune. Software e servizi basati su IA – un tempo confinati a progetti di ricerca o a grandi aziende – sono ora disponibili su larga scala: dagli assistenti virtuali negli smartphone alle piattaforme di document automation, fino ai chatbot legali. Questo significa che anche lo studio legale medio si trova davanti a strumenti che promettono salti di efficienza paragonabili a quelli visti con le tecnologie passate. Il messaggio di Quarta è che la professione forense ha già vissuto innovazioni dirompenti e le ha assorbite; l'IA rappresenta un ulteriore passo di questa evoluzione, potenzialmente ancor più trasformativo, e va affrontato con la stessa attitudine con cui si affrontò l'arrivo del computer: con apertura mentale e volontà di apprendere, senza però abdicare al senso critico professionale.

Scenari operativi: quando un'IA è ad alto rischio?

Entrando nel vivo, Quarta illustra **tre scenari concreti** in cui un avvocato potrebbe trovarsi a utilizzare sistemi di IA nel lavoro quotidiano, analizzando per ognuno se e perché si potrebbe parlare di **sistema “ad alto rischio”** secondo l'AI Act:

- **Ricerca giurisprudenziale avanzata:** immaginate un software di IA capace non solo di cercare sentenze rilevanti, ma di **identificare trend decisionali** e persino **suggerire strategie processuali** in base all'analisi dei precedenti. Un tale sistema – per quanto utile – sarebbe quasi certamente considerato **ad alto rischio**. Perché? Perché andrebbe a incidere potenzialmente sull'esito di cause legali (dando indicazioni su come impostare una difesa o un ricorso) e dunque potrebbe influenzare i diritti delle parti in giudizio. Un errore o un bias in questo contesto avrebbe conseguenze importanti. Quarta osserva che sistemi di questo tipo sono praticamente alle porte: non più fantascienza, ma prototipi già in fase di sviluppo avanzato.
- **Analisi contrattuale nelle due diligence:** un altro scenario è quello di **operazioni societarie di fusione e acquisizione (M&A)**, dove occorre analizzare migliaia di clausole contrattuali in decine di contratti. Ci sono IA addestrate per effettuare una sorta di “due diligence automatizzata”, evidenziando anomalie o rischi contrattuali. Anche questa applicazione, se usata in un contesto di decisioni economiche rilevanti, rientrerebbe nell'**alto rischio** secondo l'AI Act. Il motivo è simile al precedente: l'output dell'IA (ad esempio un report con clausole rischiose) può influenzare decisioni di investimento o di strutturazione di un affare da milioni di euro. Un errore dell'algoritmo potrebbe far perdere di vista un problema serio in un contratto, con danni ingenti. Dunque anche qui l'impatto potenziale richiede misure di controllo stringenti.
- **Chatbot di studio legale:** infine Quarta propone un esempio più semplice e quotidiano: un **chatbot installato sul sito web dello studio** che risponde alle FAQ dei clienti o che addirittura permette di fissare appuntamenti in automatico. Questo rientra tra i sistemi a **basso rischio**, in quanto fornisce informazioni generiche e non incide direttamente su diritti o decisioni legali sostanziali. Tuttavia, anche un innocuo chatbot informativo **sarebbe soggetto ad alcuni obblighi**, primo fra tutti l'obbligo di trasparenza (bisogna far sapere all'utente che sta parlando con un sistema automatizzato e non con una persona). Inoltre, andrebbero comunque rispettati requisiti di qualità del servizio: ad esempio, evitare che il bot dia consigli giuridici scorretti o violi la privacy degli utenti. Questo per dire che anche laddove l'IA non è “pericolosa”, c'è comunque una responsabilità di usarla con attenzione.

Questi esempi concreti aiutano a visualizzare cosa significhi, nella quotidianità forense, avere a che fare con sistemi di diverse categorie di rischio. La distinzione non è solo accademica: **sapere se uno strumento rientra nell'alto rischio o no fa la differenza** su ciò che l'avvocato (o il suo fornitore) dovrà fare per essere conforme alla legge.

Obblighi pratici per lo studio legale (deployer)

A questo punto Quarta passa in rassegna i principali **obblighi pratici** che ricadranno sugli studi legali in quanto **utilizzatori (“deployers”) di sistemi di IA**, specie se ad alto rischio. In parte sono adempimenti nuovi introdotti dall’AI Act, in parte sono estensioni di principi già noti:

- **FRIA – Fundamental Rights Impact Assessment:** prima di utilizzare un sistema di IA ad alto rischio, lo studio legale dovrà condurre e documentare una **valutazione d’impatto sui diritti fondamentali** (FRIA). Si tratta di un’analisi simile nella logica alla DPIA (Data Protection Impact Assessment, prevista dal GDPR per i trattamenti dati rischiosi), ma focalizzata su tutti i possibili impatti dell’IA sui diritti delle persone coinvolte. Ad esempio, se si adotta un sistema per valutare automaticamente il profilo di rischio di un cliente (pensiamo a un algoritmo che valuta la probabilità di vincere una causa), occorrerà valutare i rischi di discriminazioni, di errori e le misure per mitigarli, prima di metterlo in esercizio.
- **Sorveglianza umana:** la decisione finale **deve rimanere nelle mani di un professionista umano**. Questo non è solo buon senso, ma un obbligo esplicito: i sistemi di IA ad alto rischio devono essere affiancati da meccanismi di supervisione umana efficace. In pratica, l’avvocato non potrà dare carta bianca all’algoritmo; dovrà sempre esserci qualcuno che capisce e approva (o modifica) le conclusioni a cui giunge l’IA. Se ad esempio un software suggerisce una certa strategia legale, il legale umano dovrà vagliarla criticamente prima di proporla al cliente.
- **Trasparenza verso il cliente:** come anticipato, sarà necessario **informare in modo chiaro, specifico e comprensibile** l’assistito quando nello svolgimento della sua pratica lo studio utilizza strumenti di IA. Non basterà quindi dire in generale “usiamo software avanzati”, ma bisognerà comunicare – probabilmente già nel mandato scritto o in altri documenti – quali attività sono state automatizzate o supportate dall’IA. Ad esempio: “il contratto che le forniamo è stato anche revisionato da un algoritmo di analisi testuale per individuare clausole inusuali”. La trasparenza dovrà essere calibrata sul cliente medio: l’informazione dev’essere comprensibile e non eccessivamente tecnica.

Complessivamente, questi obblighi mirano a **rendere l’uso dell’IA uno strumento affidabile e sotto controllo** all’interno dello studio legale: c’è una fase di analisi preventiva (FRIA), una di gestione operativa (sorveglianza umana) e una di comunicazione esterna (trasparenza al cliente). Se seguiti diligentemente, garantiranno che l’introduzione dell’IA non intacchi i diritti dei clienti né la qualità del servizio legale.

L’avvocato di domani: competenze ibride

Un tema su cui Quarta pone enfasi è la necessità, per l’avvocato di domani, di sviluppare **competenze ibride**. Citando il celebre futurologo del diritto **Richard Susskind**³, afferma che i legali non dovranno diventare programmatori informatici, ma non potranno più neppure rimanere del tutto ignari del funzionamento della tecnologia. Servirà insomma una solida **alfabetizzazione sull’IA** (in linea con quanto detto da Raffiotta), affinata su

aspetti pratici: capire almeno a grandi linee **come “ragiona” un algoritmo**, dove può sbagliare, quali bias può incorporare e come leggere criticamente i risultati che produce.

Viene portato un esempio concreto e d’impatto: il caso del software **COMPAS**⁴ negli Stati Uniti. COMPAS è un algoritmo utilizzato in ambito giudiziario americano per valutare il rischio di recidiva dei detenuti. È divenuto tristemente noto perché si è scoperto che **presentava pregiudizi razziali**, tendendo a classificare come “ad alto rischio” i soggetti afroamericani più spesso di quanto facesse con altri, a parità di condizioni. Questo caso emblematico insegna che gli algoritmi possono ereditare (o amplificare) **bias** presenti nei dati di addestramento e produrre output distorti. L’avvocato del futuro deve saper fiutare queste problematiche: se, ad esempio, un software di valutazione del rischio propone un risultato sfavorevole al proprio cliente, il legale dovrà essere in grado di interrogarsi se per caso l’algoritmo possa avere una distorsione (e magari chiedere conto al fornitore del modello utilizzato). Inoltre, sottolinea Quarta, l’avvocato dovrà mantenere sempre **un approccio critico**: anche quando l’IA diventerà comune, il valore del professionista starà nella sua **capacità di validare, contestualizzare o correggere** l’output automatico. In sintesi, la figura dell’avvocato si arricchirà di queste competenze trasversali: non solo conoscenza del diritto, ma anche una certa familiarità con i concetti di base dell’IA e dei dati, in modo da fungere da **mediatore** tra tecnologia e giustizia.

