

L'IMPATTO SOCIALE E NORMATIVO DELL'INTELLIGENZA ARTIFICIALE

1^a SETTIMANA INTERNAZIONALE

(13-16 MAGGIO 2025)

UNIVERSITÀ DI OVIEDO

di Paolo Inturri

PREMESSA

Quello che segue è il resoconto della 1^a settimana internazionale su “*L'impatto sociale e normativo dell'intelligenza artificiale*” — “*El impacto social y normativo de la inteligencia artificial*” — tenutosi dal 13 al 16 maggio 2025 presso l'Università di Oviedo.

Nel tentativo di rappresentare fedelmente lo svolgimento dei lavori l'esposizione è stata suddivisa per giornate.

Di ciascuna giornata verrà illustrato il contenuto delle diverse tavole rotonde riportando sommariamente i punti salienti di ciascuna relazione.

INDICE

MARTEDÌ 13 MAGGIO 2025

<u>INAUGURAZIONE ISTITUZIONALE DEI LAVORI</u>	4
<u>IL CONTESTO DELL'IA</u>	5
SUSANA MONTES RODRÍGUEZ	5
FERNANDO HIGINO LLANO ALONSO	6
IBAN GARCÍA DEL BLANCO	7

MERCOLEDÌ 14 MAGGIO 2025

<u>LA REGOLAZIONE DELL'IA</u>	9
MIGUEL ÁNGEL PRESNO LINERA	9
JOSÉ LUIS BOLZAN DE MORAIS	11
JAVIER FERNÁNDEZ RODRÍGUEZ	12
<u>ISTRUZIONE E IA</u>	14
JAVIER FOMBONA	14
JULIO ALBALAD GIMENO	14
THOMAS CASADEI	16
<u>CYBERSICUREZZA</u>	17
MARÍA JOSEFA RIDAURA MARTÍNEZ	17
LUIS MOLINA VALBUENA	18
STEFANO PIETROPAOLI	19

GIOVEDÌ 15 MAGGIO 2025

<u>IA E LAVORO</u>	
IVAN ANTONIO RODRÍGUEZ CARDO	19
LUCIA MONTEJO	20

JAVIER CAMPA	21
ISABEL SANTOS	21
<u>IA E SANITÀ</u>	23
SERGIO CALLEJA	23
ELENA DÍAZ RODRÍGUEZ	24
PEDRO JOSÉ ESPINOSA PRADOS	25
<u>ECONOMIA E SOSTENIBILITÀ DELL'IA</u>	
LUIS MOLINA VALBUENA	27
CARLOS IBÁÑEZ SÁNCHEZ	28
 <u>VENERDÌ 16 MAGGIO 2025</u>	
<u>IA E GIUSTIZIA</u>	29
JORDI NIEVA FENOLL	29
ALFREDO RODRÍGUEZ DEL BLANCO	29
AITOR CUBO	30
<u>CONFERENZA DI CHIUSURA</u>	32
LORENZO COTINO HUESO	32

MARTEDÌ 13 MAGGIO 2025

INAUGURAZIONE ISTITUZIONALE DEI LAVORI

Comitato Organizzativo: Agustina **Bouchet** Gutierrez, Roger **Campione**, Miguel Ángel **Presno** Linera

Il comitato organizzativo ha accolto i partecipanti nella sede del **Centro di Studi sull'Impatto Sociale dell'Intelligenza Artificiale** (*Centro de Estudios sobre el Impacto Social de la Inteligencia Artificial - CEISIA*) presso la città di Lugones.

Nei loro interventi Agustina Bouchet Gutierrez, Roger Campione e Miguel Ángel Presno Linera, tutti e tre professori dell'**Università di Oviedo** (UNIOVI) e membri del CEISIA, hanno illustrato l'idea di fondo da cui ha preso le mosse l'iniziativa.

In tale ottica, il comitato organizzativo ha posto l'accento sulla necessaria **interdisciplinarietà** degli studi in materia di intelligenza artificiale (IA) e sull'importanza del dialogo tra accademia, istituzioni e società civile, testimoniato dalla qualifica dei relatori presenti nelle diverse giornate.

Nella visione degli organizzatori l'approccio all'IA non può che essere **trasversale** e, quindi, deve necessariamente giovare della mutua collaborazione tra umanisti e tecnici.

I primi si interrogano sull'impatto sociale e normativo dell'IA, in assenza, tuttavia, delle necessarie nozioni tecniche; i secondi pur essendo esperti negli aspetti tecnici, non possono prescindere dall'apporto degli umanisti per comprendere e orientare l'impatto dell'IA nella società.

IL CONTESTO DELL'IA

Relatori: Susana **Montes** Rodríguez, Fernando Higino **Llano** Alonso, Iban **García** del Blanco

Moderatore: Roger **Campione**

Contenuto: Nella prima tavola rotonda i relatori si sono confrontati sul contesto dell'IA sulla base di tre distinte prospettive di partenza: (1) tecnica (Montes); (2) etica (Llano; (3) politica (**García**).

SUSANA MONTES RODRÍGUEZ, professoressa di matematica e statistica presso UNIOVI e membro del CEISIA, ha aperto la discussione ponendo innanzi la propria **prospettiva tecnica** sull'IA.

In primo luogo, Montes ha paragonato lo sviluppo dell'IA alla ricerca medica, evidenziando come in entrambi i settori dapprima la tecnica crea gli strumenti, ma poi spetta all'etica e al diritto stabilire cosa sia lecito o etico fare.

Più in generale, la relatrice ha sottolineato come "IA" sia una espressione molto diffusa, ma troppo spesso poco compresa.

I metodi predittivi alla base dell'IA esistono da tempo, tuttavia, attualmente, si assiste ad un vero e proprio "boom" in ragione della crescente potenza di calcolo dei computer e alla sempre maggiore disponibilità di dati.

Sebbene lo sviluppo dell'IA muova dall'idea che **le macchine** possano pensare, lanciata da Alan Turing nel 1950, esse **non pensano realmente**.

Strumenti come ChatGPT, pur essendo ottimi per generare testo, basandosi su enormi volumi di dati processati velocemente, non eseguono operazioni logiche o non "ragionano" nel senso umano. Riproducono in base alla probabilità quanto hanno visto nel loro addestramento con maggiore frequenza. Proprio per questa ragione **l'atteggiamento critico e la supervisione umana** sono indispensabili. Senza supervisione l'IA può rivelarsi pericolosa a causa della sua tendenza all'errore, ma in ogni caso la responsabilità per l'uso dell'IA ricade sempre sull'operatore umano.

Sulla base degli attuali sviluppi dell'IA, che riporta risultati sorprendenti in *task* come la generazione di testo e il riconoscimento di *pattern*, la relatrice ha opinato che l'IA non eliminerà il lavoro umano, ma lo trasformerà, rendendolo più intellettuale.

Al netto di ciò, Montes ha evidenziato le principali difficoltà tecniche che, ad oggi, rappresentano ancora una sfida per lo sviluppo dell'IA:

(1) Bias: l'IA riflette i pregiudizi presenti nei dati di addestramento, potendo portare a risultati discriminatori per determinati gruppi, i quali, però, derivano non da un intento discriminatorio, ma dalla composizione dei dati stessi;

(2) Dipendenza e Vulnerabilità: tanto più deleghiamo compiti della nostra vita all'IA, tanto più ne siamo dipendenti e rischiamo di ritrovarci vulnerabili nel caso di attacchi informatici o di malfunzionamenti della rete;

(3) Privacy e Sicurezza dei Dati: l'IA apprende dai dati inseriti dagli utenti, richiedendo cautela nel fornire informazioni sensibili o non private;

(4) Disuguaglianze: l'IA presenta differenze, in termini di opportunità di accesso a questa tecnologia, tra chi dispone dei mezzi e della conoscenza per sfruttarla e chi, invece, ne è privo;

Alla luce di ciò, Montes ha concluso l'intervento sottolineando l'esistenza di una responsabilità collettiva di imparare a usare l'IA correttamente.

FERNANDO HIGINO LLANO ALONSO, professore di filosofia del diritto presso l'Università di Siviglia, ha affrontato la discussione sulla base di una **prospettiva etica e normativa**.

In particolare, ha sottolineato come l'etica dell'IA non si traduca in una applicazione particolare dell'etica tradizionale, quanto più in un insieme di principi che le **aziende** devono rispettare per essere **affidabili e credibili nel mercato europeo**.

Più nel dettaglio, per regolare il proprio mercato l'Unione europea ha tenuto in considerazione i **principi etici fondamentali** dell'IA: Giustizia, Autonomia, Benevolenza, Non-maleficenza e Spiegabilità.

Proprio per il rispetto di tali principi l'AI Act può essere considerato un modello di regolamentazione antropocentrico, a differenza del modello statalista cinese e di quello ultraliberale statunitense.

