

Resoconto integrale

Seduta del 9 ottobre 2025

Il giorno 9 ottobre 2025, alle ore 14,30, la Commissione Statuto e Regolamento, Partecipazione, Semplificazione amministrativa e innovazione digitale è convocata con nota prot. n. PG/2025/27752 del 02/10/2025, presso Sala B-C - Viale Aldo Moro 50, Bologna.

Partecipano alla seduta i consiglieri:

Cognome e nome	Qualifica	Gruppo	Voto	
PETITTI Emma	Presidente	PARTITO DEMOCRATICO - DE PASCALE PRESIDENTE	10	Presente
PRONI Eleonora	Vicepresidente	PARTITO DEMOCRATICO - DE PASCALE PRESIDENTE	5	Presente
SASSONE Francesco	Vicepresidente	FRATELLI D'ITALIA - GIORGIA MELONI	7	Presente
CALVANO Paolo	Componente	PARTITO DEMOCRATICO - DE PASCALE PRESIDENTE	5	Presente
CASADEI Lorenzo	Componente	MOVIMENTO 5 STELLE	1	Assente
CASTALDINI Valentina	Componente	FORZA ITALIA - BERLUSCONI - UGOLINI PRESIDENTE - NOI MODERATI	1	Presente
DONINI Raffaele	Componente	PARTITO DEMOCRATICO - DE PASCALE PRESIDENTE	2	Presente
EVANGELISTI Marta	Componente	FRATELLI D'ITALIA - GIORGIA MELONI	2	Assente
FIAZZA Tommaso	Componente	LEGA SALVINI EMILIA ROMAGNA - IL POPOLO DELLA FAMIGLIA	1	Assente
LEMBI Simona	Componente	PARTITO DEMOCRATICO - DE PASCALE PRESIDENTE	2	Presente
MASTACCHI Marco	Componente	ELENA UGOLINI PRESIDENTE RETE CIVICA	2	Presente
MUZZARELLI Gian Carlo	Componente	PARTITO DEMOCRATICO - DE PASCALE PRESIDENTE	2	Assente
PALDINO Vincenzo	Componente	CIVICI, CON DE PASCALE PRESIDENTE	2	Presente
PESTELLI Luca	Componente	FRATELLI D'ITALIA - GIORGIA MELONI	2	Presente
TRANDE Paolo	Componente	ALLEANZA VERDI SINISTRA - COALIZIONI CIVICHE - POSSIBILE	3	Presente
VIGNALI Pietro	Componente	FORZA ITALIA - BERLUSCONI - UGOLINI PRESIDENTE - NOI MODERATI	1	Assente
ZAPPATERRA Marcella	Componente	PARTITO DEMOCRATICO - DE PASCALE PRESIDENTE	2	Assente

Sono, altresì presenti, i consiglieri: Maria COSTI ed Elena UGOLINI.

Partecipano alla seduta Giusella Finocchiaro, professoressa ordinaria di diritto privato, di diritto di Internet e di Intelligenza Artificiale, Università di Bologna e Avv. Paola Manes, professoressa ordinaria di diritto privato, Università di Bologna.

Presiede la seduta: La Presidente Emma Petitti
Assiste la segretaria: Il Segretario Agata Serio
Funzionario estensore: Daniela Biondi

DEREGISTRAZIONE CON CORREZIONI APPORTATE AL FINE DELLA MERA COMPrensIONE DEL TESTO

- Audizione della Prof. Giusella Finocchiaro, professoressa ordinaria di diritto privato, di diritto di Internet e di Intelligenza Artificiale, Università di Bologna e della Prof. Avv. Paola Manes, professoressa ordinaria di diritto privato, Università di Bologna in merito all'intelligenza artificiale e al suo approccio strategico e consapevole.

Presidente PETITTI: Bene. Buongiorno a tutti e a tutte. Dichiaro aperta la seduta n. 11 della commissione VII che, come sapete, è dedicata all'audizione della professoressa Giusella Finocchiaro, professoressa ordinaria di diritto privato, di diritto internet e intelligenza artificiale presso l'università di Bologna e della professoressa Paola Manes, professoressa ordinaria di diritto privato dell'università di Bologna, sempre in merito all'intelligenza artificiale. Oggi parleremo anche dell'approccio strategico e consapevole.

Ricordo che abbiamo avviato come commissione VII un percorso di approfondimento, di audizioni. La scorsa seduta l'abbiamo realizzata, grazie all'intervento del professor Gardini, focalizzando l'intelligenza artificiale, l'impiego e l'utilizzo dell'intelligenza artificiale nella pubblica amministrazione e, quindi, cercando di capire insieme come si può migliorare l'efficienza, l'efficacia e la trasparenza dei servizi e ottimizzare anche la gestione dei dati nel rapporto con cittadini imprese. Voglio ricordare che come Assemblea legislativa un paio di anni fa abbiamo anche sottoscritto un protocollo con il Cineca, proprio per avviare questo percorso di approfondimento e oggi, appunto, con le professoresse Finocchiaro e Manes diamo seguito a questo approfondimento. Quindi, io cederei subito la parola alla professoressa Finocchiaro per la sua relazione e poi, a seguire, la professoressa Manes e apriamo poi il confronto, il dibattito con i consiglieri.

Grazie intanto per la vostra presenza qui, per noi è veramente un onore poter dialogare e ascoltare. Quindi, prego professoressa Finocchiaro.

Professoressa FINOCCHIARO: Grazie. Buonasera a tutti e a tutte. Grazie per questo invito che mi fa piacere e mi onora allo stesso tempo. Nella mia relazione cercherò di illustrarvi molto brevemente la legge italiana sull'intelligenza artificiale che entra in vigore domani, quindi, questa audizione è anche estremamente tempestiva perché parliamo di una legge non recente, ma addirittura recentissima, visto che la sua entrata in vigore è prevista per domani. È la legge 23 settembre 2025 numero 132.

La legge italiana in materia di intelligenza artificiale si situa in uno spazio che è già regolamentato perché, come tutti sanno, c'è un regolamento europeo sull'intelligenza artificiale, il cosiddetto AI Act, che già è entrato parzialmente in vigore. Il regolamento europeo si occupa di disciplinare l'accesso al mercato europeo dei prodotti di intelligenza artificiale. Quindi, volendo sintetizzare, questo è quello che cercherò di fare oggi, cioè, lasciarvi dei messaggi molto sintetici e chiari. Volendo sintetizzare il regolamento europeo sull'intelligenza artificiale è una legge sulla sicurezza del prodotto, cioè, indica quali requisiti deve rispettare un prodotto di intelligenza artificiale per entrare nel mercato europeo.

Al di là di questo, oltre a questo, la legge italiana cerca di entrare nel dettaglio, a un livello che naturalmente sotto il profilo delle fonti è più basso rispetto a quello del regolamento europeo e sostanzialmente si articola in quattro ambiti. Un primo ambito è ribadire i principi fondamentali. Ribadire, per esempio, che l'intelligenza artificiale e l'uso dell'intelligenza artificiale deve rispettare dei principi di trasparenza, di protezione dei dati personali, di non discriminazione, che l'approccio deve essere un approccio antropocentrico, quindi, con la persona al centro e così via. Quindi, una prima parte della legge è declaratoria, ribadisce i principi fondamentali.

Una seconda parte della legge effettua alcune scelte che sono demandate agli stati nazionali dal regolamento europeo. Per esempio, la governance dell'intelligenza artificiale. E l'Italia ha fatto una scelta piuttosto articolata che porta a una governance molto complessa, cioè, in Italia la strategia nazionale dell'intelligenza artificiale sarà gestita dalla Presidenza del Consiglio dei Ministri con l'apposito dipartimento; ci saranno due agenzie che si occuperanno di intelligenza artificiale e due autorità nazionali per l'intelligenza artificiale, che sono l'Agid (Agenzia per l'Italia digitale) e la ACN (Agenzia per la Cybersecurity) e poi ci saranno ben due comitati di coordinamento nazionali, perché bisognerà coordinare le due agenzie, con il Garante privacy perché ovviamente ci sono problemi di privacy impattati dall'intelligenza artificiale e lo abbiamo già visto sia con DeepSeek sia con ChatGPT, due provvedimenti del Garante che hanno bloccato questi due sistemi di intelligenza artificiale; ci saranno problemi di coordinamento con l'AGCOM, con l'antitrust e così via. Quindi, una prima scelta, dicevo, che effettua la legge nazionale è quella della governance.

