



UNIVERSITA' DEGLI STUDI DI ROMA TOR VERGATA

***European Digital Law of the Person, of the Contract
and of the Technological Marketplace - EUDILA
Cattedra Jean Monnet del Progetto ERASMUS +***

***“IA E RESPONSABILITA' CIVILE:
IL DIRITTO DI FRONTE ALL'ALGORITMO”
PIERFRANCESCO TUCCI
0365949***

Anno accademico 2024/2025

INDICE

Abstract

1. La responsabilità civile e l'autonomia decisionale dell'intelligenza artificiale:

- 1.1 Definizione e classificazione dell'IA
- 1.2 La responsabilità extracontrattuale nel diritto civile italiano
- 1.3 La disciplina della responsabilità da prodotto difettoso
- 1.4 Il contributo del GDPR

2. Casi di Studio: Il caso Uber e Tesla.

- 2.1 Analisi del caso Uber
- 2.2 Implicazioni giuridiche e responsabilità coinvolte nel caso Uber
- 2.3 Analisi del caso Tesla
- 2.4 Implicazioni giuridiche e responsabilità coinvolte nel caso Tesla

3. Prospettive Future e Proposte di Riforma

- 3.1 Necessità di un quadro normativo armonizzato
- 3.2 Proposte per una regolamentazione efficace della responsabilità da IA

4. Bibliografia

Abstract

Negli ultimi anni, l'intelligenza artificiale (IA) ha conosciuto un'espansione senza precedenti, divenendo uno strumento sempre più centrale in numerosi ambiti della vita economica e sociale. Dalla medicina all'agricoltura, dalla finanza ai trasporti, le applicazioni dell'IA stanno trasformando radicalmente le modalità con cui si prendono decisioni, si analizzano dati e si interagisce con l'ambiente circostante. In particolare, l'impiego dell'IA nei sistemi autonomi e automatizzati, come i veicoli a guida assistita, solleva questioni giuridiche nuove e complesse, soprattutto in relazione alla responsabilità civile per danni causati da tali tecnologie. Il quadro normativo vigente, nato in un contesto storico privo delle attuali tecnologie, mostra oggi evidenti limiti nell'affrontare situazioni in cui la responsabilità si colloca in una zona intermedia tra l'errore umano e il malfunzionamento dell'intelligenza artificiale. Le nozioni classiche di colpa, nesso causale, attività pericolosa e difetto del prodotto sono messe alla prova da tecnologie che agiscono in modo parzialmente autonomo, talvolta imprevedibile e che apprendono e modificano il proprio comportamento nel tempo. In questo contesto, il giurista si trova a fronteggiare una duplice sfida: da un lato, interpretare norme costruite per realtà radicalmente diverse, dall'altro, contribuire alla definizione di un nuovo paradigma regolatorio che sia in grado di tutelare le vittime, promuovere la trasparenza degli algoritmi e al contempo, non soffocare l'innovazione.

La tesina, quindi, si propone di analizzare il tema della responsabilità civile nel contesto dell'intelligenza artificiale, con particolare attenzione al quadro normativo europeo, in fase di evoluzione attraverso l'*Artificial Intelligence Act* e le proposte di direttive sulla responsabilità da IA e da prodotto difettoso. A supporto dell'analisi teorica, verranno presentati due casi studio: l'incidente mortale avvenuto nel 2018 in Arizona, che ha visto coinvolto un veicolo a guida autonoma sviluppato da Uber e il caso Tesla, mettendo in risalto le difficoltà pratiche nella ricostruzione della catena di responsabilità in presenza di un agente non umano, i nodi critici e le prospettive di riforma, nella consapevolezza che il diritto non può più ignorare il ruolo degli algoritmi nella costruzione della realtà, ma deve farsi carico della loro regolamentazione con strumenti adeguati, equi e aggiornati.

1) La responsabilità civile di fronte all'autonomia decisionale dell'intelligenza artificiale