Secondo Llano per comprendere appieno l'etica dell'IA è di fondamentale importanza considerare che **gli algoritmi in sé non hanno etica**, ma riflettono i bias anche inconsci di chi li progetta e li sviluppa: esseri umani che, invece, sono portatori di una specifica etica.

L'intervento è stato concluso rimarcando l'importanza dell'**allineamento dell'IA con i valori umani** e della necessità di una vasta alfabetizzazione digitale, come processo che coinvolga con responsabilità condivisa amministrazione pubblica, scienza e società civile.

IBAN GARCÍA DEL BLANCO, ex eurodeputato e negoziatore dell'AI Act, ha apportato alla discussione una **prospettiva politica e regolatoria** dell'IA con particolare riferimento al contesto europeo.

Punto di partenza dell'intervento è stata la constatazione che nonostante l'approvazione dell'AI Act, i rapidissimi sviluppi del settore dell'IA prospettano diverse questioni fondamentali che dovranno presto o tardi essere affrontate dalla nostra società.

Al netto di ciò, il relatore ha manifestato ottimismo nei confronti dell'AI Act, il quale seppur non in vigore nella sua interezza è già stato in grado di influenzare gli operatori del settore.

Questi ultimi, in particolare, sembrano vedere con sfavore tale atto normativo e attraverso un'intensa attività di lobbying mirano ad alleggerire le obbligazioni imposte dall'AI Act. In altri termini, la “competizione tecnologica” non può giustificare iniziative di *deregulation*, poiché solo il diritto può garantire la centralità dell'uomo.

Infatti, nonostante l'AI Act consideri l'IA come un prodotto, il suo **fine ultimo è costituito dalla protezione dell'essere umano** e dei diritti fondamentali.

Tra gli aspetti dell'AI Act degni di nota il relatore ha evidenziato:

- (1) La **vocazione geografica globale (extraterritoriale)**, si applica a tutti gli operatori che operano nel mercato europeo, indipendentemente dalla loro origine;
- (2) L'esclusione dall'ambito di applicazione dei settori **militare** e di **sicurezza nazionale** (in quanto di competenza esclusiva degli Stati membri), per i quali sarebbe auspicabile un quadro etico comune;
- (3) L'introduzione di **nuovi diritti per i cittadini**, come il diritto ad essere informati e a ricevere una spiegazione ragionevole sull'uso dell'IA che li riguarda;
- (4) L'obbligo di **consultare i lavoratori** prima dell'implementazione dell'IA nel luogo di lavoro.

Infine, il relatore ha evidenziato il necessario collegamento tra IA e **democrazia**. Infatti, dato il potenziale trasformativo dell'IA, i cittadini devono essere coinvolti nelle decisioni che riguardano la sua implementazione e i suoi limiti. Allo stesso tempo, l'adozione di **decisioni democratiche e collettive** consapevoli implica un livello diffuso di **istruzione digitale** tra la cittadinanza attiva, pena il crearsi di un divario di disuguaglianza digitale.

MERCOLEDÌ 14 MAGGIO 2025

LA REGOLAZIONE DELL'IA

Relatori: Miguel Ángel **Presno** Linera, José Luis **Bolzan** de Moraes, Javier **Fernández** Rodríguez

Moderatore: Miguel Ángel Presno Linera

Contenuto: La tavola rotonda su “*La regulación de la IA*” ha dato avvio alla seconda giornata della Settimana internazionale dell'IA. I relatori, riuniti presso l'aula “Severo Ochoa” dell'edificio storico dell'Università di Oviedo, hanno affrontato il tema della regolazione dell'IA da tre diverse prospettive: (1) quella europea (Presno); (2) quella brasiliana (Bolzan); quella della pubblica amministrazione asturiana (Fernández).

MIGUEL ÁNGEL PRESNO LINERA, professore di diritto costituzionale presso UNIOVI e membro del CEISIA, ha affrontato il tema della regolazione dell'IA nell'UE, focalizzando la propria esposizione su tre tematiche principali: (1) il procedimento di approvazione dell'AI Act; (2) la definizione giuridica di IA; (3) i punti critici dell'AI Act.

(1) IL PROCEDIMENTO DI APPROVAZIONE DELL'AI ACT

Prima di ricostruire il procedimento di approvazione dell'AI Act, Presno ha preso le mosse dai risultati degli studi della Professoressa Anu Bradford.

In particolare, quest'ultima nell'opera “*Digital Empires: The Global Battle to Regulate Technology*” sistematizza i modelli di regolamentazione dell'IA di USA, Cina e UE, sottolineando che nessuno di essi nella pratica si realizza nella sua versione pura.

Il modello statunitense si basa tendenzialmente sul libero mercato e, quindi, su un intervento normativo e statale minimo.

Il modello cinese prevede in linea di massima una regolamentazione statale totale, con l'obiettivo di utilizzare l'IA come strumento di controllo sociale e politico.

Il modello europeo si concentra sulla garanzia dei diritti, promuovendo lo sviluppo tecnologico nel rispetto dei diritti fondamentali.

Coerentemente con quest'ultimo modello, l'UE ha emanato l'AI Act con l'ulteriore fine di innescare quello che la professoressa Bradford definisce come “*Brussels Effect*” ovvero l'effetto per il quale Stati extra-comunitari emanano normative simili a quelle comunitarie onde poter fornire alle proprie aziende un assetto regolamentare omogeneo a livello globale.

Premesso ciò, Presno ha ricostruito il procedimento normativo che ha condotto all'approvazione dell'AI Act.

Così, ha ricordato: che il Libro Bianco sull'IA (2020) rappresenta il punto di partenza del procedimento di regolazione di questa tecnologia nell'UE; che ad esso è seguita la proposta di Regolamento della Commissione (2021); che il Parlamento Europeo ha presentato degli emendamenti (giugno 2023) volti a rafforzare la protezione dei diritti fondamentali; che alcuni degli stessi sono stati rigettati nei triloghi (dicembre 2023); che con l'approvazione dell'AI Act (2024) si è dato luogo al complesso lavoro di traduzione del testo nelle lingue ufficiali dell'Unione, di delicata importanza in ragione della diretta e obbligatoria applicabilità dello stesso in tutti gli Stati membri.

(2) LA DEFINIZIONE GIURIDICA DI IA

Successivamente, Presno ha affrontato il tema del concetto giuridico di IA, punto molto dibattuto in sede di approvazione dell'AI Act, in ragione del fatto che tra i diversi documenti ufficiali pubblicati a livello mondiale figurano almeno 55 definizioni distinte.

Il legislatore europeo ha adottato una definizione che non si discosta troppo da quelle utilizzate in altri contesti internazionali (si vedano ad es. quelle adottate dall'OCSE e dal Consiglio d'Europa).

In sostanza, il Regolamento definisce l'IA come un sistema basato su macchine capace di inferire e, a partire da questa inferenza, fare previsioni su eventi futuri, elaborare contenuti, dare raccomandazioni o persino prendere decisioni. **L'inferenza, l'intellettualità e l'autonomia** sono gli elementi chiave. Pertanto, i sistemi che non possiedono queste caratteristiche, anche se algoritmici, non saranno soggetti al Regolamento (es. VioGén, utilizzato in Spagna per prevedere il rischio di recidiva nei casi di violenza di genere).

Tale concezione di fondo dell'IA si riversa in una delle principali particolarità del Regolamento ovvero sia che lo stesso considera l'IA sia come **oggetto di regolamentazione** (ad es. con pratiche proibite o permesse) sia come **soggetto di regolamentazione**, potendo costituire o influenzare situazioni giuridiche soggettive (ad esempio, un sistema che decide l'assunzione di un lavoratore).

(2) I PUNTI CRITICI DELL'AI ACT

Infine, Presno ha evidenziato quelli che a suo parere sono i principali nodi problematici dell'AI Act.

Innanzitutto, il Regolamento è stato concepito per essere dinamico, ovvero sia con la previsione di **meccanismi di aggiornamento periodico**.

Tuttavia, da una parte, questo conferisce alla **Commissione** un **potere** significativo nell'approvazione di linee guida (es. quella sulle pratiche proibite) e atti delegati, dall'altra, non è

detto che l’emanazione di atti per interpretare un Regolamento già di per sé molto esteso possa effettivamente chiarificarne la **portata precettiva**.

Allo stesso tempo, l’AI Act **entra in vigore a scaglioni**: alcune parti sono già in vigore (es. le pratiche proibite); altre non lo saranno prima del 2030.

Tutto ciò, unito alla rapidità dello sviluppo del settore, pone il serio dubbio su quale Regolamento sarà effettivamente in vigore nei prossimi anni.

Ulteriori dubbi sono stati posti con riferimento alla previsione della necessaria **supervisione umana** per i sistemi ad alto rischio, in ragione dell’improbabile capacità umana di supervisionare sistemi complessi senza esserne influenzati.

