



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 +

Sistemi di rischio nell'IA Act

Damiano Pace

0365207

Anno accademico 2024/2025

Indice

1. Il contesto normativo dell'AI Act

- 1.1 Origine e finalità del Regolamento europeo sull'IA
- 1.2 Principi generali dell'AI Act
- 1.3 Principi di proporzionalità, trasparenza e di precauzione tecnologica
- 1.4 Ambito di applicazione

2. La classificazione dei sistemi di rischio nell'AI Act

- 2.1 Panoramica dei livelli di rischio
- 2.2 Sistemi a rischio inaccettabile
- 2.3 Sistemi a rischio elevato
- 2.4 Sistemi a rischio limitato
- 2.5 Sistemi a rischio minimo o nullo

3. Sistemi ad alto rischio: requisiti e obblighi

- 3.1 Requisiti per i fornitori di IA ad alto rischio
- 3.2 Valutazione della conformità

4. Criticità e dibattito attuale

- 4.1 Impatto sui diritti fondamentali
- 4.2 Rischi di overregulation
- 4.3 Posizione delle grandi piattaforme digitali
- 4.4 L'approccio europeo vs altri modelli

5. Conclusioni

- 5.1 Sintesi dei principali punti emersi
- 5.2 Riflessioni sul futuro dell'IA e del diritto

Abstract

Questa tesina analizza il modo in cui l'Unione Europea ha deciso di regolamentare l'intelligenza artificiale attraverso l'AI Act, con particolare attenzione al sistema di classificazione del rischio.

L'intelligenza artificiale è ormai parte integrante della nostra vita quotidiana e sta rivoluzionando settori chiave come la medicina, il lavoro, la sicurezza e la giustizia.

Tuttavia, accanto alle opportunità, emergono anche rischi e interrogativi di natura etica e giuridica.

Al centro dell'analisi si trova il modello a quattro livelli di rischio proposto dall'AI Act (inaccettabile, alto, limitato e minimo), che stabilisce obblighi diversi a seconda della pericolosità dei sistemi.

La tesina affronta inoltre due aspetti oggi molto discussi: il principio di precauzione tecnologica, necessario per tutelare i diritti fondamentali, e il rischio di overregulation, che potrebbe ostacolare l'innovazione.

L'obiettivo è comprendere come l'Europa stia cercando di bilanciare tutela dei cittadini e progresso tecnologico, costruendo un modello di regolazione dell'intelligenza artificiale che sia efficace, sostenibile e orientato ai valori democratici.

1 Il contesto normativo dell'AI Act

1.1 Origine e finalità del Regolamento europeo sull'IA

Prima di introdurre i concetti dell'AI Act, è necessario definire cosa si intende di preciso per intelligenza artificiale al fine di comprendere a pieno il seguito della trattazione.

In base al Regolamento (UE) 2024/1689 del Parlamento Europeo e del Consiglio, noto anche come AI Act, l'intelligenza artificiale (IA) è un insieme di tecnologie in rapida evoluzione, in grado di generare benefici significativi a livello economico, sociale e ambientale.

Grazie alla capacità di migliorare le previsioni, ottimizzare le operazioni e le risorse, e personalizzare i servizi digitali, l'IA rappresenta uno strumento strategico per l'innovazione sia nelle imprese sia nel settore pubblico.

L'adozione di sistemi di IA può contribuire non solo a rendere più efficienti i processi produttivi e decisionali, ma anche a migliorare la qualità della vita delle persone, favorendo uno sviluppo più sostenibile e inclusivo.

L'AI Act, quindi, nasce dalla crescente consapevolezza che l'intelligenza artificiale, pur offrendo grandi opportunità, può anche rappresentare un rischio serio per i diritti fondamentali, la sicurezza e la fiducia dei cittadini.

Già nel 2018 la Commissione Europea aveva iniziato a lavorare su una strategia per rendere l'Europa un polo di riferimento per un'IA "affidabile", e nel 2021 è arrivata la prima proposta ufficiale di regolamento.