Roadmap per la compliance nell’uso dell’IA

A conclusione del suo intervento, Quarta delinea una possibile **roadmap**, cioè un percorso a tappe, per accompagnare lo studio legale nell’**adeguamento graduale** alla nuova normativa e in generale a un uso corretto dell’IA:

- **Fase 1 – immediata: Mappare** i sistemi digitali già in uso nello studio e **verificare la compliance dei fornitori**. Questo significa fare un inventario di tutti i software e servizi basati su IA (o con funzionalità analoghe all’IA) che lo studio utilizza, dai più banali ai più sofisticati. Per ciascuno, cominciare a raccogliere informazioni dal fornitore su come il prodotto è conforme alle normative (es. ha previsto misure per la privacy? Rispetta requisiti di non discriminazione? Se è potenzialmente ad alto rischio, il produttore si sta muovendo per ottenere certificazioni?). In questa fase, insomma, lo studio prepara il terreno conoscendo ciò che ha e dialogando con chi gli fornisce tecnologia.
- **Fase 2 – entro fine 2025: Redigere una policy interna sull’IA e investire in formazione** per tutto il personale. Una volta capito quali strumenti si usano e a quali rischi possono associarsi, è opportuno scrivere delle linee guida di studio su come utilizzare l’IA. Ad esempio, la policy potrebbe stabilire che “per le ricerche giuridiche si possono usare questi strumenti, ma il risultato va sempre rivisto da un avvocato senior” oppure che “è vietato inserire dati sensibili di clienti in servizi cloud di IA non approvati dallo studio”. Contestualmente, tutti i collaboratori dovrebbero ricevere una **formazione specifica** su queste policy e più in generale sull’uso consapevole dell’IA (torna qui l’alfabetizzazione di cui

sopra). L'obiettivo di questa fase è creare **cultura e regole condivise** dentro lo studio.

- **Fase 3 – dal 2026 in poi: Implementare le FRIA** per i sistemi ad alto rischio, nonché **sistemi di logging e monitoraggio costante**. Questa è la fase operativa avanzata: quando l'AI Act sarà verosimilmente entrato in vigore, lo studio dovrà eseguire le valutazioni d'impatto (FRIA) ogni volta che introduce o utilizza un'IA ad alto rischio, documentandole per eventuali ispezioni. Inoltre dovrà mettere in piedi meccanismi di **registrazione e controllo** sull'uso dell'IA (per esempio, tenere traccia di quando il sistema viene utilizzato e con quali risultati, così da poter auditare il suo funzionamento nel tempo). È un'attività continua di **monitoraggio** e aggiustamento, che assicura che la compliance non sia vista come un evento una tantum, ma come un processo permanente.

Questa roadmap in tre fasi proposta da Quarta evidenzia un principio importante: **non si improvvisa dall'oggi al domani la conformità all'AI Act**, ma ci si arriva per gradi, partendo subito con ciò che è fattibile (mappare e formarsi) per poi farsi trovare pronti al momento dell'entrata in vigore delle norme più stringenti.

6. Governance dell'IA nello Studio Legale: Rischi e Opportunità - Prof. Pierluigi Perri

Pierluigi Perri, professore di diritto e avvocato in materia di cybersecurity e data protection, conclude la giornata con un approccio molto concreto, quasi da “manuale di bordo” per chi si appresta a utilizzare l'IA in uno studio legale. La sua esperienza nell'implementazione di strumenti di IA in una grande law firm gli permette di condividere spunti concreti su come governare questi sistemi, massimizzandone i benefici e minimizzandone i pericoli. Perri struttura la sua riflessione attorno a cinque punti: (1) l'adozione dell'IA non è più una scelta sul *se*, ma sul *come*; (2) l'IA eliminerà la “mediocrità” liberando i professionisti dai compiti ripetitivi; (3) le sette domande da porsi prima di adottare un'IA generativa; (4) i rischi professionali identificati dal CCBE; (5) alcune raccomandazioni finali per l'uso responsabile.

L'adozione dell'IA: non “se”, ma “come”

Perri esordisce con un'affermazione netta: **il problema ormai non è se adottare l'IA, ma come farlo**. L'uso di strumenti di intelligenza artificiale, infatti, è già una realtà di fatto: il mercato li offre e gli studi legali – chi più chi meno – li stanno già introducendo nei loro processi. Ignorare o bandire l'IA non è una strada percorribile, sarebbe come voler fermare il tempo. Pertanto la questione centrale diventa come **governare questo processo** di introduzione dell'IA. Perri sottolinea che ogni studio dovrebbe dotarsi di una strategia chiara per integrare l'IA nei propri flussi di lavoro in modo ponderato. Non ci si può limitare a lasciare che ogni singolo professionista inizi ad usare strumenti vari a macchia di leopardo: serve una visione coordinata. E, come evidenziato anche dai precedenti interventi, ciò passa dall'adozione di policy interne, formazione, controlli e così via. In breve, l'IA **va gestita proattivamente**, non subita. Il mercato sta “risolvendo”

il *se* (nel senso che rende inevitabile usare l'IA per restare competitivi), spetta al singolo studio risolvere il *come*.

La sfida della “mediocrità”: via le attività ripetitive

Un concetto chiave introdotto da Perri è che **l'intelligenza artificiale eliminerà la mediocrità**. Cosa intende esattamente? Parlando con una punta di provocazione, riferisce che tutte le attività professionali di **livello medio o ripetitivo** verranno gradualmente automatizzate e “spazzate via” dall'IA. In altre parole, le mansioni standardizzate, quelle a basso valore aggiunto creativo, saranno svolte in modo sempre più efficiente dalle macchine. E questo, afferma Perri, **non è necessariamente un male**. Significa che il **valore del professionista umano** si concentrerà su ciò che le macchine non sanno fare: le attività di alto profilo, strategiche, creative, relazionali. Ad esempio, l'IA potrà forse scrivere una bozza di contratto standard attingendo a modelli predefiniti, ma solo l'avvocato umano saprà capire quali clausole innovative inserire per tutelare davvero uno specifico cliente in uno specifico affare, o negoziare con la controparte tenendo conto delle sfumature commerciali. L'IA potrà riempire formulari e fare calcoli, ma il **giudizio critico** e l'empatia nel consigliare un cliente resteranno competenze umane insostituibili. Tuttavia, Perri aggiunge un corollario importante: per rimanere rilevanti, i professionisti dovranno **imparare a interagire efficacemente con lo strumento IA**. In pratica, l'avvocato vincente non sarà quello che si ostina a fare tutto “alla vecchia maniera”, né quello che delega tutto al computer senza capire; sarà invece colui che saprà usare l'IA come un **complemento** potenziante, e non come un mero sostituto. Ad esempio, un bravo avvocato potrà usare un algoritmo di analisi legale per ottenere in pochi minuti una panoramica di 100 sentenze, ma poi userà la propria esperienza per scegliere la linea argomentativa più persuasiva basata su quelle informazioni. La “mediocrità” di cui parla Perri, dunque, si riferisce al lavoro routinario fatto senza particolare valore aggiunto: sarà svolto dall'IA, e questo spingerà i professionisti a puntare sempre più in alto sulle proprie capacità distintive.

Sette domande per adottare l'IA in studio

Prima di implementare concretamente uno strumento di **IA generativa** (come potrebbero essere i sistemi tipo ChatGPT adattati all'uso legale), Perri suggerisce che un professionista o uno studio dovrebbe porsi **sette domande fondamentali**. Queste domande fungono da checklist etica e operativa per valutare la compatibilità dello strumento con i doveri professionali e la sicurezza dei dati:

1. **Qual è la base giuridica per il trattamento dei dati?** – In altre parole, su quale norma o consenso ci si basa per poter inserire eventuali dati (anche personali o sensibili) nel sistema di IA? Bisogna essere sicuri di non violare la privacy o altri obblighi di riservatezza.
2. **È stata fatta una valutazione d'impatto (DPIA/FRIA)?** – Analogamente a quanto detto da Quarta, occorre chiedersi se si è già valutato in anticipo l'impatto dell'uso di quell'IA, sia in termini di privacy (DPIA, *Data Protection Impact*

Assessment) sia in termini di diritti fondamentali in generale (FRIA). Se non lo si è fatto, va sicuramente messo in agenda.

3. **Come garantisco la trasparenza verso i clienti?** – Bisogna avere un piano chiaro su come informare i clienti che si userà (o si sta usando) l'IA e in quali attività. La trasparenza non può essere improvvisata: serve una formula e un momento preciso in cui comunicarlo (per esempio nel contratto di incarico).
4. **Come gestisco i rischi per la sicurezza delle informazioni (mie e dei clienti)?** – Questa domanda tocca il tema della **cybersecurity**: usare un sistema di IA significa spesso inviare dati a un fornitore esterno o comunque digitalizzarli. Occorre assicurarsi che il sistema sia sicuro, che i dati siano cifrati, che non vi sia rischio di violazioni o di utilizzi non autorizzati delle informazioni confidenziali dello studio.
5. **Come limito il trattamento ai soli dati necessari?** – Qui c'è un principio di **minimizzazione dei dati**: davvero devo inserire tutto il documento nel sistema di IA o posso limitarmi alle parti rilevanti? Meno dati si danno in pasto all'algoritmo, minori sono i rischi in caso di problemi. Questa è una buona prassi derivata anche dal GDPR.
6. **Come rispetto i diritti degli interessati (ad es. cancellazione)?** – Bisogna interrogarsi su come eventualmente far valere diritti come quello all'oblio o alla rettifica di dati inseriti nell'IA. Se, ad esempio, ho fatto analizzare dei dati personali di un cliente a un servizio di IA, come posso poi cancellarli se il cliente lo chiede? Il fornitore me lo consente? È una domanda importante soprattutto per i dati sensibili.
7. **Per cosa utilizzerò esattamente il sistema?** – Infine, Perri suggerisce di chiarire bene **lo scopo** per cui si adotta l'IA. Sarà un semplice supporto interno? Oppure lo si userà per generare documenti destinati ai clienti o al giudice (come una bozza di atto)? O addirittura per prendere decisioni strategiche? Definire l'uso serve a capire il livello di affidamento che si darà all'IA e quindi il grado di cautela da mantenere. Ad esempio, se il sistema è usato solo come assistente per dei riassunti, il rischio è minore; se gli si delegherà la stesura di interi atti, il controllo dovrà essere molto più serrato.