1.1 Definizione e classificazione dell'IA

L'intelligenza artificiale, un tempo confinata alla letteratura fantascientifica, rappresenta oggi una delle innovazioni tecnologiche più importanti e decisive del XXI secolo. Con il termine *intelligenza artificiale* si intende, in senso generale, l'insieme di tecnologie e sistemi informatici progettati per svolgere funzioni normalmente associate all'intelligenza umana: apprendere dall'esperienza, risolvere problemi complessi, prendere decisioni autonome, riconoscere immagini e linguaggio, interagire con l'ambiente. L'intelligenza artificiale di cui si parla più frequentemente oggi è la cosiddetta IA generativa. Questa forma di IA si basa su modelli di apprendimento profondo (deep learning), che analizzano enormi quantità di dati e imparano a generare contenuti simili a quelli umani. Il suo impatto è visibile in molti ambiti: dalla scrittura all'arte digitale, dalla comunicazione aziendale fino all'istruzione e al marketing. Non solo, si fa molto riferimento anche all'IA reattiva, che invece non si occupa di creare contenuti, ma si limita a classificare, riconoscere o reagire a stimoli esterni, secondo regole predefinite. È il caso, ad esempio, dei sistemi di sorveglianza basati su videocamere o sensori infrarossi, che individuano movimenti sospetti grazie ad algoritmi impostati dall'uomo oppure ancora, è il caso di molte funzionalità di sicurezza integrate nei nostri smartphone, come il Face ID o il rilevamento automatico di incidenti; tutti questi sfruttano IA reattive per rispondere a situazioni in tempo reale. L'intelligenza artificiale permea ormai ogni aspetto della vita quotidiana, dal monitoraggio della salute tramite app di fitness, agli assistenti vocali come Siri di Apple e Google Assistant, fino alle sofisticate soluzioni aziendali che potenziano l'efficienza di vendite, logistica e customer service. Colossi tecnologici quali Apple, Google, Microsoft e Amazon investono ingenti risorse nello sviluppo e nell'integrazione dell'IA, generando un impatto tangibile su milioni di utenti in tutto il mondo. Nella sua declinazione più avanzata, l'IA non si limita a seguire istruzioni predefinite, ma è in grado di modificare il proprio comportamento in base ai dati ricevuti, spesso sfuggendo al controllo diretto dei programmatori originari. Questo livello di autonomia solleva interrogativi profondi e inediti sotto il profilo giuridico: chi è responsabile se una decisione presa da un sistema intelligente causa un danno?

1.2 La responsabilità extracontrattuale nel diritto civile italiano

Nel sistema giuridico italiano, la responsabilità civile si fonda tradizionalmente su due pilastri: la responsabilità contrattuale, prevista dall'articolo 1218 del Codice Civile, che riguarda il mancato rispetto di obblighi derivanti da un accordo tra le parti, e la responsabilità extracontrattuale, sancita dall'articolo 2043. Quest'ultima, costituisce la base dell'intera architettura risarcitoria del diritto civile e prevede che *“qualunque fatto doloso o colposo che cagiona ad altri un danno ingiusto, obbliga colui che ha commesso il fatto a risarcire il danno”*. Il presupposto essenziale, dunque, è la presenza di un comportamento umano riconducibile al soggetto agente, da cui derivi un danno ingiusto e un nesso causale tra condotta e danno. Tuttavia, quando entriamo nel campo dell'intelligenza artificiale, questa logica mostra i suoi limiti: nei casi in cui una macchina prenda decisioni complesse in autonomia, sulla base di algoritmi e processi di apprendimento, risulta spesso difficile, se non impossibile, individuare con certezza una responsabilità umana diretta.

Nel diritto civile italiano sono stati elaborati ulteriori articoli volti ad ampliare la responsabilità oggettiva, come l'art. 2050 c.c. (*“Responsabilità per l'esercizio di attività pericolose”*), che impone un regime probatorio aggravato al soggetto che esercita un'attività intrinsecamente rischiosa. Alcuni studiosi hanno suggerito di inquadrare l'utilizzo dell'intelligenza artificiale avanzata tra le *“attività pericolose”*, attribuendo così al produttore o all'utilizzatore l'obbligo di risarcimento, salvo che venga fornita una prova liberatoria. Tuttavia, questa prospettiva non è condivisa unanimemente, poiché non tutte le forme di IA possono considerarsi intrinsecamente pericolose e il perimetro di tale categoria rischia di risultare eccessivamente ampio e generico. Un'altra possibile via interpretativa è quella offerta dall'art. 2051 c.c., che disciplina il *“danno cagionato da cosa in custodia”* e presuppone la responsabilità di chi ha la custodia di un bene da cui derivi un danno. In questo contesto, la macchina intelligente potrebbe essere assimilata a una *“cosa”*, rendendo responsabile l'utilizzatore o il proprietario per i danni causati. Tuttavia, anche questa soluzione presenta criticità: a differenza di una semplice cosa, che non modifica autonomamente il proprio comportamento, un sistema di intelligenza artificiale è in grado di evolversi nel tempo. Di conseguenza, l'imputazione oggettiva prevista

dall'articolo potrebbe non risultare adeguata a coprire pienamente i rischi emergenti legati a tali tecnologie.