Una seconda scelta riguarda la disciplina delle cosiddette Sandbox che sono molto importanti anche per questa regione. Le Sandbox sono spazi di sperimentazione tecnologica e insieme normativa e, quindi, visto che noi in questo momento storico non sappiamo ancora bene come utilizzare l'intelligenza artificiale, soprattutto non sappiamo quali saranno i problemi causati dall'intelligenza artificiale, possiamo solo immaginarli, ma non li conosciamo, sperimentiamo la tecnologia e insieme le regole, per vedere se le regole sono adeguate. Faccio un esempio per essere più concreta: vogliamo realizzare il gemello digitale di una qualche realtà, sia essa una realtà giuridica che fisica, il modello di una città per studiare come cambierebbero le condizioni se si cambiasse il traffico, come cambierebbe l'inquinamento se si cambiassero certe condizioni, quindi, supponiamo di realizzare un gemello digitale. La normativa che abbiamo, per esempio la normativa privacy, ci sostiene oppure è un ostacolo? Allora sperimentiamo: facciamo una Sandbox che riguardi un particolare problema, che riguardi una particolare questione, sperimentiamo la tecnologia e insieme le regole. Quindi, la legge italiana apre a questa sperimentazione prevedendo appunto una disciplina apposita.

Quindi, il primo ambito è quello, dicevo, dei principi, il secondo è quello delle scelte nazionali e il terzo, le nuove regole che sono introdotte dalla legge italiana sull'intelligenza artificiale. In materia penale ci sono delle aggravanti, se il reato viene commesso con l'intelligenza artificiale c'è il nuovo reato di Deep Fake. Ma ci sono anche due disposizioni molto importanti in ambito sanitario sulle quali vorrei richiamare la vostra attenzione.

Queste due norme, l'articolo 8 e l'articolo 9 della nuova legge, a cui ho dato un contributo nella redazione, servono per cercare di semplificare l'applicazione della normativa sulla protezione dei dati personali nel caso venga utilizzata l'intelligenza artificiale. Una prima semplificazione riguarda la ricerca scientifica effettuata su dati sanitari con intelligenza artificiale. Ci sono moltissimi progetti scientifici che impegnano i nostri ricercatori, non soltanto all'università, anche negli ospedali anche nella realtà mediche e sanitarie e sappiamo che uno dei problemi concreti è rendere legittimo l'utilizzo dei dati sanitari per fini di ricerca perché, fino ad oggi, non è l'unica strada possibile, ma è la strada maggiormente percorsa, si richiede il consenso del paziente nuovamente e questo spesso costituisce un problema perché i numeri necessari per una ricerca scientifica possono essere talmente ampi da rendere difficile la raccolta di un secondo consenso e quindi, in questo articolo 8, si stabilisce che, come previsto dalla Costituzione, quindi dovrei dire "si ribadisce che" la ricerca scientifica è di interesse pubblico; essendo di interesse pubblico, questa è una base giuridica sufficiente per trattare i dati personali e non è necessario richiedere un nuovo consenso, quindi, questa è una semplificazione importante che viene introdotta dall'articolo 8.

Si prevedono poi altre disposizioni molto tecniche sulle quali non entro, ma che sintetizzo come una semplificazione della normativa sulla protezione dei dati personali per scopi di ricerca in ambito sanitario.

Il quarto ambito della legge è quello che riguarda le deleghe che vengono conferite al governo sempre in materia di intelligenza artificiale, che riguarderanno la responsabilità civile, la responsabilità amministrativa, le norme sulla sull'addestramento dell'intelligenza artificiale e così via.

In estrema sintesi, ma sono stato veramente molto molto sintetica, però credo che forse questo sia di interesse per voi, questo è il nuovo quadro normativo previsto dalla legge italiana.

Due concetti sui quali forse vale la pena di soffermarsi. Uno, quello della strumentalità dell'intelligenza artificiale: viene ribadito più volte nella legge italiana che l'intelligenza artificiale è soltanto uno strumento e che, quindi, risponde dell'utilizzo dell'intelligenza artificiale chi la usa, quindi, l'avvocato che scrive un atto con l'intelligenza artificiale, se quell'atto contiene delle sentenze inesistenti - ed è successo - viene sanzionato o dovrebbe essere sanzionato, perché una sentenza (Firenze) non l'ha sanzionato e due sentenze (Torino e Latina) hanno sanzionato gli avvocati che avevano inserito sentenze che non esistevano nell'atto scritto con l'intelligenza artificiale, senza controllare. Lo stesso ovviamente vale per il medico, lo stesso vale per il giudice e, inutile a dirsi, vale anche per la pubblica amministrazione. C'è un articolo apposito che dice: la pubblica amministrazione può utilizzare l'intelligenza artificiale, ma ovviamente è soltanto uno strumento. Quindi è importante sottolineare questo concetto di strumentalità dell'intelligenza artificiale che conduce poi alla responsabilità per l'uso dell'intelligenza artificiale, responsabilità che deve essere umana.

Lo stesso vale per l'autorialità: se utilizziamo un sistema di intelligenza artificiale per scrivere un testo o anche per creare un'opera d'arte - come pure è successo - pensiamo A Refik Anadol e alle sue creazioni artistiche - l'autore o l'autrice è sempre la persona, non è intelligenza artificiale. Anche su questo ci sono tantissime sentenze, tutte in questa direzione, dalla Cina all'Australia, agli Stati Uniti, e la legge italiana ribadisce questo concetto chiarendo che l'ingegno è l'ingegno umano non dell'intelligenza artificiale. È spontaneo aggiungere "fino ad oggi" perché quello che ci riserva il futuro non lo sappiamo, però la nostra conversazione si ferma all'oggi e apre poi ad altre prospettive sul futuro.

Un'ultimissima breve considerazione, su un aspetto che io ritengo di particolare interesse per questa regione: il patrimonio informativo disponibile. Questa regione è una miniera di dati e di informazioni; ho avuto la fortuna di poterlo constatare da vicino, mi sono occupata di qualche progetto in passato che riguardava questo aspetto; l'intelligenza artificiale si nutre di dati, cioè, non si può utilizzare l'intelligenza artificiale se non si hanno dati a disposizione, dati che possono essere personali, cioè, riguardare persone fisiche o anche non personali e riguardare, per esempio, l'ambiente, il territorio, i sensori delle macchine industriali, l'inquinamento e quant'altro. Ecco, c'è un patrimonio enorme di dati e di informazioni che deve essere valorizzato anche per sviluppare applicazioni verticali di intelligenza artificiale.

Siamo in un momento storico in cui l'Italia non è più in grado, come non lo è l'Europa, di essere in prima linea per lo sviluppo tecnologico, gli Stati Uniti e la Cina ci hanno di gran lunga superato. Diciamo che in questa corsa stiamo partendo adesso, mentre altri sono già arrivati sotto il profilo strettamente tecnologico. Quello che però l'Italia e l'Europa possono fare è sviluppare dei verticali su aree specifiche in cui si può utilizzare l'intelligenza artificiale. Pensate ai beni culturali, pensate alle ricerche e alle sperimentazioni in medicina; noi abbiamo dei ricercatori dal punto di vista strettamente medico e sanitario che sono eccellenze nel mondo e che vincono progetti europei superando moltissimi altri ricercatori. Tutte queste applicazioni hanno bisogno di nutrirsi di informazioni, di dati, di riusare i dati che sono già nella disponibilità della pubblica amministrazione. Allora, siccome in questa regione le miniere di dati ci sono, il mio umile invito è di riflettere se da qui non si possano trarre effettivamente occasioni di crescita e di sviluppo anche per le nostre imprese e per i nostri ricercatori. Bene, io mi fermerei qui e vi ringrazio per l'attenzione.

Presidente PETITTI: Grazie davvero professoressa Finocchiaro per l'illustrazione, per la sua relazione e per gli spunti che ci ha offerto e sono sicura che poi ritorneremo in seguito anche grazie agli interventi dei consiglieri. Ora passo la parola alla professoressa Manes. Prego professoressa.