Infine, il **modello di governance** dell’IA basato sulla **garanzia dei diritti** si scontra con le esigenze del **mercato** unico e della competitività.

Il Regolamento tratta l’IA come un **prodotto industriale**, concentrandosi solo sui rischi significativi per i diritti fondamentali. Sempre in tale ottica appare dubbioso l’ampio margine di **autovalutazione** del rischio dei prodotti concesso in capo ai fornitori dei sistemi.

Così come appare dubbiosa l’esclusione dall’ambito di applicazione dei sistemi militari e di ricerca scientifica. Infatti, sebbene motivate rispettivamente da ragioni di competenza comunitaria o di promozione della conoscenza, tali esclusioni potrebbero trascurare l’impatto sui diritti in questi settori.

Soprattutto, conclude Presno, appare paradossale che le esigenze di mercato consentano di esportare (anche a Stati che non riconoscono il rispetto dei diritti umani) quegli stessi sistemi di IA che sono proibiti nell’UE.

JOSÉ LUIS BOLZAN DE MORAIS, professore di diritto costituzionale e digitale presso l’Università di Vitória (Brasile), ha illustrato il contesto normativo nel quale si colloca il fenomeno dell’**IA in Brasile**.

Punto di partenza della riflessione di Bolzan è stato la constatazione del **limite degli Stati** nazionali territoriali nella regolazione di fenomeni tecnologici come Internet e l’IA, che hanno un respiro non territoriale. Ciò non vuol dire per Bolzan abdicare alla regolazione: essa è necessaria per bilanciare le esigenze di mercato con la tutela dei diritti, il che non è concepibile con l’autoregolazione.

Premesso ciò, il relatore ha evidenziato come l’ordinamento brasiliano si caratterizzi per la presenza di diversi atti e disposizioni tese a regolare il fenomeno *de qua*.

Innanzitutto, è stato menzionato l’**habeas data** di cui all’art. 5, co. LXXII, della Costituzione brasiliana del 1988, che, sebbene pensato prima dell’era di Internet, costituisce un precursore della normativa di settore, garantendo il diritto di accesso alle informazioni personali in database pubblici.

La vera svolta per l'ordinamento brasiliano avrebbe dovuto essere rappresentata dal *Marco Civil da Internet* del 2014, considerata al momento dell'entrata in vigore come la "Costituzione di Internet". Tuttavia, dopo 10 anni, proprio l'esperienza di questa normativa dimostrerebbe, secondo Bolzan, quanto affermato in premessa ovverosia l'incapacità degli Stati territoriali di regolare fenomeni di questa portata.

In particolare, sotto il profilo della **responsabilità delle piattaforme digitali** il modello del *Marco Civil* impone l'obbligo in capo al *provider* di rimuovere un determinato contenuto solo in seguito ad un ricorso dell'utente al quale sia seguito un ordine giudiziale di rimozione. Tale meccanismo ha mostrato tutti i suoi limiti con i problemi di disinformazione legati alle elezioni presidenziali del 2018.

La **giustizia elettorale** in Brasile ha mostrato maggiore agilità: è titolare di poteri normativi e amministrativi che le hanno permesso di regolamentare le elezioni del 2022, promuovendo la collaborazione con le piattaforme per gestire la disinformazione, senza necessità di attendere un ordine giudiziario. È stato anche implementato uno strumento ("Pardal") per la segnalazione amministrativa di casi di disinformazione.

Al contempo esistono alcuni progetti legislativi per modificare il regime della responsabilità delle piattaforme, ma sono stati bloccati da campagne di *lobbying* da parte delle stesse piattaforme.

Per quanto riguarda la **protezione dei dati** il Brasile si è dotato nel 2020 della **Lei Geral de Proteção de Dados Pessoais** (LGPD) che ricalca la struttura del GDPR europeo.

Permane ancora in stato di discussione un **progetto di legge quadro sull'Intelligenza Artificiale (PL 2338/2023)** modellato sulla falsariga dell'AI Act europeo.

Infine, il professor Bolzan ha sottolineato l'ampio processo di digitalizzazione e informatizzazione che ha interessato il sistema giudiziario brasiliano, il quale, in questa fase, si sta caratterizzando per un uso sempre crescente dell'**IA nei tribunali**, tanto regolato dalle risoluzioni del Consiglio Nazionale di Giustizia (n. 332 e 615 del 2015), quanto informale.

JAVIER FERNÁNDEZ RODRÍGUEZ, direttore della *Dirección General de Estrategia Digital e Inteligencia Artificial* del Principato delle Asturie, ha concluso la serie di interventi della tavola rotonda affrontando il tema della **regolazione dell'uso dell'IA nell'amministrazione del Principato delle Asturie**.

Secondo Fernández, l'IA rappresenta una sfida per la p.a. in ragione dell'aspettativa della cittadinanza di utilizzarla per migliorare l'efficienza dei servizi e della generale lentezza degli apparati burocratici di adattarsi a cambiamenti tecnologici così rapidi e impattanti.

Affinché l'amministrazione pubblica possa sfruttare tutte le potenzialità dell'IA è necessario adottare una **strategia** strutturata, che includa nella propria visione obiettivi e priorità della p.a., nonché le esperienze del mercato e delle imprese private.

Solo in tal modo sarebbe possibile sfruttare e adattare l'IA alle **concrete esigenze dell'amministrazione**.

In tale ottica, l'introduzione dell'IA nella p.a. richiede **l'adattamento dell'organizzazione e del personale**: è necessario un piano di risorse umane che consideri l'impatto sui profili professionali (alcune mansioni potrebbero diventare superflue, richiedendo riqualificazione o adattamento). Servono dirigenti dedicati, un gruppo di lavoro specializzato, il supporto di esperti esterni e un adeguato sistema di governance.

Per affrontare queste sfide, il Principato delle Asturie sta adottando un **decreto interno per regolare l'uso dei sistemi di IA**, strumento per garantire il rispetto delle *good practice* e la riduzione dei rischi (come ad es. la non conformità con la normativa di rango superiore), così da fornire sicurezza giuridica ai progetti della p.a. sull'IA.

Più nel dettaglio il decreto, elaborato con l'apporto di un gruppo di lavoro multidisciplinare, mira a regolare l'uso dei sistemi di IA, fissando criteri per l'acquisizione, lo sviluppo, l'implementazione, il monitoraggio e la valutazione continua. Inoltre, ha lo scopo di definire il meccanismo di governance dell'amministrazione autonoma, nonché di creare *sandbox* normative per permettere alle aziende di testare soluzioni IA in un ambiente controllato.

In conclusione, Fernández ha illustrato i diversi **casi d'uso** e progetti pilota o in corso nell'amministrazione del Principato delle Asturie, come Hoga, un coordinatore virtuale per i diritti sociali che interagendo con gli utenti permette di fornire un livello di attenzione al cittadino personalizzato non altrimenti ottenibile con risorse umane.

ISTRUZIONE E IA

Relatori: Javier **Fombona**, Julio **Albalad Gimeno**, Thomas **Casadei**

Moderatore: Marián González Rúa

Contenuto: La seconda tavola rotonda del mercoledì ha visto confrontarsi i relatori sulle potenzialità e sull'impatto dell'IA nel settore dell'istruzione.

JAVIER FOMBONA, professore di scienze dell'educazione presso UNIOVI e membro del CEISIA, ha aperto il proprio intervento lanciando l'idea che l'IA non sia un semplice strumento, ma un compagno capace di interagire con gli esseri umani.

Ha evidenziato come l'IA sia in costante cambiamento e in rapido progresso in moltissimi settori, rispetto ai quali, tuttavia, quello dell'istruzione rimane marginale.

Nello specifico, in tale contesto l'IA permette una **gestione smart del web**, permettendo agli studenti di ottenere direttamente documenti e riassunti da diverse fonti senza accedere ai siti tradizionali.

Le attuali capacità dell'IA di generare immagini ed elaborare il linguaggio naturale, interagendo così con l'utente, permettono agli studenti di utilizzarla per ottenere traduzioni, redazione di testi accademici e creazione di materiali educativi.

Dall'**integrazione dell'IA con la realtà aumentata e virtuale** stanno sorgendo piattaforme educative in grado di generare unità didattiche nelle quali è possibile apprendere in maniera innovativa, ad esempio, conversando con personaggi storici. Inoltre, stanno emergendo **piattaforme con apprendimento adattivo**, che partendo da un'analisi individualizzata delle capacità dello studente personalizzano la strategia di apprendimento a seconda delle sue esigenze. Ciò permette di implementare **metodologie differenziate** che tengano conto del ritmo e dello stile di apprendimento individuale, anche in classi con molti studenti.