Porsi **in anticipo** queste sette domande significa evitare di trovarsi poi spiazzati da problemi non considerati. Perri le propone come una sorta di **lista di controllo etica** per l'innovazione: prima di fare il passo, guardati allo specchio e verifica di avere le risposte. Se qualcosa manca, meglio fermarsi e predisporre le tutele necessarie.

I rischi professionali evidenziati dal CCBE

Perri richiama quindi l'attenzione su un documento del **Consiglio degli Ordini Forensi Europei (CCBE)**, l'organismo che rappresenta gli avvocati a livello europeo. Questo documento ha identificato alcuni **rischi specifici** che l'uso dell'IA pone per la professione legale:

- **Rischi per la competenza professionale:** c'è il pericolo di un **affidamento acritico** sugli strumenti di IA che porti a un **depotenziamento delle competenze** (*de-skilling*) degli avvocati. Se ci si abitua a lasciar fare tutto alle macchine, col tempo

si potrebbe perdere abilità importanti (pensiamo alla capacità di fare ricerche giurisprudenziali “a mano” o di scrivere un contratto complesso da zero). Inoltre, un’eccessiva fiducia potrebbe portare a non approfondire più le questioni come si dovrebbe. Insomma, l’avvocato rischia di diventare un passacarte di output generati da altri, perdendo il suo valore aggiunto.

- **Rischi per il segreto professionale:** molti servizi di IA sono offerti da provider esterni (spesso big tech). Inserire dati di un cliente in questi sistemi potrebbe esporre informazioni coperte da **segreto professionale** o confidenzialità. Ad esempio, dare in pasto a un servizio cloud una memoria legale potrebbe violare l’obbligo di custodia, se i dati non sono protetti o se il fornitore li riutilizza per addestrare ulteriormente l’IA. C’è quindi un serio tema di **riservatezza** e protezione dei dati nello scegliere e usare questi strumenti.
- **Rischi per l’indipendenza professionale:** il CCBE evidenzia anche il potenziale di **dipendenza eccessiva dai fornitori tecnologici**. Se uno studio lega gran parte delle sue attività a un certo software o piattaforma, potrebbe trovarsi in balia di decisioni commerciali altrui (ad esempio un aumento improvviso dei prezzi di licenza, o la dismissione di un servizio). Inoltre, un forte condizionamento da parte della tecnologia potrebbe in qualche modo uniformare il modo di lavorare degli avvocati, riducendo quella autonomia di giudizio che è essenziale. Bisogna stare attenti a non far sì che “comandi” più l’algoritmo della nostra valutazione indipendente.

Menzionando questi punti, Perri ci ricorda che l’innovazione non porta solo opportunità ma anche **sfide deontologiche**. Ogni studio dovrebbe valutare tali rischi e predisporre misure per mitigarli: ad esempio, scegliere con cura fornitori che garantiscano la non commistione dei dati (per il segreto), mantenere programmi di training per i giovani legali affinché sviluppino comunque le competenze di base (per evitare de-skilling), e magari usare più di una piattaforma per non dipendere da un singolo ecosistema (per l’indipendenza).

Raccomandazioni finali per una strategia AI consapevole

In chiusura, il professor Perri fornisce alcune **raccomandazioni pratiche** che suonano come conclusioni operative rivolte ai colleghi avvocati:

- **Valutare attentamente i fornitori di IA e le loro garanzie:** prima di adottare un software o servizio di IA, è bene fare **due diligence** sul produttore. Quali misure di sicurezza offre? Che politiche ha sui dati inseriti (vengono cancellati? Riutilizzati?)? Ha qualche certificazione o aderisce a codici etici? Scegliere bene all’inizio può prevenire problemi dopo.
- **Implementare gradualmente, partendo dalle aree a basso rischio:** non c’è bisogno di rivoluzionare tutto in un colpo solo. Una buona strategia è **iniziare dagli usi più semplici e meno critici**. Ad esempio, si può partire introducendo un piccolo chatbot per le FAQ sul sito, o un tool che aiuti a smistare le e-mail in studio. Queste applicazioni, essendo a basso rischio, permettono di prendere confidenza con l’IA senza esporre lo studio a grandi pericoli. Poi, man mano che

si acquisisce esperienza, si può salire di complessità (strumenti per ricerche legali, per due diligence, etc.).

- **Non vietare l'uso dell'IA:** Perri avverte di non scegliere la strada della **proibizione totale** per timore. Vietare l'IA in studio sarebbe come vietare l'uso delle banche dati elettroniche quando comparvero decenni fa. Si rischierebbe di rimanere tagliati fuori. Meglio invece imparare a usarla in modo **produttivo e sicuro**. Questo implica, certo, stabilire regole (ad esempio: "ok usare ChatGPT, ma non per inserire dati dei clienti"), però senza chiudere la porta alle opportunità.
- **Investire costantemente in formazione:** ultima ma forse più importante raccomandazione, è quella di **formarsi e aggiornarsi di continuo**. L'IA evolve rapidamente; ciò che oggi è all'avanguardia domani potrebbe essere obsoleto, e nuovi rischi o strumenti emergeranno. L'avvocato deve quindi mettere in conto un apprendimento costante: partecipare a corsi, leggere linee guida, confrontarsi con esperti. Solo così potrà mantenere il **controllo umano** sulla tecnologia e un adeguato **senso critico** nell'adottarla. Questo investimento in conoscenza è in definitiva un investimento nella propria **capacità di restare rilevante e competitivo**.

Con queste indicazioni, Perri chiude il suo intervento e, simbolicamente, la sessione formativa.

7. Conclusioni – Non solo AI Act, ma governance dell'intelligenza artificiale

La lezione del 25 giugno 2025 ha offerto uno sguardo ricco e sfaccettato sulle implicazioni dell'intelligenza artificiale per gli avvocati. Da essa emergono alcuni **punti chiave** e linee guida operative che valgono la pena di essere ricapitolati:

- **L'AI Act non è l'unica norma da considerare:** gli avvocati dovranno continuare a tener conto di **tutte le normative applicabili**, in primis il GDPR e le regole sulla responsabilità professionale, che rimarranno pienamente in vigore accanto all'AI Act. La compliance in ambito IA sarà quindi parte di un mosaico più ampio di obblighi legali.
- **La trasparenza verso il cliente è un obbligo imminente:** indipendentemente dalle tempistiche di entrata in vigore dell'AI Act, il dovere di **informare i clienti sull'uso dell'IA** è già nell'aria – sancito dalla Carta dei Principi di Milano e presto dall'ordinamento statale – e richiede sin d'ora di pensare **come implementarlo** concretamente. Ogni studio dovrebbe definire una strategia: cosa comunicare, dove (mandato, informative) e con che livello di dettaglio.
- **La governance interna è la chiave di tutto:** adeguarsi alle normative sull'IA non sarà un semplice spuntare caselle, ma imporrà un **processo interno** fatto di mappatura degli strumenti, definizione di policy, formazione del personale e attribuzione di responsabilità. In altre parole, servirà un approccio organizzativo serio, che potrebbe alla fine migliorare l'efficienza complessiva dello studio oltre che garantirne la conformità.
- **Alto rischio = valutazione d'impatto:** se uno strumento di IA utilizzato (o previsto) rientra nella categoria "ad alto rischio", lo studio legale sarà tenuto a

condurre una **Fundamental Rights Impact Assessment (FRIA)** prima del suo impiego. Questo diventerà un passaggio obbligato simile a ciò che il GDPR già richiede per alcuni trattamenti dati: va messo in conto nella pianificazione dei tempi e delle risorse.

- **Competenze ibride e “AI literacy”**: il futuro avvocato non dovrà essere un tecnico informatico, ma dovrà possedere una **solida alfabetizzazione in materia di IA**. Ciò significa familiarità con i concetti di base (es. cosa sono i bias, cosa sono le allucinazioni, come funzionano i modelli addestrati sui precedenti) e capacità di utilizzare gli strumenti mantenendo un approccio critico. Questa sarà probabilmente una delle sfide formative più importanti per la professione nel prossimo decennio.
- **Il rischio maggiore è non agire**: infine, un monito trasversale è emerso da più voci – restare **inerti di fronte all’IA** è la scelta più pericolosa. Non cogliere le opportunità offerte da un uso consapevole e governato dell’intelligenza artificiale rischia di tradursi nel **più grande svantaggio competitivo** per uno studio legale moderno. Al contrario, chi saprà coniugare innovazione e responsabilità potrà trasformare l’IA in un prezioso alleato, migliorando la qualità e l’efficienza dei servizi legali offerti.