Professoressa avv. MANES: Buonasera. Grazie per questo invito che mi onora. Cercherò brevemente anch'io di incasellare, di innestare il mio discorso su tutto ciò che rappresenta una cornice data dalla professoressa Finocchiaro, che ha però avuto dei carotaggi già molto importanti, quindi, diciamo che se noi dobbiamo pensare all'edificazione della casa, la professoressa Finocchiaro ci ha dato, da un lato, le fondamenta e, dall'altro, quello che costituiscono i pilastri fondamentali dai quali non si può prescindere e quindi io mi muoverò nell'ambito di questo schizzo che è stato tratteggiato con delle pennellate importanti e ineliminabili.

Parto dal regolamento perché la professoressa Finocchiaro ci ha detto che la legge che domani entra in vigore riprende, amplifica, sottolinea ed esalta alcune caratteristiche del regolamento stesso; prima di tutto il fatto, e questo è l'esordio della legge europea, il pur tanto esecrato e criticato regolamento europeo, cioè, la centralità dell'intervento umano e quindi il fatto che nessuno algoritmo mai, che cirolerà nello spazio europeo, potrà confliggere con quella tassonomia di diritti fondamentali che rappresenta ormai ciò che è scolpito nella pietra dalla Carta fondamentale di Nizza. Questo, che è un movimento di pensiero che ha trovato, grazie al grande contributo, devo dire, di due maestri, cioè, Oreste Pollicino e la qui presente professoressa Finocchiaro, rappresenta quello che oggi si chiama convenzionalmente "costituzionalismo digitale", cioè l'aver ormai scolpito a caratteri maiuscoli il fatto che la tecnologia, per quanto trasformativa, non può elidere quell'autonomia dei diritti fondamentali, che ci consegnano le carte europee. Bene, ci dobbiamo muovere nel quadro di prendere questa come premessa fondamentale e capire che, al di là delle tante critiche che si possono muovere a questa regolazione Monster, perché lo è effettivamente, quasi una saga - hanno commentato alcuni - alla fine della quale non si capisce neanche come vada a finire perché - ed è stato questo sottolineato da commentatori ben più illustri di me - il regolamento non tocca palla su due fondamentali - per tutti noi, aziende, cittadini, stakeholders, consumatori, utilizzatori finali, enti pubblici - su due argomenti fondamentali, in primis, chi risponde quando qualcosa accade, quando un danno, su un bene. Su questo solo un punto interrogativo e nessuna risposta, dopo tante pagine di regolazione molto intrusiva e molto di dettaglio. Secondo, e questo è chiarissimo da quel rapporto Draghi che tutti noi dovremmo leggere tutte le mattine, c'è la pagina 26 che io infatti ho con un numero stampato in mente, il rapporto Draghi che ci ha consegnato la drammatica e schiacciante verità dei fatti, ci dice che ci sono 170 regolatori solo sullo spazio digitale, che addirittura oggi il digitale è un ambiente più normato del già pur normatissimo spazio giuridico degli operatori finanziari che, fino a ieri, pensavamo essere gli enti più normati in assoluto. Ebbene, il digitale è ancora più normato e Draghi, che non è certo sospettabile di anti-europeismo, ci ha detto che così nessuno sopravvive né chi fa Machine Learning, algoritmi, né chi li utilizza, né aziende, né enti pubblici, né il privato cittadino: nessuno può sopravvivere ad una pressione regolatoria tale.

Non solo - e qui è attualissimo il contributo della professoressa Finocchiaro su una legge che appunto ancora deve emettere i primi vagiti, domani - la possibilità che gli Stati membri abbiano di elevare ed innalzare addirittura il peso regolatorio della cornice europea. Cosa che poi effettivamente viene pienamente assunta ed eseguita e declinata dagli Stati membri, almeno il nostro l'ha sposata in pieno. Quindi, non solo il quadro regolatorio già così intrusivo e dall'altra parte, così evasivo per risposte fondamentali come quelle "chi si assume il rischio quando qualcosa va male nelle mani dell'utilizzatore finale", "da chi devo, nella lunghissima catena di tutti coloro che intervengono nel processo, da chi l'algoritmo lo pensa e scrive le righe di codice, perché l'algoritmo è una riga di

codice, e chi poi invece ce l'ha in mano, anche semplicemente per fare una ricerca con Chat GPT, quanto indietro devo risalire in questa catena e dove invece mi devo fermare per indagare". Con il requisito della colpa? Nessuno lo sa. Ad oggi non c'è, al di là di proposte interpretative su come il paradigma della responsabilità debba essere attuato, ci sono varie opzioni, ma non c'è una soluzione univoca che dia una risposta al medico che, in sala operatoria, con un blink degli occhi fa arrivare l'immagine di una tac o di una risonanza e tagli un segmento dell'addome diverso da quello che avrebbe tagliato, qualora l'input dell'intelligenza artificiale non gli fosse arrivato in tempo reale e poi però si scopre che invece quell'incisione doveva essere fatta pochi millimetri sopra o sotto.

Chi risponde? Perché i medici in una recente -- che è stata fatta da Deloitte, che ha campionato centinaia di migliaia di medici, hanno detto che, se è così, loro scappano dalla tecnologia e non ci meravigliamo, già - poveretti - sono oggetto di cause di malpractice che li soffoca da decenni, dicono proprio "ci manca pure che ci si metta l'intelligenza artificiale".

E allora, se l'esito deve essere quello della fuga della tecnologia, sicuramente abbiamo sbagliato direzione. Perché invece noi vogliamo una tecnologia che vada a braccetto con l'uomo e, in questo senso, questo è ribadito molto opportunamente dalla legge che entra in vigore domani, dove si dice che i professionisti intellettuali, non solo rispondono della prestazione e che questa prestazione rimane intellettuale se la prevalenza è della prestazione umana, cioè, se l'uomo è impegnato in prevalenza rispetto alla macchina. Che mi sembra un criterio e una norma di grande saggezza e di grande opportunità. E così, l'avvocato che vada - come ricordava la professoressa Finocchiaro - in udienza e abbia scaricato sentenze fantasma, si deve oggi aspettare - e non può più pensare di non saperlo - che quella non è più una prestazione intellettuale o non lo è mai stata.

Quale scenario o quale opzione abbiamo oggi, dato che di queste regole non possiamo fare a meno perché noi non possiamo andare in deroga a regole che ci siano state, ovviamente, indicate ed imposte a livello unionale. Non è la prima volta che i fruitori di queste regole si devono immaginare, in un contesto che oggi è, come si dice anche tra i giuristi, liquido, quindi di fonti liquide, di fonti mobili, che si possa immaginare un intervento, diciamo, di combinazione tra regole esistenti e, invece, regole di dettaglio che le aziende, in primis, questa, una realtà come quella nella quale tutti voi operate tutti i giorni, possano declinare in dettaglio regole in modo diverso da quello che è stato indicato o non è stato indicato affatto a livello unionale.

E faccio un esempio per essere più intellegibile: ci sono delle norme, due normette quasi dimenticate, ma molto importanti, verso la fine del regolamento europeo che parlano di codice di buone pratiche e di codici etici. Sono delle norme che, se sapute impugnare in modo sapiente dalle aziende, possono costituire un'ottima soluzione per i tanti problemi che oggi dobbiamo affrontare rispetto ad una legge europea che lascia insoluti problemi maggiori, come quello della responsabilità. I codici di buone pratiche e i codici etici consentono alle aziende e vivono e si nutrono - lo dice chiaramente la legge - del contributo di tutti gli stakeholders, quindi, dei cittadini, quindi di tutti coloro che da utilizzatori finali possano aver subito un'applicazione distorta, discriminatoria, non etica, non consapevole dell'intelligenza artificiale, possano contribuire a redigere delle regole che sono ritagliate sulla propria realtà e, sottoponendole al vaglio di un organismo che ne possa vagliare l'omogeneità, farsi candidati a un codice etico, un prontuario di regole che possa integrarsi con la cornice europea.

Dove sono più necessarie queste regole? Dove è necessaria quella regola di dettaglio che, al di là della corposissima regolazione che ci dice tutta la tassonomia dei rischi che, per esempio, la fascia di algoritmi che ci interessa di più, quella che espone effettivamente gli utilizzatori al rischio massimo, ci interessa avere come regola di dettaglio: sono tutte quelle regole che, per esempio, comportano cambiamenti organizzativi nell'ambito dell'ente che deve appunto utilizzare gli algoritmi e devono implementare delle pratiche di mitigazione del rischio, cosa di nuovo sulla quale il regolamento europeo, ma anche la legge che entra in vigore domani, non prende posizione, ma di

cui tutti hanno bisogno, da una piccola azienda che utilizza l'intelligenza artificiale per sostituire i tornelli all'ingresso o gestire la mobilità dei dipendenti, fino al grande istituto di ricerca, come i tanti che abbiamo, gli IRCCS, su questa regione che devono appunto utilizzare i dati che entrano in una PET tutti i giorni, per poi vedere come questi dati debbono essere lavorati dagli algoritmi per avere risultati accurati.