Quindi, come mostrato, l'applicazione di queste norme nell'ambito dell'intelligenza artificiale pone problemi significativi. I sistemi di IA, specialmente quelli autonomi e auto-apprendenti, non sono semplici strumenti meccanici: essi interagiscono con l'ambiente, prendono decisioni, apprendono dai dati e modificano il proprio comportamento nel tempo. Ciò rende problematico l'inquadramento giuridico dei danni da IA, in particolare nei casi in cui non sia possibile individuare con chiarezza un autore umano del danno o una violazione specifica di regole di condotta.

Uno degli aspetti più critici è proprio la ricostruzione del nesso causale, quando l'algoritmo compie scelte automatizzate di cui è complesso ricostruire il percorso logico, non direttamente prevedibili dai programmatori o dagli utilizzatori. Inoltre, si pone la questione della colpa: chi è colpevole in caso di errore algoritmico? Il produttore, lo sviluppatore, l'utente, oppure nessuno in senso tradizionale?

1.3 La disciplina della responsabilità da prodotto difettoso

Accanto alla tradizionale disciplina civilistica italiana in materia di responsabilità extracontrattuale, si affianca una regolamentazione di matrice europea che ha introdotto una forma di responsabilità oggettiva a capo al produttore. Si tratta della Direttiva 85/374/CEE del Consiglio dell'Unione Europea del 25 luglio 1985, recepita in Italia con il d.P.R. n. 224 del 24 maggio 1988, poi confluito nel Codice del Consumo (artt. 114-127). La norma sancisce la responsabilità del produttore per i danni causati da un difetto del prodotto, indipendentemente dalla prova della colpa e pone in capo al danneggiato soltanto l'onere di dimostrare il nesso causale tra difetto e danno.

L'articolo 117 del Codice del Consumo definisce un prodotto difettoso come quello che *“non offre la sicurezza che ci si può legittimamente attendere tenuto conto di tutte le circostanze”*, tra cui la presentazione del prodotto, l'uso al quale può essere ragionevolmente destinato e il momento della sua messa in circolazione. Il legislatore ha quindi adottato una logica anticipatoria e protettiva, fondata su un criterio di ragionevole aspettativa di sicurezza del consumatore, piuttosto che sull'effettiva colpa del produttore.

Tuttavia, questa architettura normativa, pensata per prodotti industriali tradizionali, fisici, statici e tendenzialmente prevedibili, mostra gravi limiti di applicabilità rispetto ai sistemi di intelligenza artificiale. Tali tecnologie, specialmente quelle basate su apprendimento automatico (machine learning) e apprendimento profondo (deep learning), presentano una complessità intrinseca che rende inadeguata la tradizionale nozione di “*prodotto difettoso*”. Le principali criticità che possono essere rilevate sono:

a) La definizione di “difetto” è inadeguata per sistemi dinamici

La nozione di “difetto”, intesa come mancato rispetto delle legittime aspettative di sicurezza da parte dell’utente medio, è fortemente ancorata a un contesto in cui il comportamento del prodotto è determinabile ex ante. Al contrario, i sistemi di IA, specie quelli che operano in ambienti non supervisionati o evolutivi, non mantengono caratteristiche fisse, ma sono soggetti a mutamenti progressivi nel tempo, derivanti dall’interazione con nuovi dati e contesti. Ad esempio, un algoritmo di guida autonoma potrebbe adottare strategie diverse per compiere le stesse operazioni a distanza di settimane, rendendo difficile attribuire una responsabilità fondata su un “difetto originario”.

b) La natura processuale e non statica dei sistemi di IA

I sistemi di intelligenza artificiale non sono più “beni” nel senso classico del termine, bensì processi cognitivi artificiali, capaci di apprendere e di modificare autonomamente le proprie regole operative. Questo comporta che il momento della “messa in circolazione” del prodotto, momento giuridicamente rilevante secondo l’art. 114 del Codice del Consumo, non coincida più necessariamente con il momento in cui il sistema manifesta comportamenti potenzialmente dannosi. In altre parole, l’IA può diventare “pericolosa” non al momento dell’immissione sul mercato, ma molto tempo dopo, a causa di aggiornamenti, auto-apprendimento o utilizzo in contesti non previsti.

c) La difficoltà nell’individuare il “produttore” responsabile

Un ulteriore aspetto problematico riguarda l’identificazione del soggetto responsabile. In molti casi, la produzione di un sistema di IA coinvolge una molteplicità di attori: sviluppatori del software, produttori dell’hardware, fornitori di dataset, integratori di sistema, operatori che gestiscono la manutenzione e persino gli utenti finali che personalizzano l’algoritmo con input specifici. Tale

frammentazione rende estremamente arduo risalire a un singolo “produttore” in senso giuridico, soprattutto in contesti di sviluppo collaborativo o open source. In questi casi, l’attribuzione della responsabilità rischia di diventare un’operazione arbitraria o di ricadere su soggetti che non hanno effettivo controllo sui processi decisionali dell’IA.