In definitiva, l’IA rappresenta per gli avvocati sia una sfida che una chance: **ripensare il proprio ruolo** alla luce di strumenti potentissimi ma imperfetti. La lezione ha messo in guardia dai rischi e indicato la rotta per navigare questo mare nuovo. Starà a ciascun professionista seguire queste coordinate – imparando, sperimentando e mantenendo saldo il timone dei principi etici – per approdare a una pratica legale arricchita (e non spodestata) dall’intelligenza artificiale.

Capitolo IX

Intelligenza Artificiale tra Privacy e Copyright

1. Introduzione

Il 2 luglio 2025 si è tenuta la nona lezione del Corso di Perfezionamento in Intelligenza Artificiale per Avvocati, dedicata a due pilastri della compliance nell'era dell'IA: la protezione dei dati personali (Privacy) e il diritto d'autore (Copyright). L'incontro, moderato dall'avv. Giuseppe Vaciago, era suddiviso in due sessioni – prima la Privacy, poi il Copyright – per esplorare in profondità le sfide e le opportunità che l'intelligenza artificiale pone in questi ambiti.

La sessione sulla **Privacy** ha messo a confronto il Regolamento Generale sulla Protezione dei Dati (GDPR) con il nascente Regolamento Europeo sull'Intelligenza Artificiale (AI Act), analizzando l'approccio basato sul rischio, le difficoltà di individuare basi giuridiche per l'addestramento dei modelli e le strategie di compliance che aziende e studi legali possono adottare. A seguire, la parte sul **Copyright** si è concentrata sulla tutela delle opere generate dall'IA, discutendo chi ne sia autore o titolare, come le piattaforme regolano l'uso dei contenuti e quali contenziosi stanno emergendo a livello globale. L'obiettivo formativo dichiarato era duplice: da un lato fornire agli avvocati strumenti per identificare e mitigare i rischi legali legati all'IA; dall'altro, evidenziare le nuove **opportunità professionali** che nascono dall'obbligo per tutte le imprese di adeguarsi a normative in rapida evoluzione. Di seguito ripercorriamo in ordine cronologico gli interventi dei relatori, ciascuno focalizzato su un aspetto chiave.

2. Rischi, Opportunità e il “Peccato Originale” della Privacy nell'IA - Avv. Giuseppe Vaciago

Introduzione alla doppia natura di IA e diritto: In apertura, l'avv. **Giuseppe Vaciago** ha sottolineato la duplice valenza del rapporto tra IA e diritto: l'IA è **fonte di rischi** per chi la utilizza, ma al contempo costituisce una **straordinaria opportunità** per gli avvocati. Da un lato, infatti, l'uso di sistemi di intelligenza artificiale genera nuove incognite legali – in primis sul fronte della privacy – di cui occorre essere consapevoli. Dall'altro, queste stesse incognite creano domanda di competenze legali specifiche, offrendo ai professionisti preparati la possibilità di proporsi come consulenti e formatori specializzati. Vaciago invita dunque a vedere nell'IA non solo un problema da gestire, ma anche un volano per rinnovare la professione forense.

Un obbligo che diventa opportunità: A supporto di questa visione positiva, Vaciago evidenzia un recente sviluppo normativo di immediata ricaduta pratica: l'articolo 4 dell'**AI Act**, la proposta di Regolamento UE sull'Intelligenza Artificiale, introduce per tutte le aziende l'obbligo di garantire una sorta di “*alfabetizzazione sull'IA*” (“AI literacy”) ai propri dipendenti. Ciò non si riduce a un addestramento tecnico sull'uso dei software, ma implica fornire una **formazione sui rischi legali dell'IA**. Questo obbligo – già in vigore

da febbraio 2025 secondo Vaciago – è accompagnato da sanzioni tutt’altro che simboliche in caso di inadempienza (si parla di multe fino a **35 milioni di euro o il 7% del fatturato aziendale**). Per gli avvocati, soprattutto quelli formati in materia di IA, questa novità rappresenta **un’occasione professionale concreta**. Le imprese, dovendo per legge istruire il personale sui rischi dell’IA, **avranno bisogno di esperti** che le aiutino a farlo: i legali possono quindi proporsi come formatori qualificati, trasformando un obbligo di compliance in un servizio di consulenza ad alto valore aggiunto.

GDPR e AI Act: due strade parallele? Nel proseguire, Vaciago traccia un interessante **parallelismo tra il GDPR e l’AI Act**, evidenziando come entrambi i regolamenti – pur riferiti a materie diverse – condividano **principi comuni**. In particolare, si fondano sull’**approccio basato sul rischio** (*risk-based approach*): le misure da adottare devono essere proporzionate al rischio effettivo, tenendo conto dello *stato dell’arte* tecnologico. Inoltre, entrambi pongono al centro il principio di **accountability** (responsabilizzazione): il titolare del trattamento deve essere in grado di **dimostrare e giustificare** le scelte fatte nell’uso dell’IA o dei dati. Questo schema, osserva Vaciago, di fatto trasforma molte disposizioni in “*norme penali in bianco*”, ossia norme generiche i cui dettagli applicativi devono essere “scritti” dagli stessi operatori. In altre parole, spetta al professionista (l’avvocato o il compliance officer) definire regole e cautele adeguate caso per caso, mettendo in gioco la propria competenza per colmare i vuoti lasciati intenzionalmente dal legislatore. Vaciago fa notare come ciò responsabilizzi ulteriormente chi consiglia sull’IA: bisogna saper adattare i principi a situazioni concrete, con la consapevolezza che queste scelte verranno giudicate **ex post** dalle autorità.

Il “peccato originale” dei big data nell’IA generativa: Affrontando il tema cruciale della **compliance**, Vaciago introduce quella che definisce la criticità fondamentale, il vero “*elefante nella stanza*” quando si parla di IA e privacy. I grandi modelli di **IA generativa** – come GPT-4 e simili – **non sono conformi by default alla normativa europea**. Questo perché alla base del loro addestramento c’è un “**peccato originale**” dal punto di vista privacy: **la mancanza di una base giuridica valida per l’uso massiccio dei dati personali** impiegati nell’addestramento. I modelli generativi sono stati infatti “nutriti” con l’intero scibile umano disponibile online, includendo enormi quantità di informazioni personali che milioni di utenti hanno condiviso sul web senza poter prevedere questo utilizzo. Vaciago cita un episodio concreto e illuminante: il suo collega **Matteo Flora** ha inviato a *Common Crawl* – un vasto archivio di pagine web usato come fonte di dati da aziende come OpenAI – una richiesta, ai sensi dell’art. 15 GDPR, per sapere se e quali dati personali lo riguardassero nei loro dataset di addestramento. *Common Crawl* ha opposto un netto rifiuto, **negando l’accesso ai dati** in aperta violazione della norma. Questo esperimento dimostra che **il problema è a monte**: se nemmeno chi fornisce i dati per l’IA riesce a rispettare i diritti degli interessati, viene da chiedersi se non fosse necessario ottenere un **consenso informato** o un’altra solida base giuridica **prima** di effettuare trattamenti così massivi. Del resto, il Garante Privacy italiano – ricorda Vaciago – fin dal marzo 2023 aveva sollevato proprio questo punto cruciale, avviando un dibattito che ormai è globale. Si discute infatti ovunque (UE, USA, ecc.) di come

regolamentare l'IA per **prevenire abusi**, con proposte estreme come moratorie temporanee sui progetti più avanzati. In sintesi, l'intervento introduttivo di Vaciago lancia un messaggio chiaro: **serve consapevolezza sui rischi legali intrinseci nell'IA**, ma anche **intraprendenza nel cogliere le nuove opportunità professionali**, perché chi saprà navigare questi mari turbolenti diventerà un punto di riferimento indispensabile per aziende e colleghi.

3. GDPR e AI Act: un'interazione complessa - Avv. Lucio Scudiero

L'Europa tra sfida tecnologica e iper-regolazione: Scudiero esordisce inquadrando il problema su scala macro. Richiama il *Rapporto Draghi* sulla competitività europea, che evidenzia come l'Europa stia accusando un **ritardo tecnologico** significativo nel campo dell'IA. Basti pensare – cita Scudiero – che circa *il 70% dei modelli di IA di base (foundation models)* sono sviluppati negli Stati Uniti, e che il mercato del cloud computing (infrastruttura chiave per l'IA) è dominato da tre colossi americani. Questa dipendenza dall'estero, unita alla frammentazione normativa interna (si contano oltre **100 leggi sul digitale e 270 enti regolatori** in Europa), rende difficile per le imprese europee competere in modo armonico. In altre parole, l'UE si trova nella posizione delicata di voler *regolare* a fondo le nuove tecnologie (per tutelare valori come la privacy), rischiando però di soffocare l'innovazione locale e ampliare il gap con le superpotenze tecnologiche. **L'IA rappresenta una sfida aperta:** trovare il giusto equilibrio tra **innovazione e tutela** è cruciale per non rimanere indietro.