Le aziende si nutrono e devono poter classificare e modellare i rischi perché un rischio non modellato, cioè, un rischio di cui non si conosce la magnitudo e l'entità, è per ciò stesso respinto da un ente perché l'azienda deve poter misurare i rischi; mettersi in azienda un rischio non misurabile è un controsenso, ma non per il Risk Management, non per la Compliance, non per chi deve vigilare sull'osservanza delle regole, ma proprio perché non è possibile far correre alle aziende e agli enti, pubblici o privati che siano, rischi non misurati e misurabili. Poi si potranno appostare in bilancio, ci potranno essere le poste che coprono, ci potranno essere le assicurazioni e le polizze che coprono, ma nessuna compagnia assicurativa copre un modello, un ALM di cui non si sa la portata del danno sui cittadini. Provate a trovare una compagnia che voglia rispondere o debba rispondere di un danno dalla magnitudo non classificabile, non limitabile. Quindi, anche la risposta che a volte viene data per cui, rispetto a quel buco di tutela, vuoto di tutela che il regolamento lascia su come si risponde a dei rischi, e una possibile soluzione è legata proprio alla classificazione sulla base della prossimità del rischio, una delle possibili soluzioni in termini di responsabilità, più sei vicino al rischio che quell'algoritmo ti comporta e più devi rispondere, anche su questo, la soluzione deve ancora essere trovata sul fatto di chi poi copra quel rischio, perché la catena delle responsabilità, a ritroso, può essere molto lunga e non è detto che sia, anzi, è quasi impossibile, soprattutto oggi con l'intelligenza generativa, possibile risalire a chi in effetti risponde. L'onere della prova è qualcosa che è contro intuitivo, è completamente contrario a come funziona la generativa oggi, ma anche molti modelli di regressioni non lineari.

Quindi, certamente è facile, con una prima immagine che abbiamo tutti in mente, dell'albero che ha due sole diramazioni, capire che si risale al ramo comune, ma quando poi le diramazioni diventano una foresta, è impossibile risalire all'onere della prova e allora nessuno si assumerà di nuovo una responsabilità oggettiva per sé, di cui non conosce i limiti.

Quindi, diciamo due possibilità: da un lato, ci sono ordinamenti, ci sono paesi, ci sono aziende che, come sapete, hanno fatturati pari o superiori a stati sovrani, che hanno già deciso di mettersi al di sopra delle regole e sono coloro che tutti i giorni all'-- office, vanno a negoziare, di fatto, esenzioni da responsabilità e, anche senza negoziare, le ottengono di fatto, ma le ottengono molto spesso sulla nostra pelle.

Il giustissimo problema della sovranità o indipendenza digitale, al quale accennava la professoressa Finocchiaro prima, è un tentativo di non essere dipendenti da aziende private, perché di questo si tratta, perché non sono stati sovrani, che di fatto, ancora oggi, dispongono di quei modelli, dispongono della potenza di calcolo e della rapidità per lavorare con i dati, senza i quali noi tutti non possiamo di fatto operare.

Da un lato, quindi, lo scenario di una utilizzazione da far west, cioè, di un'utilizzazione indiscriminata senza controllo. Questo l'America l'ha fatto per tantissimi anni; i grandi grandissimi operatori lo fanno e lo continueranno a fare; dall'altro, la iper-intrusività, la capillarità di regole, come diceva Draghi, già vecchie prima di entrare in vigore. Draghi dice, sempre in quel rapporto: state attenti perché tutto quello che avete scritto su quei modelli che vi fanno rischiare così tanto, è già vecchio, perché le aziende si trovano intrappolate in una tassonomia di modelli di intelligenza artificiale che o non hanno o hanno superato o non avranno. E dunque, dice lui: la definizione, già avere sclerotizzato e imbrigliato quei modelli, è già un autogol perché, come giustamente oggi, invece, le Sandbox fanno nella legge che entra in vigore domani, bisogna procedere - dice l'economista Joshua Gans - per piccoli esperimenti retrattabili, piccoli tentativi che da domani si potranno, e tutte le

aziende potranno fare nelle Sandboxes, che siano reversibili, retrattabili, cancellabili, senza avere esposto la collettività a danni il cui impatto non possiamo neanche misurare né stimare.

Reputazionale. Oggi sul Sole 24 Ore c'era un piccolo trafiletto destinato a mettere in luce quanto una famosa società di consulenza - è un dato pubblico, quindi non dico notizie riservate - Deloitte abbia dovuto retrocedere l'ultima tranche di una commissione del lavoro che il governo australiano aveva commissionato, 440.000 dollari, perché questo lavoro era basato su informazioni false al 70, al 50, al 40, non importa, il danno reputazionale è stato tale e tanto per questa società di consulenza da impegnarsi, lei stessa, a retrocedere l'ultima tranche di compensi, pur di non subire un danno reputazionale che, essendo una società globale, le avrebbe fatto pagare un prezzo molto più alto.

E allora, questo che succede dalle ricerche che noi possiamo affidare al primo dottorando che abbiamo e collabora con noi, a quanto possa accadere in una grandissima società di consulenza, una grande società, essere o avere, non è l'unico modo questo di utilizzare l'intelligenza artificiale. Chi fa, chi gestisce i rischi all'interno delle società e chi gestisce le procedure e le pratiche di compliance, è già abituato a misurare e stimare i rischi. Allora, a queste aziende si dice che non è possibile perché il rischio dell'intelligenza artificiale è un rischio illimitato e ingovernabile. Ma se tu scegli di quali meccanismi, di quali algoritmi puoi e vuoi per ora avvalerti, si parte per piccoli passi, cioè, quindi, il processo destinato ad implementare gli algoritmi all'interno di una pubblica amministrazione, di un ente, di una società, non è ON-OFF, non è da zero a cento in una settimana; è un processo e deve essere un processo graduale. Gli algoritmi che, anche piccoli piccoli, consentano anche di ottimizzare ed efficientare un solo processo dell'ente possono però essere autorizzati e certificati come accurati, cioè, i cui risultati possano, risalendo all'indietro, in un ambiente diverso, essere replicati, perché algoritmo affidabile vuol dire algoritmo replicabile ottenendo lo stesso risultato. Partiamo con quelli piuttosto che partire con l'immenso mondo dell'intelligenza artificiale sul quale nemmeno i tecnici, cioè, i fisici applicati, i matematici, gli informatici e i Data Scientist, sono ancora d'accordo sul significato. Nessuno sa questa endiadi intelligenza artificiale che cosa abbia dentro o, se lo sa, ha poco senso definirlo perché, come spesso è stato detto, quella definizione è obsoleta già quando l'abbiamo data.

E allora, forse, l'idea sarebbe proprio quella di partire con pochi modelli osservabili e replicabili, il cui processo possa essere mappato con gli strumenti che già, ne faccio uno per tutti, il metodo Montecarlo col quale gli econometrici modellano i rischi, almeno pari, se non più grandi, cioè, i grandi danni catastrofali del clima, vi faccio un esempio per tutti; già sono capaci di modellare quei rischi. Allora, nelle grandi compagnie assicurative una prima ipotesi è quella di sostituire le carte delle serie storiche di quanta acqua cada sulla regione Emilia-Romagna in venti anni o solo anche in venti giorni, con un algoritmo che efficienti quel processo, cioè, non ci obblighi a ricostruire le serie storiche in modo manuale, ma lo dia in pasto ad un algoritmo che ci dia come output un risultato accurato.

Quindi, partire con un esperimento piccolo, con un uso piccolo minimale, ma di quell'uso essere certi che l'ente, se chiamato a rispondere, possa dire "rispondo perché ne conosco l'accuratezza del risultato".