1.4 Il contributo del GDPR

Il Regolamento (UE) 2016/679, noto come GDPR, rappresenta il primo significativo tentativo di disciplina europea che, pur non essendo concepito specificamente per l’intelligenza artificiale, si rivela cruciale nell’ambito delle decisioni automatizzate basate su dati personali. Entrato in vigore nel 2018, il GDPR incide profondamente sul rapporto tra individuo e tecnologia, ponendo al centro la tutela della persona rispetto ai processi decisionali algoritmici.

Particolarmente emblematico è l’articolo 22, il quale sancisce il diritto dell’interessato a non essere soggetto a una decisione derivante unicamente da un sistema automatizzato, cioè senza alcun intervento umano diretto, a meno che non ricorrano specifiche condizioni stabilite dalla legge. Tale disposizione trova applicazione soprattutto in contesti sensibili quali il credito, la selezione del personale o la sanità, dove le scelte operate da algoritmi possono determinare conseguenze di rilievo giuridico e sociale.

Inoltre, il GDPR promuove principi imprescindibili quali la trasparenza, il diritto all’informazione e alla spiegazione delle decisioni automatizzate. Questi obblighi impongono ai titolari del trattamento di fornire una chiara e comprensibile esposizione della logica sottesa ai sistemi algoritmici e delle finalità perseguite, al fine di garantire una consapevolezza reale e sostanziale del coinvolgimento dei dati personali nei processi decisionali. Tuttavia, anche in questo caso si evidenziano limiti strutturali:

- Le norme sono orientate alla prevenzione del trattamento illecito ma non disciplinano i danni causati dall’IA.
- Il GDPR non fornisce strumenti per l’attribuzione della responsabilità in caso di malfunzionamenti o bias discriminatori.
- Mancano sanzioni civili specifiche in caso di danni derivanti da algoritmi che operano in violazione della normativa sulla protezione dei dati.

In tal senso, il GDPR costituisce una base importante ma non sufficiente per affrontare in modo sistemico il problema della responsabilità civile da IA, infatti come sottolinea V. Cuffaro *“la disciplina sulla protezione dei dati personali è fondamentale, ma non può da sola governare i rischi connessi all’uso dell’intelligenza artificiale nei processi decisionali”*.

2) Casi di Studio: Il caso Uber e Tesla.

2.1 Analisi del caso Uber

Il 18 marzo 2018, nella città di Tempe, in Arizona, si è verificato il primo incidente mortale che ha coinvolto un veicolo a guida autonoma. Protagonista dell’accaduto è stato un SUV Volvo XC90 gestito dalla compagnia Uber Technologies Inc. per operare in modalità di guida autonoma all’interno del programma sperimentale dell’azienda.

Il veicolo stava circolando in modalità automatica con un’operatrice di sicurezza a bordo, incaricata di intervenire in caso di malfunzionamento, quando il veicolo ha investito la quarantannenenne Elaine Herzberg, che stava attraversando una strada a quattro corsie fuori dalle strisce pedonali, spingendo a mano una bicicletta.

Le indagini successive hanno rivelato diverse criticità:

- Il sistema di guida autonoma aveva effettivamente rilevato la presenza della persona, ma aveva più volte riclassificato l’oggetto in modo errato prima come veicolo, poi come oggetto statico, poi come ciclista, generando confusione nell’algoritmo predittivo del comportamento.
- L’algoritmo era stato progettato per non attivare automaticamente la frenata d’emergenza durante la guida autonoma, al fine di ridurre falsi positivi rimandando questa decisione all’intervento umano.
- L’operatrice a bordo, come evidenziato da riprese video, non era attenta al traffico e non è intervenuta tempestivamente.

Questi fattori, combinati, hanno condotto all’investimento mortale. Il caso ha avuto risonanza globale, ponendo per la prima volta, in termini concreti e drammatici, la questione della responsabilità civile e penale nell’impiego di sistemi di intelligenza artificiale in ambito veicolare.