Dopo un lungo processo di negoziazione, l'AI Act è stato approvato nel 2024, diventando il primo regolamento completo sull'intelligenza artificiale a livello globale.

L'obiettivo principale dell'AI Act prevede un duplice aspetto:

- Garantire la sicurezza e il rispetto dei diritti fondamentali delle persone in tutti i casi in cui viene utilizzata l'intelligenza artificiale;
- Favorire lo sviluppo e l'adozione dell'IA in Europa, creando regole comuni per il mercato unico europeo.

Quindi, l'Unione Europea intende guidare lo sviluppo dell'intelligenza artificiale rimanendo competitiva sul piano tecnologico, ma senza rinunciare ai propri valori: etica, responsabilità e tutela dei diritti delle persone.

1.2 Principi generali dell'AI Act

Anche se l'AI Act è un regolamento tecnico e dettagliato, si basa su alcuni principi fondamentali, che guidano tutto il testo:

- Approccio basato sul rischio: non tutte le IA sono uguali. Il regolamento distingue tra diversi livelli di rischio e applica regole più severe dove necessario.
- Trasparenza e tracciabilità: i sistemi di IA devono essere spiegabili, comprensibili e ispezionabili. Le persone devono sapere quando stanno interagendo con un'IA.
- Controllo umano: i sistemi di IA non devono sostituire completamente il giudizio umano, soprattutto in ambiti delicati come la salute, la giustizia o la sicurezza.
- Protezione dei diritti fondamentali: il regolamento mette al centro valori come la privacy, la non discriminazione, la sicurezza e la dignità umana.
- Sviluppo sostenibile e innovazione responsabile: l'AI Act punta a regolamentare senza soffocare il progresso tecnologico, in particolare per le piccole e medie imprese.

Questi principi rappresentano la base su cui l'Unione Europea intende costruire una fiducia collettiva nell'uso dell'IA.

1.3 Principi di proporzionalità, trasparenza e precauzione tecnologica

I principi di proporzionalità, trasparenza e precauzione tecnologica rappresentano pilastri fondamentali del diritto dell'Unione Europea e sono presenti trasversalmente in molteplici ambiti normativi, incluso l'AI Act.

Per tale motivo, è utile richiamare i loro enunciati e contestualizzarli in relazione all'AI Act, dove assumono forme specifiche e operative.

Il principio di proporzionalità, richiamato dall'articolo 5, paragrafo 4, del Trattato sull'Unione Europea (TUE), stabilisce che ogni azione dell'Unione non deve andare oltre quanto necessario per il raggiungimento degli obiettivi dei trattati, limitando quindi l'intervento normativo affinché sia adeguato e non eccessivo rispetto agli scopi perseguiti.

Nell'ambito specifico dell'AI Act, questo principio viene recepito e declinato concretamente nell'articolo 4 del regolamento, che assicura che gli obblighi imposti ai fornitori di sistemi di intelligenza artificiale siano proporzionati al livello di rischio associato al sistema in esame.

Il principio di trasparenza, invece, è volto a garantire la chiarezza, l'accessibilità e la comprensibilità delle informazioni rivolte a cittadini, imprese e autorità.

All'interno dell'AI Act, la trasparenza assume un ruolo cruciale: i fornitori devono fornire una documentazione tecnica esaustiva, istruzioni chiare e informazioni sui limiti dei sistemi, in modo da tutelare i diritti degli utenti e facilitare il controllo da parte delle autorità competenti (Regolamento UE 2024/1689, art. 13).

Infine, il principio di precauzione tecnologica è particolarmente rilevante nel contesto delle tecnologie emergenti come l'intelligenza artificiale.

Esso impone di adottare misure preventive e di gestione del rischio in presenza di incertezze scientifiche o potenziali danni gravi, allo scopo di proteggere i diritti fondamentali e la sicurezza delle persone.