Questo processo è un processo attuabile, ma è anche un processo che già ha gli strumenti che in tutto il mondo sono già a disposizione delle aziende. Come si certificano gli ISO? Ci sono gli ISO anche per l'intelligenza artificiale; alcuni modelli di ML hanno già i loro sistemi ISO che dicono che un modello di Machine Learning con determinate caratteristiche può essere certificato; ci sono degli enti certificatori americani che sono perfettamente in grado di consegnarci, perché noi li possiamo implementare all'interno delle nostre aziende, dei meccanismi di attestazione della rischiosità dei modelli che preservino da quel buco nero di responsabilità che nessuno di noi è pronto ad assumersi. Le aziende questo lo fanno già. Assieme a questa implementazione dei modelli di rischio, direi che il cambiamento più epocale che le aziende devono abbracciare, pubbliche o private che siano, è

quello sulle persone e quindi il cambiamento di governance, che non è quella alla quale faceva riferimento la professoressa Finocchiaro prima, ma la governance sono i modelli e le procedure interne che l'azienda si dà quando implementa meccanismi di intelligenza artificiale, partono dalle persone, quindi, il cambiamento è un cambiamento culturale dove, se l'ultimo stagista assunto alla Regione impatta in un algoritmo che crede sollevi un problema di conflitto con un diritto fondamentale, ha la stessa voce del dirigente che ne ha autorizzato l'acquisto e deve poter segnalare questa stortura, quindi, la pervasività del fenomeno consente che le risorse umane che devono - si parla di reskilling, upskilling - adattarsi a questo cambiamento, siano coinvolte nel processo e quindi, solo così, l'intelligenza artificiale andrà davvero a braccetto con le risorse e questo cambiamento coinvolgerà in modo olistico l'ente, l'organizzazione e sarà quindi non un rigetto, come un corpo estraneo, ma sarà invece un innesto che non comporta rigetto perché un innesto consapevole.

Quindi, è una grande sfida oggi quella che si pone, soprattutto direi ad una regione come questa, che è una regione che ha grande maturità nel proporre alla propria collettività sfide che la coinvolgono e potrebbe essere proprio un punto di ingresso per autenticare ed implementare questa intelligenza artificiale consapevole, dove tutti possono giocare un ruolo molto importante. Il danno reputazionale non è soltanto della dirigente, del CEO, del grande manager della società e di tutti che, a cascata, subiscono il riflesso di un fallimento e dunque, questo è un modo, quello che le aziende hanno di implementare con regole e buone pratiche e con codici di autoregolamentazione, che è una grande opportunità in cui ciascuno si può scrivere le regole con le quali intende implementare l'intelligenza artificiale, nel quadro, ovviamente, delle indicazioni date dal regolamento, ma può declinare regole di dettaglio in modo customizzato e individuale per la propria collettività.

Questo ha due grandi vantaggi: uno è quello di coinvolgere l'ente stesso, le risorse dell'ente nello scrivere, nel disegnare, nell'implementare questo processo; due, ha il grande vantaggio della versatilità del modello. Queste regole di dettaglio, questi codici di buone pratiche, questi codici di autoregolamentazione seguono il muoversi dell'ente nell'adattamento all'implementazione dell'intelligenza artificiale, non sono loro che rincorrono la tecnologia, ma sono loro che, scoperto il grande beneficio che la tecnologia comporta, consentono di cambiare una regola che altrimenti, calata dall'alto in modo anonimo, un modo "one size fits all", cioè un modo buono per tutti, è una regola che, di fatto, non potrà necessariamente adattarsi alle diverse realtà.

Quindi, è una grande sfida che oggi il regolatore può dare, se noi la sappiamo cogliere, alle aziende, agli stakeholders, ai cittadini, ai corpi intermedi, agli enti pubblici, per far sì che queste regole si trasformino, da regole intrusive e del tutto inefficienti, a regole, invece che rispondono ai reali bisogni dell'ente. In questo senso, l'autoregolamentazione è già un modello molto sperimentato con grande successo, pensate soltanto ai codici di autodisciplina delle società quotate; le società quotate, che sono le società che occupano la gran parte del mercato, sono quasi interamente normate dai codici di autoregolamentazione. Il codice civile fa una parte, ma non tutto il lavoro. Il codice di autoregolamentazione se lo danno loro, certamente per cluster, per grandi gruppi, ma se lo danno in modo molto più aderente alle proprie esigenze.

Quindi la sfida, secondo me, è una sfida da cogliere, anche perché regioni come questa hanno la maturità per farla, per scriversi codici di autoregolamentazione che facilitino, integrandolo, il lavoro del regolatore in un innesto consapevole, con un gioco di fonti del diritto a cui siamo già abituati, soprattutto direi nel campo del digitale, che consente però che queste regole non siano espunte, non siano sentite come oppressive soltanto, ma siano invece regole scelte deliberatamente e, dunque, sposate con convinzione. Direi che mi fermo qui. Grazie.

Presidente PETITTI: Grazie professoressa Manes per la sua relazione molto eloquente, ricca di spunti e di elementi. Ora aprirei il dibattito oppure le domande che vogliamo porre appunto alla professoressa Finocchiaro e alla professoressa Manes. Prego consigliera Ugolini.

Consigliera UGOLINI: Cerco di non andare fuori tema perché avrei tante domande che non sono però pertinenti all'obiettivo che abbiamo come commissione. Ne faccio una generale, ma noi come Italia avremmo potuto non legiferare per evitare di moltiplicare gli enti anche se non era necessario? Visto che ci sono tante legislazioni, che il problema che abbiamo è quello della semplificazione, non della complicazione.

Professoressa FINOCCHIARO: Allora, come Italia, avremmo in ogni caso dovuto legiferare per scegliere un'autorità nazionale per l'intelligenza artificiale. Il regolamento europeo prevede che ogni Stato membro scelga un'autorità. Poi la scelta che ha fatto l'Italia è complessa perché ne ha scelte due, cioè, ha diviso le competenze fra Agid e ACN.

Il dibattito è stato un dibattito molto nutrito perché, in astratto, tante sarebbero state le scelte possibili e anche le scelte che hanno fatto altri paesi; per esempio, la Spagna ha scelto di istituire - ancor prima che il regolamento europeo entrasse in vigore - un'autorità nazionale per l'intelligenza artificiale, un'autorità nuova per l'intelligenza artificiale; altri paesi hanno affidato le competenze in materia di intelligenza artificiale al Garante privacy, che esisteva già. Anche in Italia, il Garante privacy aveva premuto per avere queste competenze. Altri paesi hanno affidato le competenze a autorità che si occupano di comunicazioni. Quindi, le scelte sono state molto varie. C'è chi ha istituito un Ministero per l'intelligenza artificiale. Comunque, un'autorità in ogni caso andava scelta e poi, ripeto, c'è stato un ampio dibattito e in Italia la scelta è stata quella di una governance che è piuttosto complessa perché, coordinare queste due agenzie con la Presidenza del Consiglio e con tutte le altre autorità, non è semplice. Poi appunto si tratta non di un'autorità amministrativa indipendente, ma di due agenzie. Quindi, anche sotto questo profilo, qualche criticità c'è. Però ecco, per rispondere alla domanda: come minimo, la legge italiana avrebbe dovuto avere un articolo sulla scelta della governance e poi un altro articolo sulla disciplina delle Sandbox. Questi erano proprio obbligatori. Tutto il resto, in realtà, secondo me, non complica molto il quadro normativo perché, una parte della legge italiana è dedicata a ribadire dei principi generali, quindi, possiamo forse dire che era pleonastico, ma non che complichino il quadro normativo; e poi ci sono alcune disposizioni nuove, come quelle sulla privacy in sanità, che sono un primo tentativo di semplificazione.

Ecco, questo è molto importante. Dovremmo cercare dal punto di vista normativo di semplificare e di razionalizzare le molte norme che abbiamo in materia digitale, come illustrava prima la professoressa Manes e come nel rapporto Draghi è detto chiaramente. Cioè, in questo momento ci sarebbe bisogno di un consolidamento e di una razionalizzazione. Questa può avvenire a livello europeo, vediamo il Digital Omnibus che cosa porterà; e può avvenire anche in parte a livello italiano, per esempio, insisto su un tema che a me sta molto a cuore, se solo si ponesse mano alla semplificazione delle regole privacy in materia sanitaria, faremmo un enorme passo avanti, però questa è materia in parte europea e in grandissima parte nazionale.