“L'IA non esiste senza dati”: **dati personali ovunque?** Entrando nel merito del rapporto tra IA e protezione dati, Scudiero afferma una verità semplice ma fondamentale: **nessuna IA funziona senza dati**. La recente esplosione dell'IA è stata resa possibile proprio dalla disponibilità enorme di dataset e dalla riduzione dei costi di calcolo. Tuttavia – ed ecco il nodo – in Europa vige il GDPR, che è una legge specifica e stringente sul trattamento dei dati personali. **Il GDPR prevale:** qualsiasi uso di dati personali deve prima di tutto rispettare i suoi requisiti. Il nuovo AI Act, per quanto importante, *non costituisce di per sé una base legale* che autorizzi a usare dati personali. Dunque, chi sviluppa o utilizza IA deve comunque trovare una giustificazione conforme al GDPR per trattare i dati. Scudiero nota che le **autorità privacy europee** (riunite nell'EDPB) tendono a interpretare il concetto di *dato personale* in modo molto ampio: un dato può considerarsi veramente *anonimo* solo se è praticamente impossibile risalire alla persona a cui si riferisce. In tutti gli altri casi, anche se i nomi sono rimossi o offuscati, prevale la prudenza e il dato va trattato **come personale**. Ne discende che, in mancanza di misure robuste di anonimizzazione, **si presume che l'IA lavori quasi sempre su dati personali**, perché qualche rischio di re-identificazione, anche minimo, è presente.

Tuttavia, Scudiero evidenzia che di recente stanno emergendo **crepe in questa visione rigida**. Alcune decisioni hanno infatti adottato approcci più sfumati:

- In **Germania**, ad esempio, un tribunale (decisione di Colonia, maggio 2025) ha ritenuto **lecito** l'uso, da parte di Meta (la società di Facebook), dei contenuti *pubblici* postati dagli utenti per addestrare modelli di IA. I giudici tedeschi hanno valorizzato vari fattori: l'interesse economico dell'azienda (considerato legittimo), le misure di mitigazione implementate (ad esempio la possibilità data agli utenti di *opt-out*, cioè di tirarsi fuori dall'uso dei propri dati, e la tokenizzazione, cioè l'offuscamento dei dati identificativi), nonché la natura pubblica delle informazioni coinvolte. In pratica, hanno accettato *ex post* una prassi già diffusa, cercando di bilanciarla con gli interessi in gioco.
- Un altro spiraglio viene da una sentenza del **Tribunale dell'UE (2023)**: in quel caso è stato giudicato che alcuni dati *pseudonimizzati* (ovvero mascherati da codici alfanumerici) trasferiti a un consulente esterno potessero considerarsi **anonimi**, poiché il ricevente (un soggetto terzo, ad es. una società di auditing) non aveva la chiave per ricollegarli ai nomi reali.

Di contro, la **posizione del Garante italiano** resta più rigorosa. In un provvedimento recente, il Garante ha sanzionato una società (Thin S.r.l.) per un'attività di analisi dati simile, ritenendo che **la possibilità di isolare gli effetti relativi a un singolo individuo** – pur non sapendo esattamente chi fosse – basti a rendere i dati utilizzati *personali*. In altre parole, se tramite i dati di training si possono trarre conclusioni su un certo utente specifico (anche identificato solo come “utente X”), quei dati non sono veramente anonimi. Questa disparità di vedute fra autorità indica che il confine tra dato personale e anonimo, nell'era dell'IA, è ancora terreno di dibattito e potrebbe evolvere con la giurisprudenza.

L'interesse legittimo: unica via, ma stretta. Posto che i dati usati dalle IA in genere sono dati personali, resta da capire **su quale base giuridica** le aziende possano fondare questi trattamenti nel rispetto del GDPR. Scudiero è chiaro: **il consenso informato degli interessati è impraticabile** in contesti così vasti (sarebbe impossibile raccogliere il consenso di milioni di persone i cui dati compaiono magari in un dataset pubblico). Anche l'ipotesi di basarsi su un contratto con l'interessato regge solo in casi limitati e viene interpretata in modo restrittivo. **Di fatto, l'unica base plausibile diventa l'interesse legittimo del titolare.** Cioè, l'azienda deve poter dire: “uso questi dati per addestrare l'IA perché ho un interesse mio, legittimo, che prevale sui possibili impatti negativi sugli interessati”. Ma Scudiero mette in guardia: **non è una scorciatoia semplice.** Anzitutto, **l'interesse invocato deve essere lecito** e concreto. E qui casca c'è il primo problema, perché diversi provvedimenti recenti delle autorità hanno dichiarato *illeciti* i trattamenti operati da note piattaforme di IA. Vengono citati i casi del Garante italiano contro **OpenAI** (per ChatGPT), del blocco a **Replika** (un chatbot AI) e di interventi sul gruppo editoriale GEDI, ammonito dall'autorità per aver sottoscritto un accordo di condivisione del proprio archivio giornalistico proprio con OpenAI, in funzione di training del modello e di suo utilizzo da parte delle testate del gruppo. . Ciò rende difficile, per una qualsiasi azienda utente, sostenere che il proprio interesse a usare quelle stesse tecnologie sia legittimo. In secondo luogo, anche ammesso che l'interesse sia riconosciuto astrattamente valido, occorre fare un **bilanciamento rigoroso** con i diritti e

le libertà degli interessati. In gergo, serve una **valutazione di bilanciamento** ben documentata, in cui si dimostri che le garanzie adottate dall'azienda riducono i rischi per le persone e che questi non superano i benefici perseguiti. È un esercizio complesso e, sottolinea Scudiero, inevitabilmente caso-specifico. Quindi sì, al netto di interventi normativi volti a preconstituire ex lege una base giuridica (si pensi a cosa fa il ddl AI italiano nell'autorizzare la ricerca clinica AI-based), l'interesse legittimo al momento pare la *sola strada percorribile* per legittimare i trattamenti di dati nell'IA, ma è una strada in salita, irta di ostacoli e con esito incerto.

Strategie di compliance e accountability: A questo punto del suo intervento, Scudiero passa a dare consigli pratici, rispondendo in sostanza alla domanda: *cosa può fare un'azienda (o uno studio legale) per usare l'IA ed essere compliant con privacy e GDPR?* La chiave, afferma, è adottare un approccio basato sull'**accountability**, cioè sulla responsabilizzazione e la proattività. In pratica, ogni organizzazione dovrebbe dotarsi di un quadro di regole interne e cautele che dimostri la propria diligenza. Scudiero suggerisce **quattro azioni** concrete:

- **Policy interne chiare:** definire esplicitamente **in quali modi è ammesso utilizzare strumenti di IA** e per quali scopi è invece vietato. Ad esempio, una policy aziendale potrebbe proibire di inserire dati sensibili o informazioni riservate del cliente in servizi di AI generativa pubblici, oppure limitare l'uso dell'AI a certe funzioni aziendali. Stabilire linee guida evita usi impropri e dimostra alle autorità che l'azienda non lascia il tema all'improvvisazione.
- **Due diligence sui fornitori:** prima di adottare una piattaforma di IA di terze parti, è fondamentale **analizzare a fondo i contratti e le privacy policy del fornitore**. Bisogna verificare, ad esempio, se il servizio offre la possibilità di **opt-out** dall'addestramento (cioè di non far utilizzare i dati immessi dall'utente per migliorare l'AI). *OpenAI*, ad esempio, nelle licenze business prevede di default che i dati dei clienti **non** vengano riutilizzati per training, ma altre piattaforme potrebbero non offrire garanzie simili. Inoltre, va considerato dove risiedono i dati (questioni di trasferimento extra-UE) e quali misure di sicurezza sono dichiarate.
- **Formazione e informative:** non basta che gli esperti conoscano i rischi – occorre **alfabetizzare l'intera popolazione aziendale** sull'uso corretto dell'IA. Scudiero propone di organizzare sessioni formative per dipendenti e collaboratori, in modo che capiscano quali dati possono inserire in uno strumento di IA e quali no, quali sono i pericoli (es. fughe di dati, errori grossolani, dipendenza dall'AI) e come segnalare eventuali problemi. Parallelamente, servono **informative privacy trasparenti** rivolte sia al personale interno che, se del caso, ai clienti, per spiegare che l'azienda utilizza strumenti di IA in certe attività. Ad esempio, in ambito di lavoro ciò tocca temi di controllo a distanza: il dipendente va informato se le sue attività sono monitorate da sistemi automatizzati.
- **Valutazione d'impatto (DPIA):** infine, per utilizzi significativi dell'IA sarebbe buona prassi (e spesso obbligatorio) effettuare una **Data Protection Impact Assessment** – in italiano, *valutazione d'impatto sulla protezione dei dati*. Si tratta di un'analisi documentata dei rischi che uno specifico progetto di IA comporta

per i diritti delle persone e delle misure messe in atto per mitigare tali rischi. La DPIA serve sia come strumento interno di analisi, sia come prova di accountability in caso di controlli: mostra che l'azienda ha riflettuto ex ante su rischi e benefici, prendendo decisioni consapevoli.

Seguendo queste linee guida, sostiene Scudiero, anche una **PMI italiana** può iniziare a utilizzare l'IA in modo responsabile, riducendo al minimo i rischi legali. Certo, resta inteso che il quadro definitivo dipenderà anche dall'evoluzione normativa (l'AI Act è ancora in divenire) e dalle posizioni che assumeranno le autorità di controllo.