Nell'AI Act, questo principio si traduce nell'obbligo per i fornitori di sottoporre i sistemi ad alto rischio a una rigorosa valutazione di conformità e a un monitoraggio continuo anche dopo l'immissione sul mercato.

In sintesi, pur essendo principi generali e trasversali nel diritto dell'Unione Europea, proporzionalità, trasparenza e precauzione tecnologica sono stati esplicitamente integrati e adattati nel quadro normativo dell'AI Act, assumendo un ruolo centrale per garantire un approccio regolatorio equilibrato, affidabile e orientato alla tutela dei diritti fondamentali e della sicurezza dei cittadini.

1.4 Ambito di applicazione

L'AI Act si applica a tutti gli operatori pubblici e privati che sviluppano, immettono sul mercato o utilizzano sistemi di intelligenza artificiale nell'Unione Europea, anche se il fornitore è situato fuori dall'UE.

Questo significa che anche aziende esterne, come quelle americane o asiatiche, devono rispettare le regole europee se vogliono operare nel mercato europeo.

Il regolamento si applica a tutti i settori, tranne quelli legati alla difesa e alla sicurezza nazionale. Rientrano quindi nel suo campo di applicazione sistemi di IA usati, ad esempio:

- nella sanità (diagnosi automatizzate);
- nella pubblica amministrazione (selezione per bandi);
- nel lavoro (sistemi di selezione del personale);
- nella sorveglianza (riconoscimento facciale);
- nella giustizia (software predittivi per decisioni giudiziarie).

Questa ampiezza rende l'AI Act una normativa centrale per il diritto digitale europeo, perché stabilisce le regole di base per l'uso dell'IA in quasi ogni ambito della vita sociale.

2 La classificazione dei sistemi di rischio nell'AI Act

2.1 Panoramica dei livelli di rischio

Nell'AI Act viene adottato un approccio basato sul rischio per regolare l'impiego dei sistemi di IA.

Ma cosa è esattamente un “rischio”? Il rischio viene definito come la combinazione tra la probabilità che si verifichi un danno e la gravità del danno stesso.

Partendo da questa definizione, l'AI Act classifica i sistemi di IA in quattro categorie principali: rischio inaccettabile, rischio elevato, rischio limitato e rischio minimo o nullo.

Vediamo quindi nel dettaglio le varie categorie dei sistemi.

2.2 Sistemi a rischio inaccettabile

Questa prima categoria comprende quei sistemi considerati intrinsecamente pericolosi per i diritti fondamentali delle persone.

Si tratta di tecnologie che, per finalità o modalità d'uso, risultano incompatibili con i valori dell'Unione Europea, come il rispetto della dignità umana, della libertà individuale e della non discriminazione.

L'AI Act stabilisce che i sistemi rientranti in questa categoria devono essere vietati e non possono essere immessi sul mercato o utilizzati nell'Unione Europea.

Un esempio paradigmatico è il cosiddetto social scoring: si tratta di sistemi che attribuiscono un punteggio ai cittadini sulla base dei loro comportamenti, delle loro interazioni sociali o delle loro abitudini di consumo.

Questi punteggi possono poi essere utilizzati per decidere se concedere o meno determinati vantaggi, come prestiti bancari, accesso a servizi pubblici o opportunità lavorative.

È facile immaginare le distorsioni che questo tipo di meccanismo può produrre: discriminazioni sistemiche, esclusione sociale, sorveglianza pervasiva, e un profondo condizionamento delle libertà personali.

La logica del divieto in questi casi non è basata solo sulla possibilità di errore tecnico, ma sulla violazione strutturale della libertà e della privacy, considerate non negoziabili.

2.3 Sistemi a rischio elevato

Al di sotto del rischio inaccettabile si colloca una vasta gamma di sistemi che, pur essendo ammessi, possono comportare danni significativi in caso di malfunzionamento, uso improprio dei dati o degli algoritmi.