Presidente PETITTI: Grazie professoressa Finocchiaro. Altri interventi o domande? Prego consigliere Mastacchi. Magari raccogliamo un po' di interventi e poi lasciamo concludere e rispondere alle due professoresses.

Consigliere MASTACCHI: Grazie presidente, come sempre, per questa giornata molto interessante e grazie anche alle due professoresses che ci aprono nuovi orizzonti di riflessione, sia sul fronte dell'utilizzo, cioè, per quello che possiamo fare noi per far sì che i nostri concittadini utilizzano

l'intelligenza artificiale, ma anche per quello che è l'impatto che questo strumento può avere sul nostro lavoro. Quindi, come ho detto già nella precedente audizione con il professor Gardini, se non ricordo male, per noi c'è una doppia responsabilità che oggi è stata messa bene in luce, perché noi la volta scorsa, credo, almeno io mi sono molto focalizzato sull'aspetto di responsabilità personale, anche sul piano etico e sul piano della credibilità della persona, qui invece andiamo più sul pesante quando parliamo di responsabilità e di diritto, quindi con tutti i rischi che questo comporta e, quindi, questo ci invita in qualche misura ad alzare il livello di attenzione per garantire che, quello che noi facciamo, l'utilizzo che noi facciamo e che ne faremo fare rispetto ai provvedimenti che andremo ad emanare, faccia sì che riduca più possibile i rischi tecnici, legali e anche reputazionali. Quindi sul doppio binario innestato sulla volta precedente.

Quindi, è molto importante investire in formazione nostra interna e far sì che questo avvenga anche per chi la utilizzerà all'esterno, per affinare appunto l'implementazione di questa tecnologia. Per esempio, oggi è uscito il tema della certificazione, che è un passaggio molto importante, che io fino ad oggi non avevo mai preso in considerazione come ipotesi, cioè, immaginavo che l'entità IA fosse -- che lo utilizzi e ti assumi le responsabilità o meno, però in effetti, avendo dei codici di scrittura di software, possono essere in qualche misura verificati e certificati.

Quindi, questo alza di molto il nostro livello di attenzione e ancora una volta ci sottolinea come succede regolarmente, bisogna uscire dalle polarizzazioni del "è buona o non è buona". È buona nella misura in cui ne facciamo un uso corretto e nella misura in cui abbiamo degli strumenti per verificare che questo uso corretto lo sia realmente, quindi, si parla di etica, ma in questo caso siamo scesi anche molto nel piano del diritto. Più riusciremo, nella nostra mente, a trasformare questo strumento realmente in uno strumento, quindi, non lo vedremo come un'entità terza che è buona o cattiva, ma riusciremo a capire come poterlo utilizzare, e più avremo la possibilità di usufruire anche dei benefici che questa ci può dare e rafforzando il piano della nostra responsabilità e mettendo sempre al centro la centralità dell'intervento umano nel suo utilizzo, chiaramente.

Quindi, alla fine l'uomo deve comunque prevalere sempre rispetto alla macchina. Faccio sempre l'esempio del coltello che, di per sé, è uno strumento che può essere buonissimo o cattivissimo: buonissimo, se lo usiamo per affettare una pagnotta di pane per sfamare delle persone o cattivissimo, se lo diamo in mano a qualcuno che vuol fare del male a altre persone; e qui credo che siamo in questo campo.

Quindi, il grande paradigma che - mi sento di poter sostenere - dobbiamo affrontare nel prossimo futuro è che l'intelligenza artificiale non deve essere vista, come succede in molti casi, sostitutiva del lavoro di qualcuno o dell'essere umano in toto, ma l'aspetto umano deve comunque rimanere prevalente, quindi, non bisogna fidarsi, ma utilizzarla appunto come strumento, cercando il più possibile di controllarne gli esiti rispetto alle fonti che siano reali. Questo è un po' quello che mi sento di dire. Però è molto interessante questa sovrapposizione di punti di vista, diciamo, una volta politici, una volta etici e nel caso di oggi giuridici, che ci consentono piano piano, man mano che solleviamo lo scolino della pasta, vediamo quale pasta rimane in cottura. Grazie.

Presidente PETITTI: Grazie consigliere. Raccogliamo altri interventi. Consigliere Calvano, prego.

Consigliere CALVANO: Grazie presidente e grazie a entrambe le professoresse per l'approfondimento che ci hanno consentito e ci stanno consentendo di fare in questa sede e questo tema penso che attraverserà le attività delle istituzioni sempre di più nei prossimi anni, a tutti i livelli, dai livelli più piccoli fino arrivare ovviamente all'Unione Europea e alle istituzioni internazionali, che sono già impegnate nelle attività di regolazione di quello che è un nuovo fenomeno, chiamiamolo così, una nuova realtà con la quale ci troviamo ad avere a che fare.

Io sono uno di quelli che vede nelle innovazioni di carattere tecnologico, dalla prima rivoluzione industriale in avanti, opportunità in più; ovviamente, opportunità in più nella misura in cui il soggetto pubblico sia nelle condizioni di mettere in campo un'adeguata regolamentazione. È un po' come per il mercato: c'è chi ritiene che si aggiusti da solo e chi ritiene che gli aggiustamenti di quel mercato debbano essere in un qualche modo guidati anche dall'intervento pubblico. Io appartengo alla seconda categoria in questo caso, senza eccedere, ovviamente, nelle intromissioni rispetto a come le cose possono andare.

L'intelligenza artificiale, a mio avviso, apre un tema di competenze anche all'interno della pubblica amministrazione. Ci sono stati fatti esempi di come probabilmente oggi noi dobbiamo modificare le competenze verso, richiamo l'intervento che ha fatto anche il collega Mastacchi, verso le attività di controllo anziché di produzione degli atti, della costruzione dei processi; nel senso che l'intelligenza artificiale può aiutarci, può sostituirsi a noi, ma sulla sua infallibilità, io credo che i dubbi siano tanti, devono essere tanti, devono continuare ad essere tanti anche col passare del tempo perché sappiamo che l'intelligenza artificiale si migliorerà man mano che la usiamo perché acquisisce sempre più dati, che poi dovrebbero mettere nelle condizioni di offrirci un prodotto sempre migliore. Però penso che all'interno della pubblica amministrazione dovremmo rivedere anche i nostri processi formativi di chi opera nella pubblica amministrazione, forse più nella logica del controllo che nella produzione di atti o di processi. Quindi, volevo chiedere, su questo, cosa ne pensate.

Dall'altro lato, oggi come oggi, nell'utilizzo molto ridotto che ne faccio personalmente, una delle cose che sto capendo è che, se si sbaglia la domanda si ha anche la risposta sbagliata. Quindi, su questo, l'intelligenza artificiale è utile nella misura in cui, chi la usa ha contezza della domanda che sta facendo perché se sbaglia la domanda, si va in una direzione completamente sbagliata. Quindi, anche forse questo è un altro versante, passatemi il termine, di carattere formativo su cui dobbiamo aggiornare le nostre competenze. Grazie.

Presidente PETITTI: Grazie consigliere Calvano. Consigliera Proni, prego.

Consigliera PRONI: Grazie Presidente. Ringrazio anch'io le due docenti. Solo un paio di sottolineature e riflessioni. Ci è stato detto che stiamo sperimentando la tecnologia insieme alla normativa che la disciplina; che se da un lato, a un primo ascolto, può sembrare una cosa che preoccupa, nel senso che quindi non siamo pronti, stiamo facendo una cosa un po' pericolosa; dall'altro, però, quando affermate che tutto quello che scriviamo su una materia così complessa e che evolve così velocemente, rischia di essere già vecchio appena terminato di scriverlo, quindi, restituisce invece una cosa che ha un dato di rischio come è inevitabile che sia, ma anche una sfida molto affascinante rispetto al fatto che per forza di cose la devi sperimentare vivendola e quindi anche tutto il tema di segmentare per processi, per piccoli movimenti, tutto quello che avviene, tutto quello che faccio, tutto quello che scopro, le varie connessioni e nel mentre, normo nuovamente oppure correggo il tiro rispetto a eventuali limiti di azione, come da mettere inevitabilmente in conto.