Una eventuale evoluzione dell'istruzione in questa direzione rischia, al contempo, di ridurre l'autonomia dei docenti, che, ove non intendano delegare in toto determinate decisioni di apprendimento all'algoritmo devono essere soggetti ad un aggiornamento costante.

JULIO ALBALAD GIMENO, Direttore del “*Instituto Nacional de Tecnologías Educativas y de Formación del Profesorado*” (INTEF) presso il *Ministerio de Educación, Formación Profesional y Deportes* spagnolo, ha focalizzato il suo intervento su tre aspetti principali: (1) l'utilizzo di alunni e

docenti dell'IA; (2) i rischi dell'IA nell'istruzione; (3) le iniziative del Ministero dell'Istruzione sull'IA.

(1) L'UTILIZZO DI ALUNNI E DOCENTI DELL'IA

Con riferimento all'utilizzo dell'IA da parte degli **alunni**, il relatore ha spiegato che la nuova legge educativa (LOMLOE 2022) ha riformato i *curricula* degli studenti basandoli sulle competenze, tra le quali figura anche quella digitale, da acquisire entro la fine della scuola secondaria obbligatoria. La **competenza digitale** implica un uso salutare e sostenibile delle tecnologie, inclusa l'IA, che però, è menzionata in questo contesto solo con riferimento a materie specifiche negli indirizzi superiori. L'obiettivo del Ministero è realizzare una **istruzione sulla IA** (sul funzionamento) e **per la IA** (per il suo utilizzo critico).

Per quanto attiene ai **docenti**, il Quadro di Riferimento della Competenza Digitale Docente (2022) menziona specificamente l'IA tanto per lo sviluppo delle competenze digitali degli alunni quanto per la promozione del suo uso critico da parte dei docenti.

(2) I RISCHI DELL'IA NELL'ISTRUZIONE

Il Ministero dell'istruzione ha pubblicato una **guida sull'uso dell'IA nell'educazione non universitaria** (2024) che analizza rischi e sfide per studenti, docenti e centri educativi, proponendo delle strategie.

I rischi maggiori per gli studenti sono la **mancanza di competenze digitali** e di **possibilità di accesso tecnologico**, specie per quelli provenienti da famiglie svantaggiate.

La preoccupazione maggiore, però, è rappresentata dai **dati**: l'uso dell'IA da parte degli studenti, se incentivato dai docenti, deve essere accompagnato dalla consapevolezza dei rischi legati alla fornitura di dati personali, ai bias e alle allucinazioni algoritmiche.

Per i docenti il rischio principale è la **mancanza di formazione**, data l'età media elevata.

In generale, il docente dovrebbe essere incoraggiato a vedere l'IA come uno **strumento di supporto, non sostitutivo**.

In particolare, il relatore sottolinea che l'uso dell'IA nell'istruzione è considerato ad **alto rischio** dall'AI Act, sicché, da una parte, ogni valutazione deve rimanere sotto il controllo umano del docente, dall'altra, pratiche come il riconoscimento facciale delle emozioni degli studenti, utilizzato in Cina, sono **vietate**.

(3) LE INIZIATIVE DEL MINISTERO DELL'ISTRUZIONE SULL'IA

Tra i progetti pilota di IA del Ministero dell'Istruzione il relatore ha menzionato: lo sviluppo di **chatbot per le famiglie** (per supportarli nei processi amministrativi maggiormente complessi); la

creazione di strumenti per la **creazione di risorse didattiche per i docenti**; basati su contenuti validati per garantirne la qualità; l'adattamento dei materiali per studenti con bisogni specifici.

THOMAS CASADEI, professore di filosofia del diritto presso l'Università di Modena e Reggio Emilia, ha condiviso delle riflessioni a partire dai risultati di un progetto da lui guidato focalizzato sull'**uso consapevole della rete, dei social media e delle nuove tecnologie**, come l'IA, da parte delle giovani generazioni.

Il progetto diretto da Casadei prende le mosse da una rilevazione statistica di Eurostat secondo la quale in Europa il **10% dei giovani può sviluppare comportamenti problematici o a rischio quando naviga in rete**, una percentuale contenuta ma significativa che richiede prevenzione.

Il progetto si caratterizza per un taglio interdisciplinare, mettendo insieme diritto, informatica e scienze sociali e si concentra sul rapporto tra cybersecurity, nuove tecnologie e protezione dei diritti.

In particolare, le quattro direttrici di ricerca sono: (1) piattaforme digitali, con focus su **TikTok**, la più usata dai minori; (2) **discriminazioni digitali** e *online hate speech*, con lo scopo di individuare i gruppi più vulnerabili; (3) educazione dei minori ai rischi legati ai **reati informatici**; (4) le buone pratiche e i progetti di scuole, città e università sulla **protezione dei dati**.

I principali risultati attesi sono: (1) la redazione di una **guida per l'uso consapevole delle nuove tecnologie** rivolta ad ogni tipologia di educatore, nella quale si affronteranno temi come *fake news*, *revenge porn*, cyberbullismo, *hikikomori* e *dark web*; (2) la creazione di un osservatorio; (3) l'apertura di uno sportello di aiuto.

Alla luce dei risultati dei propri studi, Casadei ha concluso sottolineando la necessità di un approccio sistemico all'impatto delle tecnologie, considerando tanto l'aspetto giuridico, quanto quelli scolastico, psicologico, relazionale, sanitario e politico.

Inoltre, si è mostrato dubbioso rispetto alla poca attenzione dedicata dall'AI Act al settore dell'istruzione, che potrebbe, invece, giovare di un distinto documento **sulla cultura e sulla consapevolezza digitale a livello europeo**.

CYBERSICUREZZA

Relatori: María Josefa **Ridaura** Martínez, Luis **Molina** Valbuena, Stefano **Pietropaoli**

Moderatore: Josè Manuel Redondo

Contenuto: La tavola rotonda pomeridiana in tema di cybersicurezza, tenutasi presso la sede del CEISIA di Lugones, ha affrontato le questioni problematiche legate: ai sistemi di riconoscimento facciale (Ridaura), all'identità digitale (Molina) e al rapporto tra diritto e nuove tecnologie (Pietropaoli).

MARÍA JOSEFA RIDAURA MARTÍNEZ, professoressa di diritto costituzionale presso l'Università di Valencia, ha aperto la sessione affrontando il tema dei **sistemi di riconoscimento facciale**.

In particolare, Ridaura ha evidenziato il diverso impatto sui diritti fondamentali tra sistemi di verifica e sistemi di identificazione.

Infatti, i sistemi di **verifica** (uno a uno), come quelli usati per lo sblocco del telefono, presentano meno problemi di diritti fondamentali rispetto ai sistemi di **identificazione** (uno a molti), in quanto mentre i primi confrontano l'immagine inquadrata con un'altra già precedentemente inserita, i secondi la confrontano con un intero *database*.

La questione fondamentale che pone il riconoscimento facciale di identificazione è quella del **bilanciamento** tra i **diritti fondamentali** dei cittadini coinvolti dall'utilizzo dello strumento e la **pubblica sicurezza**.

Secondo Ridaura libertà e sicurezza non sono in tensione dialettica, ma entrambe necessarie per l'esercizio dei diritti fondamentali. Allo stesso tempo, gli strumenti impiegati non si devono limitare a garantire genericamente la sicurezza, ma una **sicurezza democratica**, rispettosa dei diritti fondamentali.

I sistemi di riconoscimento facciale di identificazione risultano essere attualmente molto potenti e abbastanza diffusi (es. in Cina, USA, Colombia), ma pongono gravi problemi per i diritti fondamentali, specie quando impiegati per le pubbliche vie.

Infatti, in tale circostanza possono essere violati, non solo il diritto alla **privacy**, ma anche quello alla **libertà di movimento** e al **diritto di riunione**, in quanto la presenza di questi strumenti può creare un **effetto dissuasorio** nella cittadinanza.

Inoltre, questi sistemi possono violare la **presunzione di innocenza** in caso di **errori di identificazione**, molto spesso correlati a casi di **discriminazione** algoritmica dovuti ad un addestramento su *database* non sufficientemente rappresentativi di specifiche categorie (es. donne di

colore). Database creati con immagini dalla dubbia **provenienza e affidabilità**, spesso ottenute da aziende private o dai *social network*.

Nei casi di errore macroscopico all'errore di identificazione potrebbe seguire una detenzione, con conseguente e ingiustificata limitazione della **libertà personale**.

Questa incidenza sui diritti è aggravata dal fatto che l'uso di questi sistemi in luoghi pubblici avviene **senza il consenso** dei soggetti coinvolti. Nei casi più estremi, ciò può condurre a scenari di **vigilanza di massa**, difficilmente giustificabili in una società democratica.