Tre rischi principali da tenere a mente: In chiusura del suo approfondimento, Scudiero sintetizza i **principali rischi** legali che un avvocato dovrebbe sempre considerare quando consiglia sull'adozione di sistemi IA:

- **Rischio regolamentare:** ovvero il rischio di **non riuscire a trovare una base giuridica solida** per i trattamenti di dati personali nell'IA, con la conseguenza di esporsi a violazioni del GDPR e relative sanzioni. È il rischio di fondo derivante da tutta l'incertezza normativa descritta sopra.
- **Rischio cyber:** i sistemi di IA e i dati che elaborano possono diventare bersaglio di attacchi informatici. Occorre dunque **proteggere sia i dati che i modelli** dagli attori maligni. Ad esempio, si teme il *data poisoning*: un agente esterno potrebbe inserire di proposito dati falsi o dannosi nel set di training per "avvelenare" il modello e farlo funzionare male o con bias voluti. Oppure, attacchi per estrarre informazioni riservate memorizzate nel modello. La sicurezza informatica diventa parte integrante della compliance.
- **Rischio di asimmetria informativa:** infine, Scudiero richiama un profilo etico-legale spesso trascurato, ossia il **dovere di informare correttamente gli utenti** quando interagiscono con sistemi automatizzati. Sia per i dipendenti interni sia per i clienti o consumatori, usare l'IA senza trasparenza crea un'asimmetria (il sistema sa molto di loro, ma loro non sanno di aver di fronte un sistema e non un umano, né come vengono prese le decisioni). Il rischio, quindi, è di violare principi di **trasparenza e correttezza**: per evitarlo, bisogna sempre comunicare – in modo comprensibile – che c'è un'IA in gioco e quali logiche generali la guidano, così che l'individuo possa regolarsi di conseguenza.

Con questa disamina, l'intervento di Scudiero fornisce una visione a 360 gradi della *complessa interazione* tra GDPR e AI Act: da un lato le tensioni e i gap normativi da colmare, dall'altro gli strumenti pratici per navigare l'incertezza in attesa che il diritto si assesti. È un invito a **non restare paralizzati dai dubbi legali**, ma a **muoversi con cautela** e intelligenza, perché l'IA è ormai una realtà inevitabile anche per il settore legale.

4. Opere generate dall'IA: tutela, piattaforme e contenziosi - Avv. Lucia Maggi

Creazione artificiale e diritto d'autore: serve un autore umano. La prima domanda posta da Lucia Maggi è quasi filosofica: *un'opera prodotta da un'intelligenza artificiale può essere protetta dal diritto d'autore?* In altri termini, **chi è l'autore** di un contenuto

generato dall'IA e quel contenuto gode di tutela come opera dell'ingegno? Maggi spiega che, a livello internazionale, la risposta che si sta affermando è **quasi unanimemente “antropocentrica”**: per avere tutela, nell'opera deve esserci un **contributo creativo umano** significativo. Diversi esempi lo confermano:

- Negli **Stati Uniti**, la posizione ufficiale del Copyright Office (l'agenzia federale sul diritto d'autore) è che **solo un'opera con un apporto creativo di un essere umano** può ottenere copyright. Se un'immagine, un testo o una musica sono generati interamente da un algoritmo senza scelte creative umane prevalenti, non possono essere registrati come opere protette. Questa linea è stata ribadita in un caso clamoroso, *Thaler v. Perlmutter*, in cui l'inventore Stephen Thaler cercava di accreditare un algoritmo (“Creativity Machine”) come autore di un'opera d'arte: la giustizia USA ha negato la registrazione, affermando che **la paternità umana è un requisito imprescindibile**.
- In **Italia**, sebbene la legge vigente non menzioni espressamente l'IA, il legislatore si sta muovendo nella stessa direzione. Un disegno di legge sull'IA (*DDL IA*) propone di modificare l'art. 1 della Legge sul Diritto d'Autore inserendo la parola “umano” accanto a “opera dell'ingegno”. L'intento è chiarire *a priori* che la creatività tutelata è solo quella umana. Anche la giurisprudenza italiana ha dato segnali interessanti: Maggi cita un caso di qualche anno fa (Cassazione *Biancheri vs. RAI*), in cui si discuteva la proteggibilità di un **fiore digitale** creato con l'ausilio di un software. Ebbene, la Cassazione riconobbe la tutela d'autore sull'opera, ma pose l'accento sul fatto che **l'autrice umana aveva fornito istruzioni complesse e dettagliate** al programma, usando quest'ultimo come mero strumento esecutivo della propria idea. In sostanza, *anche in Italia è tutelabile il risultato di un'IA solo quando dietro c'è una persona che dirige e plasma il processo creativo*. Se invece un'opera nascesse al 100% dall'autonomia generativa di una macchina, senza apporti umani originali, molto probabilmente non sarebbe considerata proteggibile.

Chi possiede l'output dell'IA? Le regole delle piattaforme. Dopo aver delineato il principio generale, l'avv. Maggi sposta l'attenzione su un aspetto pratico: *cosa dicono i termini di servizio delle principali piattaforme di IA generativa riguardo alla titolarità dell'output?* In altre parole, quando un utente usa strumenti come ChatGPT per testi o Midjourney per immagini, **di chi sono i diritti sul contenuto generato?** Maggi evidenzia che ogni provider ha fissato proprie regole contrattuali, spesso poco note agli utenti, e fa un breve confronto:

- **OpenAI** (produttore di ChatGPT) e **Anthropic** (produttore di Claude) hanno politiche simili: **riconoscono sempre all'utente la titolarità dei diritti sull'output** generato, sia che il servizio sia usato gratuitamente sia a pagamento. Quindi, se chiediamo a ChatGPT di scrivere un articolo e otteniamo un testo originale, secondo OpenAI *noi* ne siamo gli autori/titolari e possiamo usarlo liberamente (fatte salve ovviamente le possibili violazioni di diritti altrui, tema su cui torneremo).

- **Midjourney**, piattaforma popolare per generare immagini tramite prompt, adotta una distinzione netta: **solo gli utenti paganti diventano titolari pieni delle immagini create**. Chi usa la versione gratuita, invece, accetta che le immagini generate siano rilasciate sotto una licenza Creative Commons non commerciale (CC BY-NC 4.0). In pratica, l'utente free può usare l'output solo per scopi personali e non commerciali, e deve anche attribuirlo secondo le regole CC, mentre l'utente a pagamento acquisisce un diritto esclusivo di sfruttamento (come se fosse l'autore, nei limiti in cui l'opera è tutelabile).
- Un'altra differenza cruciale sta nell'uso degli **output per addestrare ulteriormente l'IA** (*feedback loop* di training):
 - OpenAI, per default, **riutilizza i contenuti generati dall'utente e le sue interazioni per migliorare il modello**. Tuttavia, offre un'opzione di **opt-out**: gli utenti (soprattutto business) possono chiedere che le proprie conversazioni non vengano conservate né utilizzate per training futuri.
 - **Runway** (piattaforma di generazione video/immagini) adotta una politica più invasiva: tutti i contenuti che l'utente crea o carica vengono **sempre usati per l'addestramento**, senza possibilità di opt-out.
 - **Midjourney** pubblica di default tutte le immagini generate sulla propria community gallery, rendendole visibili e *riutilizzabili per training*. Solo acquistando la cosiddetta **"modalità stealth"** (un abbonamento speciale) l'utente può generare immagini private e non condivise, evitando così che finiscano nel dataset pubblico.
 - **Anthropic** dichiara al contrario di **non usare per l'addestramento** i contenuti prodotti dall'utente, a meno che sia l'utente stesso a fornire feedback o segnalazioni utili.

Queste clausole, spiega Maggi, sono importanti da conoscere perché incidono sui **diritti pratici** degli utenti: chi lavora con output generati dall'IA deve capire se può sfruttarli commercialmente, se deve tenerli riservati e soprattutto se i dati che inserisce o i risultati che ottiene finiranno a loro volta nel "circolo" dell'addestramento alimentando l'intelligenza artificiale (magari anche a beneficio di altri utenti). In ottica legale, leggere i **Termini di Servizio** equivale a leggere il **contratto** che regola l'uso di queste tecnologie – un tema che tornerà nelle conclusioni.

Plagio involontario e uso illecito dei dati: cause in corso. L'ultima parte dell'intervento di Lucia Maggi offre una panoramica sui principali **contenziosi legali** esplosi recentemente attorno all'IA generativa nel campo del copyright. Si può notare che le dispute si sviluppano su **due fronti contrapposti**, in base a chi subisce il possibile danno:

- Dal **lato degli utenti/creatori** che usano l'IA c'è il rischio del *plagio inconsapevole*. Un utente potrebbe chiedere a un'IA di creare un'opera originale – per esempio un brano musicale, un disegno o un articolo – e ottenere invece (senza saperlo) un risultato che **viola i diritti di terzi**, perché troppo simile a un'opera esistente. Questo genere di violazione è subdolo: chi usa l'IA magari non riconosce la somiglianza e diffonde l'opera generata, ma se essa riproduce parti sostanziali di un'opera altrui protetta, l'utente ne risponde legalmente. **La**

responsabilità è oggettiva, osserva Maggi: il fatto di non averlo fatto apposta o di essersi affidati a una macchina non scusa l'infrazione. Un caso ipotetico: un illustratore chiede a un generatore di immagini di creare un personaggio "alla Topolino" e ottiene un disegno molto simile a Mickey Mouse; utilizzandolo in un libro, violerebbe il copyright Disney anche se è stato l'algoritmo a produrlo. Questo timore di "plagio da parte dell'IA" sta già causando allarme tra professionisti e aziende che vorrebbero usare contenuti generati automaticamente.