Sono definiti ad alto rischio (o rischio elevato) tutti quei sistemi che interagiscono con ambiti particolarmente delicati della vita sociale, come:

- l'istruzione e la formazione professionale (es. IA per la valutazione degli studenti);
- l'occupazione e la gestione del personale (es. strumenti per selezione automatica dei CV);
- l'accesso a servizi essenziali, come sanità e giustizia;
- la gestione della sicurezza pubblica (es. riconoscimento facciale);
- la mobilità e il trasporto (es. guida autonoma);
- il credito e la finanza (es. scoring creditizio).

Per questi sistemi, l'AI Act impone una serie di obblighi molto rigidi.

Tra i principali obblighi e doveri vi sono: la necessità di effettuare valutazioni d'impatto prima dell'immissione sul mercato, l'obbligo di progettare sistemi trasparenti e comprensibili, l'introduzione di misure di gestione del rischio, la supervisione umana e la tracciabilità delle operazioni svolte dal sistema.

Le aziende che sviluppano o utilizzano questi strumenti devono inoltre registrare i sistemi in una banca dati europea accessibile alle autorità competenti.

Un caso emblematico in questa categoria è il riconoscimento facciale in tempo reale nei luoghi pubblici.

Immaginiamo una metropoli in cui sono installate videocamere con IA in stazioni, aeroporti o piazze principali.

Queste telecamere confrontano i volti delle persone in tempo reale con un database di individui ricercati.

Anche se ciò potrebbe aiutare nella prevenzione del crimine, comporta comunque gravi rischi: violazioni della privacy, errori di identificazione, uso sproporzionato delle forze dell'ordine.

Per questo motivo, l'AI Act ne consente l'uso solo in circostanze eccezionali e con autorizzazioni giudiziarie specifiche.

2.4 Sistemi a rischio limitato

I sistemi di IA classificati come a rischio limitato sono quelli che non interferiscono direttamente con i diritti fondamentali, ma che possono comunque creare disagi, fraintendimenti o problemi minori se utilizzati in modo scorretto.

Per questa categoria non sono previsti divieti o requisiti stringenti come nei casi precedenti, ma l'AI Act impone comunque alcuni obblighi di trasparenza.

In pratica, chi utilizza un sistema di questo tipo deve assicurarsi che l'utente sia consapevole di stare interagendo con un'intelligenza artificiale; è un principio che può sembrare banale, ma che è cruciale per evitare malintesi e per garantire un uso informato e responsabile.

Esempi tipici sono i chatbot generici utilizzati per fornire assistenza online.

Pensiamo a un sito e-commerce che implementa un assistente virtuale per rispondere a domande frequenti su prodotti, spedizioni o resi.

Questo sistema non accede a dati sensibili e non prende decisioni critiche, ma è comunque necessario avvertire l'utente che sta dialogando con una macchina.

2.5 Sistemi a rischio minimo o nullo

Infine, il livello più basso della scala riguarda i sistemi con un impatto trascurabile sui diritti, la sicurezza o la dignità delle persone e per questo motivo non sono soggetti a obblighi specifici, se non quelli già previsti per qualunque prodotto o software (come sicurezza informatica, protezione dei dati, ecc.).

In questa categoria rientrano applicazioni molto comuni e semplici: calcolatrici intelligenti, app per organizzare la lista della spesa, giochi digitali, o strumenti didattici interattivi per bambini.

Questi sistemi non elaborano dati sensibili, non prendono decisioni autonome e non incidono sulla vita delle persone in modo rilevante.

Proprio perché il loro impatto è minimo, l'AI Act adotta un approccio "soft", lasciando piena libertà d'uso e sviluppo, in linea con il principio di proporzionalità (BicoccaNews, 2024).

3 Sistemi ad alto rischio: requisiti e obblighi

3.1 Requisiti per i fornitori di IA ad alto rischio

I sistemi di intelligenza artificiale ad alto rischio rappresentano le tecnologie soggette alla maggiore regolamentazione all'interno del quadro normativo europeo, a causa dei potenziali impatti significativi sui diritti fondamentali, sulla sicurezza e sull'interesse pubblico.