Questione cruciale è quella relativa alla **proporzionalità** di un sistema di controllo generalizzato finalizzato all'individuazione di singoli individui. In uno stato di diritto, le intromissioni nella *privacy* per sicurezza devono avere giustificazione legale, obiettivo legittimo, necessità e proporzionalità, come stabilito dalla giurisprudenza della Corte EDU e della Corte di Giustizia dell'UE.

Con l'entrata in vigore dell'AI Act l'uso di questi strumenti da parte della polizia è consentito **solo in casi specifici e con molte garanzie**, ma risulta problematica l'esclusione dall'ambito di applicazione dei settori di: difesa, sicurezza nazionale e *intelligence*.

LUIS MOLINA VALBUENA, consulente legale di Castroalonso specializzato in diritto digitale, ha dedicato il proprio intervento al tema dell'**identità digitale** e dei pericoli derivanti dai **contenuti sintetici** iperrealistici, i quali possono attualmente essere creati facilmente e senza conoscenze tecniche specifiche.

Per identità digitale si intende l'insieme di tutti i dati attribuiti ad una persona nell'ambiente digitale, siano stati essi condivisi volontariamente o raccolti senza consenso.

Sulla base dei dati dell'identità digitale, anche quelli apparentemente più innocui come un audio o un video di pochi secondi o una foto, è possibile clonare la voce di una persona o creare **video iperrealistici**.

Tali contenuti possono essere potenzialmente generatori di danni enormi, che, laddove non integrano apposite fattispecie di reato, sono comunque idonei a distruggere reputazioni personali e carriere professionali.

Ci si troverebbe, dunque, in presenza di un **vuoto legale** nella protezione dell'**identità digitale autentica**. Nonostante le prescrizioni del GDPR, le garanzie complessive appaiono scarse e la repressione delle condotte criminali inefficace.

Così nel **diritto penale** manca una legislazione specifica contro la generazione fraudolenta di identità sintetiche. Allo stesso tempo, rimane aperta la questione della **responsabilità** di chi usa fraudolentemente le tecnologie in esame, che, ad opinione di Molina, dovrebbe ricadere in capo

all'autore della condotta, non al creatore. Per di più, si pone il problema della prova della generazione dei contenuti sintetici, in quanto potrebbe essere effettuata senza lasciare tracce digitali.

Le soluzioni proposte da Molina per far fronte ai problemi sollevati sono: (1) il **riconoscimento legale dell'identità digitale autentica**, da riconoscere come un bene giuridico da proteggere; (2) l'istituzione di un **registro volontario di identità verificate**; (3) l'**aggiornamento del codice penale** per tipificare specificamente l'uso malizioso di contenuti sintetici; (4) l'istruzione della cittadinanza; (5) l'adozione di buone pratiche di sicurezza.

STEFANO PIETROPAOLI, professore di filosofia del diritto presso l'Università di Firenze, ha chiuso la tavola rotonda con un intervento incentrato sulle sfide poste dalle nuove tecnologie per il diritto.

Punto di partenza dell'intervento è stato la constatazione che le nuove tecnologie, sebbene si manifestino in una dimensione digitale, generano un impatto reale, e non virtuale, sulla vita e i diritti fondamentali delle persone.

Tuttavia, nell'affrontare queste nuove sfide il diritto si trova a fare i conti con una **crisi dei concetti giuridici tradizionali** (e.g. proprietà, soggettività, nesso di causalità) non più adeguati a rappresentare i fenomeni della nuova dimensione digitale.

Secondo Pietropaoli tale crisi dei concetti si traduce in una **crisi del diritto stesso**. Il diritto, come ricordava Ermogeniano è una creazione dell'essere umano per l'essere umano ("*Omne ius hominum causa constitutum est*"). I codici normativi sono tali perché è possibile disobbedirvi, a differenza dei codici informatici ai quali le macchine non possono manifestare dissenso.

Lo spazio cibernetico è un ambiente in rapida evoluzione dove i diritti dei cittadini si trovano in una condizione di vulnerabilità in quanto privi di adeguata formazione e in balia dei grandi attori privati del mondo della tecnologia.

In breve, non può esservi tutela effettiva dei diritti se non si comprendono i meccanismi, i rischi e logiche dell'informatica. La **cittadinanza digitale** richiede di sapere come e perché ci si connette e quali conseguenze possono derivare da interazioni apparentemente neutrali.

Infine, Pietropaoli ha affrontato il tema della **spiegabilità algoritmica**, sostenendo l'insensatezza della ricerca dei motivi dei risultati degli algoritmi. Essi sono privi di una razionalità in senso umano: sono basati sulla statistica, sulla logica induttiva, non deduttiva. A partire da ciò è possibile polemizzare contro l'idea stessa di *decisione* algoritmica. Infatti, secondo il relatore, le decisioni algoritmiche non sarebbero decisioni: le macchine non decidono, perché non possono esprimere una volontà. Non bisogna antropomorfizzarle. Semmai, vi è, a monte, la decisione umana di utilizzare uno strumento per risolvere un problema.

GIOVEDÌ 15 MAGGIO 2025

IA E LAVORO

Relatori: Ivan Antonio **Rodríguez** Cardo, Lucia **Montejo**, Javier **Campa**, Isabel **Santos**

Moderatore: Maria Antonia **Castro** Arguelles

Contenuto: I lavori di giovedì 15 maggio 2025 sono stati aperti nell'aula "Severo Ochoa" dell'edificio storico dell'Università di Oviedo con la tavola rotonda su "*IA y Trabajo*", nella quale si sono confrontati rappresentanti del mondo accademico (**Rodríguez**), del mondo sindacale (**Montejo**, **Campa**) e datoriale (**Santos**).

IVAN ANTONIO RODRÍGUEZ CARDO, professore di diritto del lavoro presso UNIOVI e membro del CEISIA, ha offerto una panoramica generale dalla **prospettiva giuridica** dell'**impatto** dell'IA nelle relazioni di lavoro.

Nel perseguimento di tale scopo, la relazione del professore Rodríguez ha preso le mosse da alcune considerazioni di economia dell'IA.

In particolare, dal punto di vista imprenditoriale l'IA è percepita come opportunità per efficientare ed economizzare i processi produttivi. In tale ottica, essa può essere utilizzata per organizzare le mansioni dei lavoratori e per adottare decisioni autonome (c.d. "*gestione algoritmica del lavoro*"). Inoltre, esisterebbe una certa aspettativa che possa migliorare l'occupazione generale, dando avvio alla creazione di nuovi posti di lavoro.

Tuttavia, dal punto di vista giuridico l'IA costituisce un fenomeno da regolare per tutelare i diritti dei lavoratori, per evitare la disumanizzazione delle decisioni lavorali e soprattutto per realizzare una transizione giusta che non porti alla perdita improvvisa di posti di lavoro.

A ciò si aggiungono i rischi legati ad un controllo eccessivo dei lavoratori, al trattamento dei dati personali, ai bias che conducono a discriminazioni difficilmente verificabili in ragione dell'opacità dei sistemi di IA.

Rispetto a siffatte questioni, l'AI Act appare "timido" da una prospettiva lavoristica. Infatti, esso non presenta norme lavoristiche *stricto sensu*, nel senso di regolare il rapporto tra lavoratore e datore di lavoro, ma una norma sulla sicurezza del prodotto IA. In altri termini, si limita a disciplinare la produzione di un sistema di IA, indipendentemente dalla qualifica di lavoratore dell'utente finale.

Tuttavia, alcune delle pratiche vietate dal regolamento hanno un diretto impatto sul mondo del lavoro. Così, non è ammesso l'utilizzo dei sistemi di **riconoscimento delle emozioni** e di **categorizzazione biometrica per dedurre dati sensibili** (come ad esempio l'affiliazione sindacale).

Allo stesso tempo, le norme sulla trasparenza e sulla supervisione umana dettate per i sistemi ad alto rischio non prevedono (direttamente) che i lavoratori o i loro rappresentanti possano partecipare alle attività di controllo dell'IA.

Differente è il caso della disciplina della **Direttiva europea sul lavoro nelle piattaforme digitali**, che prevede, invece, il coinvolgimento dei lavoratori negli adempimenti sulla trasparenza e sulla supervisione dell'IA.

Infine, l'**articolo 64 dello Statuto dei lavoratori spagnolo** richiede che il comitato aziendale ha diritto di essere informato dall'azienda sui parametri, sulle regole e sulle istruzioni su cui si basano gli algoritmi o i sistemi di IA che influenzano il processo decisionale che può avere un impatto sulle condizioni di lavoro, sull'accesso e sul mantenimento dell'occupazione, inclusa la profilazione.

Al netto di ciò, Rodríguez ha manifestato dei dubbi rispetto ad un approccio di tutela incentrato principalmente sulla trasparenza. Infatti, da una parte, permane il problema dell'opacità algoritmica, dall'altra, spesso gli imprenditori acquistano il sistema da terzi o ne esternalizzano la gestione. Piuttosto che concentrarsi sull'apertura della scatola nera, il relatore ha auspicato dei controlli incentrati sui risultati che verifichino la correttezza dei risultati dell'IA, ricorrendo, ad esempio, alla prova statistica.