- Dal **lato dei titolari di opere originali** (artisti, case discografiche, fotografi, ecc.) c'è invece il tema dell'**uso non autorizzato delle opere protette per addestrare le IA**. Molti modelli generativi sono stati allenati su database pieni di immagini, testi, musica raccolti dal web, *spesso senza chiedere alcun permesso* agli autori. Ora che questo fatto è emerso, vari detentori di diritti si stanno muovendo legalmente sostenendo che le tech company hanno violato il loro copyright (o diritti connessi) incorporando le loro opere nei dataset. Maggi cita tre esempi di cause pendenti a livello internazionale:
 - **Universal Music vs. Anthropic**: la major musicale Universal ha citato in giudizio la startup Anthropic accusandola di aver utilizzato i testi di canzoni protetti da copyright per addestrare il suo modello linguistico (Claude). La controversia è interessante perché riguarda il *testo* delle canzoni – protetto da copyright d'autore – e pone la domanda se "leggere" quei testi per addestrare un'IA costituisca una violazione. Al momento, riferisce Maggi, le parti hanno raggiunto un accordo temporaneo di "*non belligeranza*", sospendendo azioni ostili mentre cercano una transazione.
 - **GEMA vs. OpenAI**: in Germania, la GEMA (l'equivalente tedesco della SIAE italiana) ha intentato causa contro OpenAI per l'uso non autorizzato di testi musicali protetti nei dataset di ChatGPT. OpenAI, da parte sua, si difende invocando l'**eccezione di "text and data mining"** prevista dal diritto UE: una norma che, a certe condizioni, permette di analizzare opere protette senza chiedere permesso, per finalità di ricerca e sviluppo. Sarà il tribunale a dover decidere se addestrare un sistema AI rientra in quella eccezione o no.
 - **RIAA vs. Stable Audio (Suno/Udio)**: negli USA, la Recording Industry Association of America ha fatto causa a due servizi di IA generativa musicale (noti come Suno o Udio) accusandoli di aver utilizzato **registrazioni audio protette** da copyright e diritti connessi (diritti degli artisti interpreti e dei produttori discografici) per addestrare algoritmi che poi generano brani musicali. Questo caso allarga la prospettiva ai **diritti connessi**, sostenendo che anche usare la *performance* registrata da un cantante o musicista nei dataset senza licenza costituisce violazione, non solo usare lo *spartito* o la melodia.

Nel complesso, la rassegna di Maggi evidenzia come l'avvento dell'IA stia innescando un'ondata di contenziosi inediti. Da un lato autori e detentori di diritti temono di **perdere controllo sulle proprie opere**, sia perché vengono imitate dall'IA, sia perché vengono sfruttate per addestrarla senza compenso. Dall'altro lato, **gli utenti delle IA devono stare all'erta**: affidarsi ciecamente a un algoritmo per creare contenuti può metterli in guai legali se l'algoritmo "bara" pescando troppo dal lavoro altrui. Il diritto d'autore, insomma, è diventato un campo minato che richiede nuove competenze: gli

avvocati dovranno saper orientare sia i creatori tradizionali, preoccupati di tutelarsi nell'era delle macchine creative, sia i nuovi creatori aumentati dall'IA, ignari dei rischi che corrono.

5. Il prompt e la paternità dell'opera: prospettiva USA - Avv. Cristiano Bacchini

La posizione ufficiale USA: serve l'autore umano. Bacchini parte dal quadro normativo americano, ricordando che il **Copyright Act** statunitense tutela le *“opere originali di autore fissate in un mezzo tangibile di espressione”*. Questa formulazione, per quanto risalente, viene generalmente considerata sufficientemente flessibile da poter includere anche opere generate con l'IA, ma **a una condizione fondamentale**: che *dietro l'opera vi sia un autore umano con sufficiente apporto creativo*. Il Copyright Office USA – l'ente competente per la registrazioni delle opere – ha chiarito in varie occasioni di considerare *inammissibile* la registrazione di contenuti interamente generati da processi automatici. **Un prompt semplice non basta** a rivendicare la titolarità e/o paternità di un'opera generata: se scrivo due parole su un generatore di immagini e ottengo un capolavoro visivo, non posso pretendere una tutela autoriale perché il mio input è troppo generico e l'output non riflette un contributo diretto della mia creatività. Questa impostazione, osserva Bacchini, è ormai la prassi negli USA (come confermato anche dal noto caso *Thaler vs Perlmutter* già citato da Maggi). Ma cosa succede se il prompt non è *semplice* bensì molto *articolato*?

Prompt complessi e creatività: un dibattito aperto. Bacchini spiega che la dottrina americana è attualmente **divisa** sul valore creativo dei prompt molto elaborati. Alcuni studiosi e operatori sostengono la tesi *“pro-prompt”*: un **prompt dettagliato e originale**, che incide in modo determinante sul risultato, dovrebbe poter conferire all'utente la paternità dell'opera generata. In quest'ottica, se descrivo minuziosamente una scena immaginaria fornendo all'IA uno script creativo e l'output riflette la mia visione, il contributo umano esiste ed è proteggibile. Diversa è la posizione di altri attori, che pone l'accento sull'**imprevedibilità del processo generativo**. Anche un prompt sofisticato, dicono i critici, non garantisce un risultato univoco: lo stesso identico prompt, dato due volte, può produrre immagini o testi differenti. Inoltre, il rapporto causa-effetto tra input dell'utente e output dell'IA è mediato da un modello complesso che *interpreta* a suo modo le istruzioni. **Il controllo creativo dell'utente è pertanto indiretto e insufficiente**, rendendo quantomeno dubbia l'imputazione della relativa paternità.

Ergo non si può parlare di vera paternità autoriale. Questo scambio di opinioni indica che **non c'è ancora consenso**: la legge formale non è cambiata, ma come verrà applicata in futuro ai prompt evoluti è materia in evoluzione.

La teoria della “paternità per adozione”: Nel panorama delle teorie emergenti, Bacchini menziona la proposta denominata *“authorship by adoption”*. Secondo questa teoria, **l'utente diverrebbe autore nel momento in cui sceglie e fa proprio uno tra i molti**

output generati dall'IA. Ad esempio, se chiedo a un'IA di generare dieci variazioni di un logo e poi **seleziono** quella che mi piace, adottandola come opera finale, la mia scelta costituirebbe un atto creativo assimilabile alla selezione di uno scatto da parte di un fotografo. Bacchini riferisce però come questa teoria sia fortemente **criticata** dalla maggior parte della dottrina, posto che rischia di estendere la tutela a risultati casuali o mere preferenze, snaturando i principi del diritto d'autore che richiedono una forma espressiva riconducibile a un'attività creativa umana. Inoltre, riconoscere la tutela autoriale solo perché si è *scelto* un output potrebbe portare a proteggere ogni minima variante generata, snaturando anche in questo caso i principi del diritto d'autore (che tutela forme espressive ben individuate, non la semplice preferenza di chi le utilizza). Al momento, dunque, la *paternità per adozione* non sembra aver convinto i legislatori né i giudici, ma è indicativa dello sforzo di trovare un modello interpretativo adatto all'IA.

Quando anche l'input è creativo: gli *expressive inputs*. Bacchini prosegue osservando come il quadro cambi quando **l'utente non fornisce solo un prompt testuale, ma anche un contenuto creativo proprio come input**. Ad esempio, alcuni sistemi permettono di caricare un **disegno fatto a mano, o una foto, o un brano musicale** e poi di farli rielaborare dall'IA con istruzioni testuali. In questi casi l'output è il risultato misto di un *contenuto umano preesistente* e di una trasformazione operata dall'IA su di esso. Ebbene, c'è almeno un caso negli USA in cui il Copyright Office ha accettato di registrare un'opera ottenuta con questo metodo. Si trattava di un'immagine generata partendo da un'altra immagine frutto di umana creazione e modificata via prompt: l'USCO ha concesso la registrazione, **ma** ha specificato che la tutela copre esclusivamente gli elementi riconducibili all'autore umano originario, rimasti *inalterati o identificabili* nell'opera finale. In altre parole, viene protetta la *“paternità pittorica umana inalterata”*, mentre tutto ciò che è frutto della rielaborazione dell'IA è escluso dalla tutela. Questo approccio “ibrido” riconosce dunque che l'IA può essere usata come strumento di editing o potenziamento di un'opera umana: la parte di creatività umana preponderante resta protetta, l'apporto della macchina no. Bacchini sottolinea che questa distinzione richiede comunque analisi caso per caso e non risolve tutte le incertezze, ma offre un modello interessante per gestire scenari in cui l'IA è integrata nel processo creativo artistico.