Quindi, io credo che sia un approccio che se è spiegato, perché c'è un tema comunque di condivisione e di informazione e formazione rispetto anche all'opinione pubblica che deve comunque vederci responsabilizzati, noi stiamo facendo un lavoro che riguarda le competenze della Regione e, parlo per me, dove c'è una forte necessità di prendere più possesso e conoscenza della materia che stiamo trattando, però è anche vero che dobbiamo avvertire una responsabilità verso questa corresponsabilità verso la cittadinanza, verso l'opinione pubblica, perché non ne sia spaventata, ma nemmeno di contro che la -- con una leggerezza non corrispondente al vero.

L'altra riflessione riguardava la consapevolezza e la responsabilità. Se penso alle materie, alle sfide che come amministrazione regionale ci siamo dati, ne dico due, tutto il tema della salute, della sanità e del miglioramento della sicurezza territoriale, li avete citati entrambi, sono le competenze principi della Regione e sono quelle su cui abbiamo delle sfide importanti e assolutamente non rinviabili e quindi, nella gestione di tutto quello che riguarda il cambiamento climatico, che riguarda la protezione civile, c'è una possibilità di esplosione di opportunità, di saperi, di previsioni assolutamente necessaria. Sulla salute altrettanto. Però è vero che, se da un lato la normativa ma dall'altro anche una consapevolezza e un riconoscimento collettivo dell'importanza di questi due frangenti, non possiamo neanche pensare che questa responsabilità ricada solo sulla responsabilità individuale. Tutto quello che può fare in medicina uno strumento di questo tipo, non può essere che si accoli il rischio il professionista, quindi, credo che anche su questo ci sia un mondo di ampliamento di sguardo che può aiutare l'opinione pubblica, perché su questi temi è inevitabile che ci sia sempre lo spauracchio, ma anche il condizionamento che può avere anche un origine ben tracciabile, si muove un mondo economico immenso, quindi, avere un pochino più di strumenti di approccio laico, di aiuto da parte delle istituzioni nei confronti della popolazione, credo che ci aiuti a vederne le opportunità, non scivolare da rischi, ma anche ad aiutarci ad avere un impianto che tenga, che sia il più solido possibile.

L'ultima riflessione l'hanno già toccata i colleghi, Calvano in modo particolare, il tema della formazione. Se un processo, un sistema che contiene un rischio di esclusione, quindi, anche un rischio democratico, se vogliamo, è però vero che, se guardata anche con le parole che avete usato voi nei vostri interventi, di questo coinvolgimento anche dal basso delle persone, non soltanto il dirigente, ma il semplice - tra virgolette - ma non meno fondamentale collaboratore, che diventa parte anche non soltanto di un ne faccio un pezzettino, ma ne sono parte, quindi lo cogestiono, ne sono corresponsabile, ci si può leggere anche una responsabilità dal basso e un coinvolgimento, penso a tutto il tema delle materie scientifiche legate al genere, alla formazione delle ragazze, penso al tema delle giovani generazioni in particolare, e può essere anche un'occasione di ampliamento delle possibilità di essere più coinvolti e responsabilizzati nei processi gestionali che attengono alla propria professione. Quindi, il tema della professione, lo diceva anche il relatore della commissione precedente, della transizione professionale è un tema che apre tante tante possibilità di ragionamento, ma più in termini di possibilità positive che di spauracchi e limiti. Grazie.

Presidente PETITTI: Grazie consigliera Proni. Consigliera Costi.

Consigliera COSTI: Grazie presidente. Grazie alle professoresse, in particolare. Visto che molti temi sono stati trattati dai miei colleghi, volevo fare un paio di domande e sollecitare anche questa commissione rispetto all'osservazione che ha fatto la professoressa Finocchiaro sulla privacy e sulla semplificazione della privacy in materia sanitaria, che sarebbe veramente un tema importante per poter avere una maggiore efficienza ed efficacia perché vediamo che spesso ci scontriamo su questo tema molto pregnante.

Chiedeva due cose. Il tema della regolamentazione rispetto al resto del mondo e come si possono incrociare? L'ultima, il tema degli agenti AI, come sono e quali modifiche ancora possono portare rispetto al tema del lavoro della quale commissione sono presidente. Grazie.

Presidente PETITTI: Grazie. Consigliere mastacchi.

Consigliere MASTACCHI: Una cosa velocissima che ho dimenticato prima, che è più una domanda che una considerazione: l'impatto legato al tema del Codice dei Contratti Pubblici, che nel decreto legislativo 36 del 23, lo prevede chiaramente, fra le righe, l'utilizzo dell'intelligenza artificiale. Quindi

dice: prevede l'introduzione di soluzioni tecnologiche avanzate. Sicuramente può essere uno strumento di velocizzazione grandissima rispetto alle procedure di affidamento che, nella mia esperienza, per esempio, in alcuni casi, sono durate anche oltre un anno, quando ci sono le comparazioni a coppie, situazioni molto complicate che probabilmente nell'utilizzo dell'intelligenza artificiale possono essere velocizzate tantissimo. Quindi, capire anche dal punto di vista normativo, in quale misura questa procedura può essere utilizzata e se è verificabile fino in fondo una sua reale affidabilità. Grazie.

Presidente PETITTI: Grazie. Io non ho altre richieste di interventi o domande. Solo per dire alla consigliera Costi che stiamo, sul tema della semplificazione, anche in accordo con la giunta, preparando un percorso di approfondimento e questo tema, invece, specifico della privacy in relazione all'intelligenza artificiale sarà, anche questo, oggetto di un ulteriore approfondimento. Io ora passo, per alcune risposte, considerazioni finali, la parola alla professoressa Finocchiaro e poi alla professoressa Manes o viceversa. Decidete voi. Io chiederei ad entrambe comunque di rispondere ai tanti temi emersi da questa commissione e dal dibattito con i nostri consiglieri. Quindi a entrambe chiedo di intervenire in merito ai temi che sono stati sottolineati negli interventi. Prego professoressa Manes.

Professoressa avv. MANES: Allora, tento un po' di riallacciarmi alle varie interessantissime sollecitazioni che sono state date. Sicuramente l'impatto che coinvolge le risorse è un impatto che ha bisogno dell'elemento culturale. Questo lo abbiamo detto. Quindi è da prendere come grande opportunità di formazione, per chi già lavora nell'ente e per chi ci lavorerà. È chiaro che questo si attaglia generazionalmente a coloro che questo problema lo avranno nella loro dimensione lavorativa, per 20-30 anni o 40 e quindi problema nuovo da guardare con occhi nuovi e la soluzione a questi problemi ci verrà da gente di vent'anni e non di 50, a mio avviso, neanche di 60, ma questa è un'ottima notizia, non una cattiva notizia. Ma questo impone che tutta la governance, quindi le risorse dell'ente, siano pronte a recepire lo stimolo e la sollecitazione che arriva da qualunque punto dell'organizzazione stessa, quindi, questo lo vogliamo chiamare Change Management, con un'espressione molto in voga? Lo vogliamo chiamare invece adattamento dell'ente, quindi resilienza dell'ente rispetto all'impatto della novità dirompente? Comunque lo vogliamo chiamare questo è chiaramente un processo che deve partire dalla formazione, cioè, il rilievo epistemologico di questo cambiamento è tutto nella testa dell'uomo. Cioè, quindi, il fatto di non andare a discarico di responsabilità verso una macchina che non ce l'ha la responsabilità, perché per quanto parliamo di agente AI, come giustamente nel suo libro conclude la professoressa Finocchiaro, ci deve essere sempre un patrimonio dietro, non è che l'agente AI ha un suo gruzzoletto di cui poi risponde come noi rispondiamo col 2740, che è l'articolo che parla di responsabilità patrimoniale del debitore. Non c'è e neanche ci sarà la famosa polizza dietro che dice che assicura ogni agente AI. Magari ci fosse, ma io ci lavoro da tanti anni in un'assicurazione e tutti noi sappiamo che l'assicurazione, se non conosce ciò di cui rischia, anzi, delle due, vuole sempre rischi un po' minori e non rischi maggiori. Questo per dire, semplificando, banalizzando e con una risata, che chiunque si muova di artificiale attorno e accanto a noi, non so avvertirà mai il paradigma della responsabilità personale, qualunque soluzione si dia alla responsabilità civile che è chi risponde per un danno. Ma questo è un altro punto. La responsabilità personale va avanti, matura con la competenza, con l'aumento della familiarità con questa materia. Quindi, con l'aumento dell'elevazione culturale, l'allitterazione dell'IA di cui tanto si parla.