LUCIA MONTEJO, membro di “*Confederación Sindical de Comisiones Obreras*”, ha partecipato alla discussione fornendo la prospettiva dei lavoratori che subiscono gli effetti dell'impiego dell'IA.

Nell'affrontare la questione Montejo ha ricordato che l'IA non va considerata come un blocco monolitico, ma ricomprende tecnologie diverse. Inoltre, ha sottolineato che le imprese impiegano da svariato tempo algoritmi, ma solo di recente si sta assistendo ad una vera e propria “*iper-automazione*”.

Poiché le applicazioni dell'IA nel mondo del lavoro sono diverse, coinvolgono variamente i diritti dei lavoratori e, pertanto, necessitano **strategie sindacali dedicate**.

Secondo Montejo l'attuale quadro normativo fornisce delle tutele lavorative inadeguate: da una parte, l'AI Act non si concentrerebbe adeguatamente sulla protezione dei diritti dei cittadini; dall'altra, l'articolo 64 dello Statuto dei lavoratori spagnolo fornisce una tutela limitata. In breve, è inutile ricevere dall'impresa 200 pagine di codice sorgente dell'algoritmo, se le associazioni sindacali non dispongono dei mezzi e del personale per renderlo **intelligibile**. Al contrario, è necessario che

l'obbligo di trasparenza venga adempiuto fornendo una spiegazione effettiva del funzionamento algoritmico.

Pertanto, non tutti i problemi possono essere risolti conoscendo i meccanismi decisionali algoritmici.

A questi **problemi** che derivano specificamente dall'impiego dell'IA nel contesto lavorale, si aggiungono quelli **strutturali** del sistema produttivo e del mercato del lavoro che dall'IA vengono, però, amplificati. Ad esempio, anche in presenza di un dataset di addestramento robusto, l'IA potrebbe replicare i bias (discriminatori) preesistenti nella società.

Attualmente, ha sottolineato Montejo, i sindacati appaiono impreparati nella gestione di cambiamenti così rapidi, poiché l'alfabetizzazione digitale dei rappresentanti richiede tempo e supporto istituzionale.

In termini più generali, l'ascesa dell'IA si lega ad un processo di accelerazione della concentrazione della ricchezza a livello mondiale, con ricadute sulle attuali relazioni lavorative, sul futuro del lavoro e della società in generale. Ciò, ha concluso Montejo, imporrà di affrontare la questione della fiscalità delle tecnologie.

JAVIER CAMPA, membro di “*Unión General de Trabajadoras y Trabajadores*”, riconosce che l'IA può essere portatrice sia di opportunità che di rischi per i lavoratori.

Può facilitare le condizioni lavorative, ma anche sacrificare numerosi posti di lavoro in nome del profitto. In altre parole, l'IA ha il potenziale per realizzare una **riconversione globale** del mondo del lavoro, simile alle precedenti rivoluzioni industriali.

Proprio per questo, Campa ha sottolineato l'esigenza di un quadro normativo che regoli l'implementazione dell'IA nelle aziende in modo trasparente, che coinvolga i lavoratori nella governance e che tuteli i loro diritti. Soprattutto, ha concluso Campa, è necessario che si realizzi una transizione giusta, che permetta una riqualificazione dei lavoratori tramite la loro formazione continua e che non lasci indietro chi non si trovi nella condizione di potersi riqualificare.

ISABEL SANTOS, membro di “*Federación Asturiana de Empresarios*”, ha arricchito il dibattito con l'apporto del punto di vista datoriale.

Per le imprese, ha sottolineato Santos, l'IA non è il futuro, ma il presente. Sono costrette ad investire in questa tecnologia per non essere lasciate indietro. Per gestire al meglio questo rapido cambiamento è necessario un coordinamento tra impresa e lavoratori.

L'IA costituisce una naturale evoluzione dei processi produttivi, intervenendo in settori specifici.

Alcuni potrebbero essere del tutto sostituiti come i lavori amministrativi basati sul cartaceo o quelli manuali di base.

Altri potrebbero utilizzare l'IA come un supporto migliorativo (lavori creativi).

Infine, l'IA creerà nuovi posti di lavoro nel settore tecnologico.

Santos ha concluso prospettando alcuni dei rischi della transizione digitale, come il divario tra paesi capaci di convertire e formare la propria forza lavoro qualificata, eventualmente con il supporto di apposite sovvenzioni pubbliche, nonché quello di affidarsi ad imprese terze per l'approvvigionamento e la gestione dell'IA.

IA E SANITÀ

Relatori: Sergio **Calleja**, Elena **Díaz** Rodríguez, Pedro José **Espinosa** Prados

Moderare: Daniel Jove Villares

Contenuto: La seconda tavola rotonda della mattina del 15 maggio ha visto confrontare settore sanitario (**Calleja** e **Espinosa**) e accademia (**Díaz**) sulle più recenti applicazioni dell'IA nella medicina.

SERGIO CALLEJA, neurologo presso “*Hospital Universitario Central de Asturias*” (HUCA), ha contribuito alla discussione riportando la sua personale **esperienza clinica**.

Calleja ha aperto l'intervento descrivendo la medicina come una relazione tra chi presenta un problema (il paziente) e chi cerca di risolverlo (il medico). Una relazione che si basa sull'**incertezza**, ovviamente, del paziente rispetto alla propria patologia, ma anche del medico circa la diagnosi e il trattamento. Concetto che, ricorda Calleja, è sintetizzato efficacemente dalle parole di William Osler: “*La medicina è la scienza dell'incertezza e l'arte della probabilità*”. Questa incertezza grava pesantemente sulla nostra società: gli errori medici, ha ricordato Calleja, sono la terza causa di morte più comune in un paese come gli USA. Tuttavia, tale circostanza appare fisiologica. Infatti, moltiplicando l'incertezza per il numero di pazienti che ogni medico visita ogni giorno appaiono evidenti le difficoltà della professione medica.

Pertanto, la medicina da sempre accoglie le innovazioni capaci di aiutarla a ridurre questo margine di incertezza. L'IA potrebbe potenzialmente aiutare in questa direzione, aiutando a ridurre quegli errori che non sono il frutto di negligenza, ma della complessità della professione stessa.

L'IA pretende di fornire direttamente soluzioni. Tuttavia, è necessario riflettere su cosa si cela dietro la domanda di cura del paziente e la risposta del medico. Quest'ultimo genera ipotesi, le confronta con la realtà, raccoglie dati, ne valuta l'origine e i possibili bias. Si confronta con una complessità grazie alla propria esperienza, che non è né predeterminata né insegnabile all'università. Un medico esperto agisce per una intuizione che si fonda dalla ripetuta esposizione a casi reali.

Ricevendo risposte direttamente dall'IA, l'acquisizione della conoscenza per esperienza dell'operatore sanitario rischia di annacquarsi.

Il metodo basato sull'evidenza empirica ha trasformato la medicina da sapere prescientifico a scienza, portando sia a indubbi progressi sia a sottovalutare quelle capacità non numeriche intraducibili in formule matematiche come l'empatia, l'abilità di raccogliere dati qualitativi e l'intuizione. Si stima che quasi il 50% dei sintomi dichiarati dai pazienti siano medicalmente

inspiegabili. Inoltre, per molti pazienti non si raggiunge mai una diagnosi. In questi casi, il ruolo del medico non è trovare la risposta, ma sostenere il paziente.

Nel corso degli ultimi anni, ha ricordato Calleja, la medicina è finita per essere influenzata dalle pressioni degli interessi economici: delle industrie farmaceutiche, della sanità privata, dei titolari dei brevetti. Tutti interessati ad ottenere la maggiore deregolazione possibile nel settore.

La qualità della scienza medica risente della richiesta di generare profitti a breve termine, ad esempio con prescrizioni farmaceutiche inutili o con *screening* indiscriminati.

Questa divagazione di contesto appare fondamentale per interrogarci su chi finanzia e a quali interessi sono subordinate i sistemi di IA nel settore della sanità.

Infine, ha concluso Calleja, bisogna chiedersi: chi controllerà i bias? Chi sarà responsabile in caso di errore? Come si formeranno i nuovi professionisti sempre meno esposti all'esperienza clinica empirica?

ELENA DÍAZ RODRÍGUEZ, professoressa di fisiologia presso UNIOVI e membro del CEISIA, apre il proprio intervento evidenziando come nel campo della sanità siano già in uso moltissime applicazioni di IA.