2.2 Implicazioni giuridiche e responsabilità coinvolte nel caso Uber

Il caso Uber ha sollevato una pluralità di interrogativi giuridici, tuttora oggetto di dibattito, sia nella dottrina europea che nella prassi regolatoria:

- a) **Responsabilità penale e personale dell'operatrice di sicurezza:**
Nel 2020, l'operatrice del caso Uber è stata incriminata per omicidio colposo. Tuttavia, il procedimento ha messo in luce l'ambiguità del suo ruolo: pur essendo incaricata della supervisione, operava all'interno di un sistema pensato per ridurre al minimo l'intervento umano, in evidente tensione con i doveri di vigilanza attiva. Il caso ha evidenziato i rischi di "overtrust" nella tecnologia, ovvero l'eccessiva fiducia nei sistemi automatizzati, che può ridurre la prontezza cognitiva degli operatori.
- b) **Responsabilità civile del produttore del software e del costruttore del veicolo:**
Il software di guida autonoma era sviluppato da Uber mentre la piattaforma veicolare era prodotta da Volvo. La combinazione tra software e hardware solleva interrogativi su una possibile responsabilità plurima o concorrente: da un lato il difetto dell'algoritmo, incapace di classificare correttamente l'ostacolo e attivare misure correttive, dall'altro le caratteristiche strutturali del veicolo.
Secondo il diritto europeo vigente al momento del fatto, la Direttiva 85/374/CEE sulla responsabilità per danno da prodotti difettosi non era pienamente applicabile, mancando criteri certi per qualificare i sistemi di IA come "prodotti" ai sensi della normativa. Questo vuoto ha spinto la Commissione Europea a proporre nuove forme di responsabilità ex lege, come previsto nella proposta di Direttiva COM (2022)496.
- c) **Profili di governance e responsabilità delle autorità pubbliche:**
Non meno rilevante è l'aspetto regolatorio: l'incidente è avvenuto in uno Stato, l'Arizona, che aveva deliberatamente adottato una politica di de-regolamentazione per attrarre imprese tecnologiche, riducendo al minimo i requisiti di autorizzazione e supervisione. Ciò ha generato un contesto normativo lacunoso, che ha permesso lo svolgimento di test su strada senza garanzie adeguate alla sicurezza pubblica.

2.3 Analisi del caso Tesla

Il 7 maggio 2016, nei pressi di Williston, in Florida, una Tesla Model S con sistema Autopilot attivato e condotta da Joshua Brown, ex ufficiale della Marina statunitense e appassionato sostenitore della tecnologia Tesla, si schiantò contro un rimorchio articolato che stava attraversando l'autostrada. Il sistema non riconobbe l'ostacolo e proseguì la marcia a velocità sostenuta, circa 120 km/h, passando sotto al rimorchio e causando la morte immediata del conducente. Secondo la ricostruzione della National Transportation Safety Board (NTSB), il camion stava effettuando una svolta a sinistra, attraversando la corsia sulla quale sopraggiungeva la Tesla. A causa delle condizioni di luce e della colorazione bianca del rimorchio, il sistema Autopilot non fu in grado di distinguere correttamente l'ostacolo. Né il radar, né la videocamera, interpretarono il rimorchio come un pericolo imminente, e di conseguenza non venne attivata alcuna frenata automatica.

La stessa indagine accertò che il conducente non aveva le mani sul volante da diversi minuti, nonostante Tesla raccomandasse espressamente di mantenerle sempre sullo sterzo, anche durante l'uso dell'Autopilot. Inoltre, non risultarono tentativi di frenata manuale o sterzata correttiva, segno di una fiducia eccessiva nel sistema.

2.4 Implicazioni giuridiche e responsabilità coinvolte nel caso Tesla

L'incidente solleva interrogativi cruciali sul piano della responsabilità civile in relazione a tecnologie di intelligenza artificiale applicata alla mobilità. Le questioni principali riguardano:

a) Il grado di autonomia del sistema:

Il sistema Autopilot, al momento dell'incidente, era classificato come livello due secondo lo standard SAE (Society of Automotive Engineers). Questo significa che, pur essendo in grado di controllare direzione e velocità, il sistema richiede la costante supervisione del conducente umano, che deve essere sempre pronto a intervenire. Tuttavia, il nome stesso del sistema "Autopilot" e l'esperienza utente fornita, possono aver indotto in errore, generando un livello eccessivo di affidamento e una percezione falsata delle capacità del veicolo.