Verso una normativa italiana: il DDL e le soluzioni tecniche. Concludendo, l'avv. Bacchini riporta l'attenzione sull'Italia, dove – come accennato anche da Maggi – esiste un disegno di legge sull'intelligenza artificiale in discussione. Oltre alla modifica dell'art.1 della legge sul diritto d'autore già citata, questo **DDL IA** contiene alcune proposte concrete per affrontare le sfide emerse, cercando un equilibrio tra innovazione e tutela dei diritti:

- Si ribadisce che **saranno protette soltanto le opere con un intervento creativo umano rilevante e dimostrabile**. Ciò allineerebbe formalmente l'Italia al principio antropocentrico già descritto, dando base legale a quanto oggi è frutto di interpretazione.

- Si introduce l'**obbligo di watermark** per i contenuti generati dall'IA. In pratica, chi sviluppa sistemi di IA dovrebbe incorporare nei materiali prodotti (immagini, video, audio) un **marchio digitale** o altra traccia che indichi che l'opera proviene da un algoritmo. Questo per garantire *trasparenza* e permettere a chiunque di riconoscere se un contenuto è stato creato artificialmente – misura che aiuterebbe sia a contrastare le fake news o le manipolazioni (deepfake), sia a tutelare i creatori umani, evitando che contenuti AI vengano spacciati per umani o viceversa.
- Si prevede che **l'estrazione di dati da Internet per addestrare le IA sarà consentita solo nel rispetto della legge sul diritto d'autore** (sul punto si veda anche l'art. 70-quater della Legge Autori), quindi non in modo indiscriminato, e rispettando l'eventuale *opt-out* esplicito dei titolari dei diritti. Ciò significa ad esempio che, se un fotografo inserisce un'indicazione (metadato) che vieta l'uso delle proprie foto per training di IA, i progettisti dovranno onorarla. Questa proposta mira a dare ai creatori un potere di controllo sui propri contenuti in relazione all'addestramento, un po' come il **robots.txt** permette ai siti web di dire ai crawler dei motori di ricerca cosa possono indicizzare e cosa no.

Bacchini commenta che queste misure, se adottate, potrebbero fare dell'Italia uno dei primi paesi a dotarsi di una cornice specifica su IA e diritto d'autore. Resta da vedere come si coordineranno con le normative europee (AI Act e direttive copyright) e internazionali, ma indicano una direzione: **rendere riconoscibile l'IA e incanalarne l'uso nel rispetto della creatività umana preesistente**. In definitiva, l'intervento di Bacchini evidenzia quanto sia sfidante definire concetti tradizionali come "autore" e "opera" nell'era algoritmica, ma mostra anche che il diritto sta iniziando a reagire, tra prassi delle autorità e progetti di legge, per **colmare il vuoto giuridico** e dare certezza sia agli innovatori sia ai creatori tradizionali.

6. Conclusioni – Implicazioni pratiche per gli avvocati

In chiusura della lezione, i relatori hanno condiviso alcuni **punti chiave** emersi dalla giornata e le **implicazioni pratiche** per gli avvocati che operano (o intendono operare) nell'ambito dell'IA. Si tratta di concetti che riassumono il filo conduttore di tutti gli interventi, offrendo una bussola per orientarsi tra doveri e opportunità.

- **Privacy by design come obbligo (non opzione):** quando si implementano o si utilizzano strumenti di IA in qualunque contesto professionale, occorre **integrare da subito la tutela della privacy nel progetto**. Questo significa, ad esempio, effettuare precocemente una *DPIA* (valutazione d'impatto sulla protezione dei dati) per mappare rischi e rimedi, scegliere fornitori che garantiscano misure adeguate (come l'**opt-out** dal riutilizzo dei dati per training) e configurare i sistemi in modo da minimizzare l'esposizione di dati personali. La privacy by design non è solo buona prassi: diventa un requisito legale stringente e ignorarla potrebbe portare a violazioni difficilmente sanabili a posteriori.
- **Il legittimo interesse non è una scorciatoia facile:** come spiegato da Scudiero, usare il **legittimo interesse** quale base giuridica per trattare dati personali con

l'IA richiede **grande cautela e documentazione solida**. Non bisogna illudersi che sia un passe-partout: va svolta una rigorosa analisi di bilanciamento e soprattutto bisogna tenere conto dell'orientamento restrittivo delle autorità di controllo, che finora hanno mostrato scetticismo verso molti trattamenti legati all'IA. In pratica, un avvocato deve consigliare il cliente a *non abbassare la guardia*: il legittimo interesse va maneggiato come si maneggia nitroglicerina – con estrema attenzione – e preparandosi a difenderne l'uso in caso di ispezioni o reclami.

- **Paternità degli output generativi: prudenza e prove:** sul fronte del copyright, uno dei messaggi forti è che **la paternità di ciò che esce da un'IA è incerta** fino a prova contraria. Salvo poter dimostrare che c'è stato un apporto creativo umano significativo, dettagliato e identificabile, l'output di un'IA rischia di essere considerato **di pubblico dominio** (o comunque *non monopolizzabile* da chi l'ha generato). Questo ha enormi implicazioni: un'azienda che investe per creare contenuti tramite IA (per esempio design o testi per marketing) potrebbe scoprire di non avere diritti esclusivi su di essi, se contestati. L'avvocato deve quindi far capire ai clienti che la tutela delle creazioni AI è un terreno scivoloso: conviene, quando possibile, **integrare sempre un contributo umano creativo** (es. ritoccare manualmente un'immagine generata, o partire da materiale proprio) per rafforzare la proteggibilità del risultato.
- **Leggere (e capire) i Termini di Servizio:** come evidenziato da Maggi, le condizioni d'uso delle piattaforme di IA **sono il "contratto" vincolante** tra l'utente e il fornitore. Ignorarle espone a sorprese sgradite: si pensi a chi crede di essere proprietario di un output e invece scopre di aver solo una licenza limitata, o a chi inserisce dati riservati in un servizio cloud che poi li riutilizza. Diventa quindi imprescindibile, per gli avvocati, **analizzare a monte i TOS** delle piattaforme che il cliente vuole adottare, spiegarglieli in parole povere e, se necessario, **negoziare clausole specifiche** (nel caso di contratti enterprise) o orientare verso soluzioni più garantiste. In altri termini, l'avvocato deve fare da interprete e da filtro, assicurandosi che il cliente sappia a cosa va incontro quando clicca "Accept" su quei lunghi testi legali online.
- **Il rischio di plagio è reale:** usare sistemi di IA generativa **espone concretamente al rischio di violare senza volerlo diritti d'autore altrui**. Le piattaforme, come abbiamo visto, scaricano questa responsabilità sugli utenti tramite clausole che di solito escludono il loro liability se l'output infrange copyright di terzi. Dunque, all'atto pratico, chi utilizza un'IA per contenuti che poi pubblica o sfrutta commercialmente deve adottare un approccio simile a quello che avrebbe con un collaboratore umano: **controllare e revisionare** il risultato prima di utilizzarlo. Per un testo, ad esempio, si potrebbe sottoporlo a un controllo antiplagio; per un'immagine, verificare che non riproduca marchi o soggetti riconoscibili; per un codice software, accertarsi che non includa porzioni copiate da repository protetti. Consigliare questi step aggiuntivi ai clienti rientra nel dovere di diligenza del legale consulente.
- **L'opt-out è uno scudo fondamentale:** sia in ambito privacy sia in ambito copyright, è emerso come la possibilità di **negare il consenso all'uso dei propri dati o opere per l'addestramento** – quando esiste – sia uno strumento di tutela cruciale. Gli avvocati dovrebbero sensibilizzare aziende e creatori su questo

punto: ad esempio, un'impresa potrebbe inserire clausole contrattuali o impostazioni tecniche per esercitare il diritto di opt-out presso i fornitori di AI (come fatto con OpenAI nelle versioni business), in modo da proteggere i propri dataset interni o i dati dei clienti. Allo stesso modo, un fotografo o uno scrittore potrebbero voler indicare espressamente che le proprie opere **non** siano utilizzate per training, e avere mezzi legali per far valere questa scelta. L'opt-out diventa, metaforicamente, *lo scudo* che ogni titolare di dati o contenuti può alzare per non essere inglobato inconsapevolmente nella macchina dell'IA.

In conclusione, la lezione del 2 luglio 2025 ha mostrato come **intelligenza artificiale e professione legale** siano ormai intrecciate: all'avvocato si richiede di **comprendere la tecnologia**, anticiparne i rischi e guidare imprese e colleghi attraverso normative complesse e in divenire. Ma in questa sfida c'è anche una promessa: quella di **nuove opportunità professionali** inedite. Ogni azienda, grande o piccola, dovrà confrontarsi con la compliance AI, la protezione dei dati e la gestione dei diritti sulle nuove forme di creatività automatica. Chi saprà padroneggiare questi temi diventerà un consulente imprescindibile. Come evidenziato sin dall'introduzione da Vaciago, l'IA è al tempo stesso un campo minato e una terra fertile da colonizzare: *il segreto sta nel prepararsi, con la curiosità del pioniere ma anche con la prudenza del giurista*. In tal modo, gli avvocati potranno trasformare i rischi dell'IA in **valore aggiunto** per i clienti e in **evoluzione della propria professione**, anziché subirli passivamente.



ORDINE DEGLI
AVVOCATI DI MILANO

www.ordineavvocatimilano.it