L'allitterazione dell'IA parte dalla prima base larghissima della piramide e deve arrivare a Top Management. Questo movimento, che deve essere non per frecce su e giù, ma circolare, deve coinvolgere, e tante sperimentazioni si stanno facendo nelle organizzazioni a livello di governance,

di cambi di governance. Questo vuol dire che dapprima sentivamo che coloro che scrivevano le righe di codice dicevano “io ho scritto la riga di codice, questo è quello che so fare, non rispondo se poi questa riga di codice, avendo introiettato i dati, ha dato un risultato totalmente... non è mestiere mio, io ho fatto il mio pezzetto”. Non funziona così, cioè, il Data Scientist, che elabora la soluzione algoritmica con la riga di codice, deve essere in diretto contatto con chi lo controlla - Compliance e Risk Manager - e con il top manager che deve sentire quando si siede in CDA il racconto di chi quell'algoritmo lo ha scritto, anche se non ne sa niente o meglio, molto meglio se ne sa un po'. Quindi un altro passaggio fondamentale è rieducare, tutti noi che abbiamo già un'età matura per ricominciare con questo processo, di nuove competenze accanto a quelle che già abbiamo. Quindi, per esempio, tutto questo movimento che parla di una formazione/educazione specifica sui temi dell'intelligenza artificiale e che non tocca soltanto le risorse che si occupano direttamente di algoritmi, ma tutto l'ente - ripeto - dalla piramide dal top, dalla sommità della piramide fino alla base e viceversa.

Presidente PETITTI: Grazie professoressa Manes. Prego professoressa Finocchiaro.

Professoressa FINOCCHIARO: Grazie per i tanti stimoli, le vostre domande, a cui non so se riusciremo a dare una risposta perché sono veramente tanti. Allora, se dovessi sintetizzare, direi governare l'innovazione fra rischi e opportunità, è quello che si deve cercare di fare in questo momento e che non possiamo non fare, ciascuno dalla sua posizione, dal suo ruolo perché è un'innovazione che naturalmente non è rifiutabile e non è neanche discutibile e che però va governata.

La nostra cultura ci porta a accentuare i rischi. L'intelligenza artificiale da sempre, perché l'accompagna una certa cultura cinematografica, artistica, letteraria, ci ha portato a esaltare i rischi, abbiamo paura fondamentalmente dell'intelligenza artificiale, però l'intelligenza artificiale la vediamo nel nostro quotidiano e quindi secondo me dovremmo cercare di bilanciare appunto i pericoli che neanche conosciamo e le opportunità che sono moltissime.

Io sono per natura innovatrice e ottimista e quindi mi sento più sbilanciata verso le opportunità che verso i rischi. Come si può gestire: l'educazione, l'avete detto tutti, la formazione, che tra l'altro è un obbligo previsto dal regolamento europeo, che non è facile, che bisogna inventarsi, poi se avremo tempo vi racconto quello che faccio io in aula con i miei studenti, e un approccio assolutamente sperimentale.

È stato detto prima dalla professoressa Manes, lo riprendo: non bisogna pensare di fare un progetto globale internazionale per governare l'intelligenza artificiale, tra l'altro questo l'hanno proposto i cinesi, naturalmente perché chiaramente si stanno ponendo come leader globali dell'intelligenza artificiale, si propongono come regolatori globali, hanno l'obiettivo di essere leader entro il 2030 e questo lo avevano deciso nel 2017 e quindi, stanno marciando verso questo obiettivo.

Ma insomma, tornando a noi quello che noi possiamo fare come europei e come italiani, è governare alcuni processi determinati. Decidere, sperimentiamo l'intelligenza artificiale in medicina, supponiamo, in un ambito medico dove i temi di responsabilità sono alti ma le opportunità lo sono altrettanto, soprattutto se pensiamo alla ricerca scientifica. Che cosa serve per sperimentare - uso questo come esempio - in medicina? Alcune regole che bisogna cercare di semplificare. Questo però è un tema nazionale, su questo, il legislatore, la competenza è fondamentalmente nazionale, se pensiamo alla privacy, è una legge nazionale che va semplificata.

Nell'ospedale che cosa serve: una policy che dica chiaramente a tutti gli operatori che, primo, si può utilizzare l'intelligenza artificiale in certi processi; per esempio, non si deve utilizzare in altri processi, quindi, che chiarisca gli ambiti; che chiarisca chi è responsabile e che ci sia eventualmente una assunzione di responsabilità da parte dell'organizzazione che immette determinati programmi o

servizi all'interno dell'ospedale, quindi, è l'organizzazione che decide di utilizzare l'intelligenza artificiale e che quindi si assume una parte di responsabilità; che decide chi deve supervisionare, controllare e come; in quali specifici perimetri si utilizza l'intelligenza artificiale e poi sperimentando, se le regole non bastano o se non sono adeguate, provoca un cambiamento a livello delle regole.

Si parte dal basso. In tutte le sperimentazioni e in tutti gli utilizzi che io conosco di intelligenza artificiale si parte dal basso, non bisogna pensare di partire dai massimi sistemi. Adesso regoliamo l'intelligenza artificiale, no, regoliamo il processo in banca di concessione del mutuo, in un certo range, attraverso l'intelligenza artificiale. In Estonia regoliamo i processi fino a 5.000 euro con l'intelligenza artificiale. Cioè circondiamo l'ambito e poi regoliamo quell'ambito; cioè, bisogna partire dai casi concreti con delle regole che sono le regole che possono darsi in via privatistica sia gli enti pubblici sia gli enti privati, cioè, le policy, i contratti e così via. L'editore che voglia utilizzare l'intelligenza artificiale deve guardare i contratti con gli autori che gli hanno fornito i testi e che lui vuole rielaborare e usare per l'addestramento con l'intelligenza artificiale.

Cioè, il messaggio che sto cercando di dare e che voglio sottolineare è che dobbiamo partire dal basso e da casi concreti e poi in questo modo sperimentare, se volete, con un'altra parola, dai verticali. Non so, per esempio, sul patrimonio artistico di un certo tipo di musei in Emilia-Romagna. Ecco, sulle elaborazioni che si possono fare di patrimonio artistico intellettuale in quei musei, percorsi, sperimentazioni, perimetri, casi d'uso, chiamiamoli come vogliamo, però partire dal basso. Questo credo che sia assolutamente fondamentale perché il rischio, altrimenti, è quello di non fare, aspettando che ci siano delle regole che rispondono a tutte le domande e queste regole, secondo me, non le avremo mai o se volete una risposta diversa, le risposte alle molte domande le avremo attraverso i principi.

Cioè, è vero che non abbiamo delle regole ad hoc sulla responsabilità dell'intelligenza artificiale, però se ci fosse un processo domani, utilizzeremo le regole del codice civile e della responsabilità del produttore per individuare chi deve rispondere, chi deve rimborsare un certo danno che è stato causato dall'intelligenza artificiale e, quindi, sostanzialmente in questa cornice molto complessa, secondo me, ci dobbiamo muovere partendo dai casi d'uso e sperimentando e dando agli operatori quegli strumenti che chiariscano e assicurano perché, qual è il problema oggi dell'operatore - torniamo al medico - che vuole utilizzare intelligenza artificiale: sapere che cosa succede, chi risponde se qualcosa non funziona. Non deve trovarsi nella situazione in cui qualcuno gli dice "non potevi utilizzare l'intelligenza artificiale, perché l'hai fatto? Chi ti ha autorizzato?" E allora dobbiamo creare il percorso perché questo possa succedere.

In estrema sintesi, direi che questo è il contributo che posso dare in risposta alle vostre domande. Ecco, è un'altra cosa che mi sta molto a cuore: la semplificazione non vuol dire deregolamentazione, perché nell'era Trump semplificazione è diventata deregulation. No, non è questo quello che intendevo dire. Semplificazione significa rendere le regole applicabili perché oggi abbiamo il problema opposto, ne abbiamo troppe e quindi sono inapplicabili.

Presidente PETITTI: Grazie davvero alle professoresse Finocchiaro e Manes. Credo sia stata molto importante l'audizione di oggi, un tassello prezioso del nostro percorso di approfondimento. Ricordo anche che ci sarà una raccolta dei lavori di queste commissioni, in modo che potremo tornare sopra i vari temi che sono stati oggetto di queste riflessioni. Grazie a tutti voi e ci aggiorniamo presto per le prossime audizioni. Buon pomeriggio a tutti.