b) Il difetto di progettazione o di informazione:

Se si considera il veicolo come un prodotto ai sensi del Codice del Consumo italiano (D.Lgs. 206/2005, artt. 114-121), sorge il dubbio se vi sia stato un difetto di progettazione nel mancato rilevamento dell'ostacolo o un difetto informativo nella comunicazione del livello effettivo di autonomia. In entrambi i casi, il produttore potrebbe essere ritenuto responsabile per danno da prodotto difettoso, qualora il bene non garantisca la sicurezza che ci si può legittimamente attendere in base all'uso prevedibile.

c) Il ruolo del conducente e la colpa concorrente:

Un ulteriore aspetto critico riguarda il comportamento del conducente. La mancata supervisione attiva e il mancato intervento tempestivo possono configurare una colpa concorrente, anche se mitigata dal livello di fiducia che il sistema stesso sembrava indurre.

d) L'attività pericolosa e l'art. 2050 c.c. :

Secondo la dottrina civilistica italiana, l'uso di sistemi automatizzati a guida assistita potrebbe rientrare nell'ambito delle attività pericolose, ai sensi dell'articolo 2050 del Codice civile, per la loro complessità tecnica e il rischio intrinseco. Il produttore o fornitore del sistema, in tal caso Tesla, sarebbe tenuto al risarcimento del danno salvo prova di aver adottato tutte le misure idonee ad evitarlo.

3. Prospettive future e proposte di riforma

3.1 Necessità di un quadro normativo armonizzato

Il caso Uber ed il caso Tesla, insieme ad altri episodi simili emersi nel panorama globale, hanno evidenziato con chiarezza l'urgenza di sviluppare un quadro giuridico armonizzato e coerente a livello sovranazionale per regolare le conseguenze giuridiche delle attività svolte da sistemi di intelligenza artificiale. L'attuale frammentazione normativa, sia sul piano interno agli Stati membri sia nel confronto tra ordinamenti nazionali e norme dell'Unione Europea, rischia infatti di creare incertezza giuridica e disparità nella tutela dei diritti.

A fronte di ciò, appare necessario delineare un corpus normativo armonizzato che:

- definisca chiaramente i ruoli e le responsabilità dei diversi attori coinvolti (sviluppatori, produttori, fornitori di servizi, utilizzatori);
- stabilisca standard tecnici minimi per l'affidabilità e la sicurezza dei sistemi IA;
- preveda strumenti processuali e probatori idonei a superare le asimmetrie informative tra vittime e detentori della tecnologia;
- garantisca una tutela uniforme dei diritti fondamentali e un accesso effettivo alla giustizia.

3.2 Proposte per una regolamentazione efficace della responsabilità da IA

Alla luce delle insufficienze del quadro normativo esistente, in particolare per quanto riguarda la responsabilità civile in ambito algoritmico, l'Unione Europea ha avviato un ambizioso progetto di riforma legislativa volto a costruire un ecosistema giuridico coerente, capace di garantire sia la tutela dei diritti fondamentali sia lo sviluppo responsabile dell'innovazione tecnologica. In questo contesto si collocano due strumenti normativi complementari: il Regolamento sull'Intelligenza Artificiale (AI Act) e la Proposta di Direttiva sulla responsabilità civile da IA.

a) Il Regolamento sull'Intelligenza Artificiale (AI Act)

Il Regolamento (UE) sull'Intelligenza Artificiale, approvato in via definitiva nel 2024, costituisce il primo tentativo, a livello mondiale, di disciplinare in modo organico l'intelligenza artificiale con un approccio basato sulla gestione del rischio. L'AI Act non si limita a dettare principi etici o raccomandazioni tecniche ma stabilisce obblighi giuridicamente vincolanti, classificando i sistemi di IA in quattro categorie di rischio:

1. Rischio inaccettabile: Sistemi vietati perché contrari ai valori fondamentali dell'UE (es. social scoring da parte di autorità pubbliche, manipolazione cognitivo-comportamentale).
2. Rischio alto: Sistemi sottoposti a rigidi obblighi di conformità, ad esempio, la IA impiegata nella sanità, giustizia, trasporti, istruzione, finanza, selezione del personale. Sono richiesti: valutazioni di impatto, tracciabilità, gestione del rischio, documentazione tecnica e supervisione umana.

3. Rischio limitato: Obblighi informativi per sistemi che interagiscono con gli utenti (es. chatbot, sistemi di raccomandazione).
4. Rischio minimo o nullo: Libera circolazione senza obblighi specifici (es. videogiochi, filtri spam, IA generativa priva di impatti significativi).