Uno dei primi utilizzi dell'IA nel settore, ricorda Díaz, è stato in ambito oncologico, nel quale gli algoritmi sono in grado di analizzare le biopsie ottenendo dei risultati superiori a quelli degli specialisti. Dunque, un simile strumento potrebbe sia incrementare l'accuratezza delle diagnosi sia ridurre il carico di lavoro dei medici.

Díaz presenta due studi di applicazione dell'IA alla **cronobiologia della riproduzione**.

Per cronobiologia si intende la disciplina che studia i ritmi circadiani: quei processi fisiologici, biochimici, genetici che si ripetono ogni 24 ore (es. sonno-veglia, pasti). Il ritmo è dettato da un orologio biologico nel sistema nervoso centrale, influenzato dalla luce attraverso la melatonina. L'alterazione di questi ritmi è chiamata *cronodisruption*. Cause comuni: sono l'uso di schermi di notte, i viaggi e gli orari di lavoro irregolari. La ricerca si concentra su come la *cronodisruption* influenzi la salute riproduttiva, in particolare il parto prematuro.

Il **primo studio** mirava a identificare nuovi fattori di rischio per il parto prematuro legati alla *cronodisruption* e a sviluppare un modello matematico per predire il rischio. Sono stati raccolti dati da 481 parti, di cui 257 prematuri. Sono state considerate variabili cliniche abituali e variabili legate alla cronobiologia ottenute tramite questionari telefonici (abitudini e qualità del sonno, orari di sonno nei giorni feriali e festivi, dormire con o senza luce in camera). Il campione finale per l'algoritmo è stato di 223 parti. La affidabilità del modello è risultata alta, con una percentuale di successo del 97,3%.

Il risultato più importante dello studio è la classificazione delle variabili. Il fattore più importante è risultato il peso stimato, elemento già noto, che ha così confermando la validità dell'algoritmo. Invece, la seconda variabile più importante individuata dall'IA è stata, sorprendentemente, la luce nella stanza durante la notte, superando in importanza variabili classiche come l'età della madre. L'albero di decisione costruito dall'algoritmo aiuta i clinici a stimare il rischio di parto prematuro partendo dalle variabili: l'IA ordina i **fattori di rischio** per importanza, fornendo un **aiuto per una rapida diagnosi**.

Il **secondo studio** ha esplorato marcatori di *cronodisruption* nella salute riproduttiva generale delle donne, concentrandosi su una popolazione molto esposta a determinate malattie: le infermiere che lavorano per turnazione. I dati sono stati raccolti tramite questionari su otto aree sanitarie e nove ospedali nelle Asturie. Le domande riguardavano dati socio-lavorativi, salute riproduttiva, qualità del sonno e orari di alimentazione. I problemi studiati erano: tempo di concepimento, durata della gestazione, problemi durante la gestazione e malattie riproduttive generali. Sono stati utilizzati tre diversi algoritmi. Non sempre questi ultimi conducevano ai medesimi risultati, pertanto, si è provveduto ad una loro integrazione. Per il tempo di concepimento, la variabile più importante è risultata essere il punto medio di alimentazione nei giorni di riposo, mentre la seconda variabile più importante era la qualità del sonno. Per la durata della gestazione, la variabile principale era il “*jet lag alimentare*” (la differenza negli orari dei pasti tra i giorni lavorativi e quelli di riposo).

Alla luce dei propri studi, Díaz ha concluso affermando che l'IA costituisce uno strumento di sicura utilità per condurre ricerche simili nel settore.

PEDRO JOSÉ ESPINOSA PRADOS, capo progetto dello spazio dati salute di “*Servicio de Salud del Principado de Asturias*” (SESPA), dedica il proprio intervento all'esposizione dei risultati di un progetto di creazione di due **repository centralizzati di dati clinici dei pazienti, aggiornati in tempo reale**.

La provenienza dei dati da diversi sistemi informativi sanitari ha richiesto un intenso lavoro di normalizzazione degli stessi.

Il **repository primario** serve per l'uso assistenziale; quello **secondario** è dedicato all'uso dei dati per ricerca, sfruttamento, estrazione dati e analisi avanzata.

Il **repository primario** permetterà l'alimentazione di un portale che fornirà una **visione globale e riassuntiva del paziente**, adattata all'utilizzo da parte di molteplici profili professionali che potranno visualizzarne storia e note cliniche con una ricerca simile a quella Google.

Il repository secondario segue un modello a tre livelli: *Bronze* (dati grezzi), *Silver* (dati trasformati e orientati al modello), *Gold* (dati validati e standardizzati).

I dati sono archiviati localmente nel SESP, con la possibilità di utilizzare un *cloud* per progetti di ricerca su larga scala o collaborazioni esterne.

La re-identificazione dei pazienti viene ridotta al minimo tramite processi di pseudo-anonimizzazione.

In sostanza, ai ricercatori viene fornito uno **spazio dati sicuro** per elaborare i loro *dataset*.

ECONOMIA E SOSTENIBILITÀ DELL'IA

Relatori: Luis **Molina** Valbuena, Carlos **Ibáñez** Sánchez

Moderatore: Augustina Bouchet

Contenuto: La tavola rotonda, tenutasi presso la sede del CEISIA di Lugones, si è concentrata sull'impatto dell'IA nella società in termini sia di sostenibilità in un senso esteso, che include l'etica, i diritti fondamentali e l'equità sociale (Molina) sia in termini di economia, con particolare riferimento ai costi e le sfide delle imprese emergenti nel settore del *legal tech* (Ibáñez).

LUIS MOLINA VALBUENA, consulente legale di Castroalonso specializzato in diritto digitale, apre la discussione sulla questione della **sostenibilità** dell'IA sottolineando che tale concetto non deve ritenersi limitato all'ambiente, ma anche inclusivo dei diritti umani fondamentali.

In altri termini, sostiene Molina, la sostenibilità non può misurarsi solo in termini di impronta carbonica, ma anche di etica, giustizia algoritmica e responsabilità.

In particolare, l'IA genera moltissimo valore, ma a vantaggio solo di pochissimi individui, che concentrano su di sé un potere non democratico e privo di contrappesi.

La fonte di tale potere sono i dati, "**il petrolio del XXI secolo**", estratti direttamente dagli utenti, che li forniscono consapevolmente o inconsapevolmente, e, poi, impacchettati, monetizzati e rivenduti.

Un'economia digitale, sostiene Molina, non può essere sostenibile se ignora la giustizia sociale; non può essere verde se è insaziabile di energia, opaca nella governance e inefficace nei diritti. Così, in caso di errori algoritmici, l'assenza di responsabilità conduce ad un'economia ingiusta e immorale. L'addestramento di un LLM come GPT-4 produce un impatto ambientale paragonabile al consumo di energia di 100 famiglie europee in un anno e il suo utilizzo produce delle "impronte di carbonio invisibili".

Premesso ciò, Molina esclude la perseguibilità della sostenibilità algoritmica in assenza di una **regolamentazione giuridicamente vincolante**.

Primi passi in tale direzione sono rappresentati dall'AI Act e dalla **Direttiva sulla Corporate Sustainability Due Diligence** che impone di prevenire i danni ambientali e sociali lungo la catena del valore causati da sistemi automatizzati e di adottare politiche di *due diligence* effettive ed efficaci.

Molina riconosce l'indubbio costo di conformità per le aziende, ma ritiene che per garantirla debbano implementarsi delle sanzioni dissuasive.

Al netto di ciò, Molina evidenzia la possibilità di un **rischio sistemico**: il diffondersi globale di modelli opachi e addestrati su dati distorti. Così, ad esempio, un modello di IA può bloccare l'accesso al credito o al lavoro, normalizzando prassi discriminatorie.

Sulla base di ciò, conclude Molina, governare l'IA è l'unico modo per garantire che la stessa persegua finalità benefiche per la società.

CARLOS IBÁÑEZ SÁNCHEZ, CEO di ADEA, contribuisce alla discussione concentrandosi sulla sostenibilità dell'IA dal punto di vista economico-aziendale.

L'obiettivo della società ADEA diretta da Ibáñez è quello di fornire le migliori soluzioni tecnologiche agli operatori legali. Per fare ciò sono stati identificati i cinque punti chiave che tali soggetti ricercavano in un servizio di questa tipologia: **efficienza, accuratezza, costi contenuti, accessibilità e trasparenza**.

Per una impresa come ADEA è fondamentale mantenere il ritmo dell'evoluzione tecnologia, senza però appiattirsi su un aggiornamento fine a sé stesso, bensì considerando quale sia la migliore tecnologia possibile in relazione ai casi d'uso reali.

Inizialmente la azienda di Ibáñez utilizzava reti neurali per processare documenti legali. Ciò richiedeva l'accesso a grandi volumi di dati, la loro classificazione ed etichettatura e l'impiego di un'alta capacità computazionale.