Sebbene l'AI Act non disciplini direttamente la responsabilità civile, esso svolge un ruolo determinante nella definizione degli standard di diligenza e delle buone pratiche tecniche. In un eventuale giudizio per risarcimento danni, la violazione degli obblighi di conformità previsti dal Regolamento potrà costituire un indizio di colpa o negligenza, facilitando l'attribuzione della responsabilità al fornitore, al produttore o all'operatore del sistema.

In tal senso, l'AI Act si configura come uno strumento ex ante, volto a prevenire gli effetti dannosi dell'IA ma con importanti implicazioni ex post, nel momento in cui occorra valutare se vi sia stata una condotta colposa o imprudente da parte degli attori coinvolti nello sviluppo o nell'uso del sistema.

b) La proposta di Direttiva sulla responsabilità civile da IA (COM/2022/496)

Accanto al regolamento tecnico, l'UE ha presentato nel 2022 una Proposta di Direttiva sulla responsabilità civile extracontrattuale per l'IA, con l'obiettivo di rendere più agevole l'accesso al risarcimento per le vittime di danni causati da sistemi intelligenti. Si tratta di un intervento mirato a colmare i vuoti delle normative nazionali, che spesso risultano inadeguate nel contesto della complessità tecnologica e decisionale dei sistemi automatizzati.

Tra le principali novità introdotte dalla proposta vi sono:

a) Introduzione di regimi di responsabilità oggettiva per le IA ad alto rischio:

Analogamente a quanto avviene in ambiti come la responsabilità per danni nucleari o da aeromobili, si è proposto di introdurre una forma di responsabilità oggettiva, cioè, prescindente dalla colpa per i danni causati da sistemi di IA, classificati come ad alto rischio. Questo approccio, già in parte suggerito dalla proposta di Direttiva europea sulla responsabilità da IA, garantirebbe una tutela efficace per le vittime, compensando le difficoltà probatorie legate alla complessità algoritmica.

b) Presunzioni legali di nesso causale e onere della prova attenuato per le vittime:

Per rendere effettivo il diritto al risarcimento, è essenziale introdurre presunzioni di causalità, a carico

dell'operatore o del produttore, qualora si verifichi un danno in contesti di utilizzo dell'IA. Tale meccanismo, già previsto nella proposta COM/2022/496, mira a riequilibrare l'asimmetria informativa che penalizza la vittima, spesso priva delle competenze tecniche e dell'accesso alla documentazione necessaria per dimostrare il malfunzionamento o l'errore del sistema.

c) Obblighi di tracciabilità e conservazione dei dati decisionali dei sistemi IA:

Un'efficace regolamentazione della responsabilità richiede la possibilità di ricostruire ex post la sequenza decisionale che ha condotto al danno. A tal fine, è necessario prevedere l'obbligo per sviluppatori e fornitori di sistemi IA di mantenere log e audit trail accessibili e comprensibili, anche ai fini della prova giudiziale. La trasparenza algoritmica non deve essere intesa solo come requisito etico ma come presupposto tecnico-giuridico essenziale per l'attribuzione della responsabilità.

d) Introduzione di fondi di compensazione o assicurazioni obbligatorie:

Per i casi in cui non sia possibile individuare con certezza il soggetto responsabile o in cui l'autore del danno non sia solvibile, è stato ipotizzato l'istituto di fondi di garanzia pubblici o privati per l'indennizzo delle vittime. In alternativa o in aggiunta, si può prevedere l'introduzione di assicurazioni obbligatorie per l'uso di sistemi IA in settori ad alto rischio, simili a quelle già esistenti per la circolazione stradale o le attività mediche.

e) Adattamento delle categorie giuridiche alla specificità algoritmica:

Nel panorama giuridico e dottrinale europeo, si è diffusa l'idea di "responsabilità algoritmica", intesa come l'insieme di strumenti normativi e concettuali volti a garantire l'attribuzione della responsabilità in presenza di danni causati da sistemi autonomi o semi-autonomi. Il dibattito si articola intorno a diverse proposte teoriche, volte a colmare il vuoto lasciato dalle categorie tradizionali della responsabilità civile.

Una prima ipotesi, oggi largamente abbandonata, riguardava l'introduzione di una personalità giuridica elettronica per i sistemi di IA particolarmente autonomi, in modo da renderli direttamente responsabili. Tale proposta, contenuta inizialmente nella Risoluzione del Parlamento Europeo del 2017, ha suscitato ampie critiche per la sua astrattezza e incompatibilità con i principi fondamentali del diritto civile, che legano la responsabilità alla volontà e alla capacità di discernimento del soggetto agente.