Con l'arrivo dei modelli fondazionali e dell'IA generativa, si è passati a un **sistema ibrido**: le reti neurali sono impiegate per gli ambienti deterministici (modelli, template) perché più veloci ed economiche; l'IA generativa per quelli non deterministici, che richiedono interpretazione e analisi del contesto.

I **costi** del *machine learning* sono molto più contenuti di quelli per IA generativa: il primo richiede solamente una CPU, la seconda costose GPU e un tempo di addestramento maggiore

Tuttavia, l'IA generativa permette di generare dati sintetici a partire da un piccolo set di partenza, consentendo di aumentare artificialmente il volume dei dati di addestramento.

Dopo aver descritto le sfide di sviluppo, le caratteristiche e il funzionamento del modello di processamento di documenti legali sviluppato da ADEA Sánchez conclude affrontato il tema della **singularità legale**.

In particolare, a differenza della singularità tecnologica, Sánchez ritiene che vi siano già gli strumenti per la completa automatizzazione di moltissimi procedimenti legali. Intravede la possibilità di un diritto completo, senza lacune, auto-eseguibile e autonomo. Tuttavia, ciò richiede una certa gradualità, onde guidare al meglio un eventuale cambiamento in tal senso.

VENERDÌ 16 MAGGIO 2025

IA E GIUSTIZIA

Relatori: Jordi **Nieva** Fenoll, Alfredo **Rodríguez** del Blanco, Aitor **Cubo**.

Moderatore: Marcos Loredó Colunga

Contenuto: L'aula "Severo Ochoa" dell'edificio storico dell'Università di Oviedo ha accolto l'ultima tavola rotonda in tema di "*IA y Justicia*". I relatori hanno esplorato la sempre crescente integrazione dell'IA nel sistema giudiziario, illustrandone: il potenziale innovativo per i processi (Nieva); i rischi per i diritti fondamentali nel processo penale (Rodríguez); le recenti sperimentazioni nell'amministrazione della giustizia spagnola (Cubo).

JORDI NIEVA FENOLL, professore di diritto processuale presso l'Università di Barcellona, avvia la discussione occupandosi del tema dell'**impiego dell'IA nei processi** giudiziari.

In tale ottica, Nieva sottolinea la differenza tra **automazione** dei procedimenti e uso dell'IA nei procedimenti. Attualmente, ricorda Nieva, il sistema giudiziario si concentra sulla prima, con lo scopo di ridurre i costi. Tuttavia, a tale cambiamento deve necessariamente accompagnarsi una **modifica delle norme procedurali**, altrimenti obsolete.

I procedimenti non coinvolti da questo cambiamento (la minor parte) si caratterizzano per il necessario impiego di tutte le prerogative del giurista: ricostruzione del fatto, valutazione della prova, interpretazione del diritto.

Allo stesso tempo, è proprio in questo ambito che l'IA potrebbe dispiegare tutto il suo potenziale. Attualmente la ricostruzione dei fatti e la valutazione delle prove si basa sull'esperienza del giurista, che risente di bias e pregiudizi che questa tecnologia può aiutare a mitigare. Ancora, l'IA potrebbe **aiutare a valutare** perizie tecniche, ad analizzare prove documentali in modo rapido e a suggerire ipotesi ricostruttive.

Tutto ciò, ha concluso Nieva, deve sempre avvenire sotto il **controllo** dei giuristi e dello Stato, non di aziende tecnologiche terze.

ALFREDO RODRÍGUEZ DEL BLANCO, avvocato penalista e membro del CEISIA, contribuisce alla discussione affrontando i problemi dell'uso dell'IA nel processo penale per le garanzie processuali fondamentali.

Le caratteristiche dell'IA, sottolinea Rodríguez, la rendono particolarmente utile nelle attività di **indagine**. Tuttavia, per evitare una pericolosa compressione delle garanzie processuali, è necessario marcarne i confini, onde evitarne una inaccettabile violazione.

Per fare ciò è imprescindibile, a parere di Rodríguez, analizzare la giurisprudenza nazionale, costituzionale e internazionale sui principali diritti fondamentali coinvolti nella fase delle indagini: uguaglianza, libertà personale, *privacy*.

L'**uguaglianza** è pregiudicata dal possibile utilizzo di risultati dell'IA viziati da bias.

La **libertà personale** da quei sistemi opachi impiegati per l'applicazione di una misura cautelare.

La **privacy** da un impiego indiscriminato e generalizzato dei dati personali, nonché dall'impiego dell'IA nelle attività captative.

A ciò si aggiungono i rischi legati all'introduzione di prove false attraverso l'IA generativa, di difficile valutazione con i mezzi attualmente a disposizione dell'autorità giudiziaria.

Ciò suggerirebbe che in assenza di un adeguato controllo umano, la giustizia potrebbe disumanizzarsi.

Per evitare questo rischio è imprescindibile per Rodríguez una adeguata **alfabetizzazione digitale** degli avvocati, nell'ottica della miglior difesa possibile delle parti.

AITOR CUBO, direttore generale della Trasformazione Digitale del Ministero della Giustizia spagnolo, ha chiuso la tavola rotonda illustrando il ruolo del Ministero della Giustizia spagnolo nel processo di transizione digitale.

La digitalizzazione della giustizia è vista da Cubo come un modo per rendere un servizio pubblico più efficiente e accurato per il cittadino.

In tale ottica, il Ministero della Giustizia spagnolo sta lavorando tanto a progetti di automazione della giustizia, quanto di impiego dell'IA.

Con specifico riferimento all'IA, Cubo propone una distinzione basata sul diverso impatto nelle decisioni giurisdizionali: senza impatto, ad impatto indiretto, ad impatto diretto. Quest'ultima ipotesi, naturalmente, richiede le maggiori cautele.

Per realizzare la transizione è di cruciale importanza l'orientamento del dato al settore giustizia. I dati devono essere generati già nel momento iniziale dei processi.

Cubo ha concluso richiamando alcune delle possibili applicazioni nel settore.

Per quanto riguarda i **processi di automazione**, Cubo ha ricordato: la **cancellazione d'ufficio degli antecedenti penali alla scadenza dei termini legali**; l'automazione dei **processi monitori**; la **gestione delle cauzioni e pagamenti**.

Per quanto riguarda gli **applicativi di IA**, Cubo ha menzionato: l'analisi e la classificazione documentale; la semplificazione dei testi legali, la trascrizione automatica delle registrazioni; l'assistenza giudiziale nell'analisi dei fascicoli di grandi dimensioni e nell'individuazione di relazioni tra le parti.

CONFERENZA DI CHIUSURA

Relatore: Lorenzo **Cotino** Hueso

Interventi: Miguel Ángel Presno Linera, Ignacio Villaverde (rettore di UNIOVI).

LORENZO COTINO HUESO ha concluso la settimana di lavori portando la sua duplice prospettiva di professore di diritto costituzionale e direttore della “*Agencia Española de Protección de Datos*” (AEPD).

Fondamentale per Cotino è la **collaborazione tra accademia e istituzioni** pubbliche nel settore dell’IA.

Quest’ultima deve essere impiegata per aumentare la produttività della p.a., nonché come strumento per realizzare l’effettività dei diritti.

Inoltre, Cotino ha auspicato il riconoscimento di un “**diritto costituzionale all’innovazione tecnologica**” cui corrisponderebbe un dovere dello Stato a rimuovere gli ostacoli e ad avviare iniziative per un simile sviluppo.

Come istituzione pubblica l’AEPD si sta muovendo in prima linea per identificare le possibili “finalità di uso” dell’IA in tutte le unità, valutandone, al contempo, i rischi. L’obiettivo per l’AEDP è aumentare la propria produttività, ma soprattutto diventare un esempio virtuoso per il resto del settore pubblico, nel quale l’IA viene implementata nel rispetto di *privacy* e sicurezza.

A tal riguardo Cotino ha proposto l’istituzione di **registri pubblici degli algoritmi**: strumenti essenziali per la trasparenza e la supervisione dell’AI da parte della pubblica autorità.

Da un punto di vista dogmatico Cotino ha criticato la scarsa attenzione dei costituzionalisti verso le consolidate **metodologie di analisi dei rischi** e di studi di impatto sulla protezione dei dati, le quali permetterebbero una maggiore accuratezza nelle valutazioni di impatto sui diritti fondamentali.

Ancora, Cotino evidenzia la necessità di coordinamento interpretativo tra GDPR e AI Act, i quali presentano alcuni ambiti di sovrapposizione, primo fra tutti quello sanzionatorio. Ciò si traduce in un ampliamento delle competenze dell’AEPD esteso a quegli aspetti dell’IA strettamente legati al trattamento dei dati personali.

Più in generale, ha concluso Cotino, appare necessario dotare l’impiego dell’IA in particolari settori sensibili (e.g. giustizia) di una **adeguata copertura giuridica**, che ne stabilisca gli standard minimi di applicazione.