Più realistico è invece il ricorso a modelli di responsabilità oggettiva, incentrati sul rischio tecnologico: chi trae beneficio dall'uso di un sistema di IA ne sopporta anche i rischi, indipendentemente dalla colpa. Questo schema, ispirato al diritto dei trasporti o alla responsabilità per danni ambientali, permetterebbe di garantire pronta tutela alle vittime, anche in assenza di prove certe sulla dinamica del danno.

Ad oggi, l'orientamento prevalente a livello europeo è quello di evitare scelte radicali, preferendo una via prudente ma ambiziosa, basata su tre pilastri:

1. Rafforzamento della responsabilità umana. Si insiste sulla tracciabilità delle decisioni, sulla supervisione umana e sulla responsabilizzazione degli sviluppatori, operatori e fornitori.
2. Standard tecnici e normativi vincolanti. Si introducono regole ex ante (come l'AI Act) che definiscono ciò che è accettabile e ciò che non lo è nell'uso dell'IA.
3. Tutela giudiziaria ed equità probatoria. Con la proposta di Direttiva sulla responsabilità, si cerca di rendere concretamente accessibile il diritto al risarcimento, pur senza introdurre regimi speciali di responsabilità.

In questa prospettiva, la responsabilità algoritmica non implica la creazione di nuove entità soggettive ma la ridefinizione dei confini della responsabilità umana nell'era dell'intelligenza artificiale, favorendo una cultura della compliance e dell'etica tecnologica, in cui l'uso dei sistemi intelligenti sia coniugato con la tutela dei diritti fondamentali. Solo attraverso un approccio consapevole, equilibrato e orientato ai valori sarà possibile trasformare l'innovazione in uno strumento al servizio della persona piuttosto che in una minaccia alla sua dignità, ai suoi valori e diritti. In questo delicato contesto storico, il diritto è chiamato quindi, non solo a regolamentare, ma a guidare il cambiamento, tracciando la rotta verso un futuro in cui l'intelligenza artificiale e la responsabilità civile possano coesistere armonicamente. Una sfida complessa, certo, ma necessaria.

4) BIBLIOGRAFIA

- Codice civile, artt. 2043, 2050, 2051, 1218.
- Codice del Consumo (D.lgs. 206/2005), artt. 114–127.
- d.P.R. 24 maggio 1988, n. 224 – Attuazione della Direttiva 85/374/CEE.
- Direttiva 85/374/CEE del Consiglio, 25 luglio 1985 – Responsabilità per danno da prodotti difettosi.
- Regolamento (UE) 2016/679 (GDPR) – Artt. 5, 13, 14, 22, 82.
- Proposta di Regolamento del Parlamento europeo e del Consiglio sull'intelligenza artificiale (AI Act), COM(2021) 206 final.
- Proposta di Direttiva del Parlamento europeo e del Consiglio sulla responsabilità civile da IA, COM(2022) 496 final.
- Carta dei diritti fondamentali dell'Unione Europea, art. 8 (protezione dei dati personali).
- U. Ruffolo, *Diritto dell'intelligenza artificiale. Principi, responsabilità, tutela dei dati e governance algoritmica*, Giappichelli, 2021.
- G. Comandé, *La responsabilità da algoritmo tra regole, principi e processi*, in “Riv. Dir. Civ.”, 2020.
- V. Cuffaro, *Responsabilità e intelligenza artificiale: riflessioni in attesa della normativa europea*, in “Danno e Responsabilità”, 2022.
- A. Mantelero, *AI and responsibility: A European legal perspective*, Computer Law & Security Review, 2020.
- Commissione Europea, *Libro Bianco sull'Intelligenza Artificiale – Un approccio europeo all'eccellenza e alla fiducia*, COM(2020) 65 final.

- National Transportation Safety Board. (2017). *Collision between a car operating with automated vehicle control systems and a tractor-semitrailer truck near Williston, Florida, May 7, 2016* (Highway Accident Report No. NTSB/HAR-17/02). Washington, D.C.: U.S. Government Publishing Office.
- National Transportation Safety Board. (2019). *Collision between vehicle controlled by developmental automated driving system and pedestrian, Tempe, Arizona, March 18, 2018* (Highway Accident Report No. NTSB/HAR-19/03). Washington, D.C.: U.S. Government Publishing Office.