Per questo motivo, il Regolamento AI Act impone ai fornitori di tali sistemi una serie di requisiti stringenti volti a garantire affidabilità, trasparenza e tracciabilità.

1) Gestione del rischio:

È l'elemento centrale e consiste nell'obbligo di adottare un sistema di gestione del rischio continuo e dinamico, valido per l'intero ciclo di vita del sistema.

Questo sistema deve identificare e valutare i rischi legati all'uso dell'IA, adottare misure di mitigazione efficaci e monitorare costantemente le prestazioni, anche dopo l'immissione sul mercato.

L'approccio alla gestione del rischio deve tener conto non solo delle dimensioni tecniche, ma anche delle implicazioni etiche e sociali.

2) Qualità e governance dei dati:

La qualità dei dati rappresenta un prerequisito essenziale per l'affidabilità dei sistemi IA.

I dati utilizzati per addestramento, validazione e test devono essere rappresentativi, accurati, privi di distorsioni e gestiti con processi documentati e tracciabili.

La governance dei dati deve inoltre garantire la conformità al GDPR, prevedendo strumenti per rilevare e correggere eventuali lacune, con un'attenzione particolare alla protezione dei dati sensibili.

3) Documentazione tecnica e trasparenza:

Ogni sistema di IA ad alto rischio deve essere accompagnato da una documentazione tecnica completa, che consenta alle autorità di vigilanza e agli utenti di comprendere il funzionamento del sistema, verificarne la conformità e valutarne i rischi.

Tale documentazione include informazioni su fonti dei dati, tecniche di elaborazione, metriche di prestazione e strategie di mitigazione.

Allo stesso tempo, i fornitori devono garantire la trasparenza nei confronti degli utenti, attraverso istruzioni chiare, linee guida per l'interpretazione dei risultati e indicazioni sui limiti del sistema.

4) Logging e supervisione umana:

Il regolamento richiede che i sistemi registrino automaticamente eventi significativi durante l'utilizzo, come avvio, arresto, input/output, modifiche apportate e interazioni critiche.

In parallelo, è obbligatorio prevedere meccanismi di sorveglianza umana, che consentano un intervento tempestivo in caso di anomalie o rischi, fino all'arresto del sistema se necessario.

5) Accuratezza, robustezza e cybersicurezza:

I sistemi devono garantire risultati accurati e affidabili, essere progettati per resistere a guasti tecnici e input errati, e disporre di adeguate misure di protezione contro attacchi informatici e manipolazioni malevoli. (negg Group, 2025)

3.2 Valutazione della conformità

Il Regolamento sull'Intelligenza Artificiale, entrato in vigore il 1° agosto 2024, riprende e sviluppa i principi già enunciati nella proposta iniziale pubblicata dalla Commissione Europea nel 2021, imponendo specifici obblighi per i fornitori di sistemi ad alto rischio (Commissione Europea, Q&A, 2021).

L'adempimento di questi requisiti deve essere dimostrato attraverso una valutazione di conformità, che rappresenta uno degli strumenti centrali previsti dal regolamento per assicurare l'affidabilità dei sistemi IA ad alto rischio.

Tale valutazione, obbligatoria prima dell'immissione sul mercato o della messa in servizio nell'UE, consiste in una procedura strutturata per verificare che il sistema soddisfi tutti i requisiti tecnici e organizzativi imposti.

A seconda della tipologia di sistema, la valutazione può essere effettuata internamente dal fornitore oppure, nei casi più sensibili, come i sistemi biometrici, deve essere condotta da organismi terzi notificati.

La valutazione non è un atto isolato, ma deve essere ripetuta qualora il sistema subisca modifiche sostanziali o cambiamenti nella sua finalità originaria, per garantire che ogni evoluzione resti conforme alla normativa.

Inoltre, i fornitori devono dotarsi di sistemi interni di gestione della qualità e del rischio, in grado di monitorare la conformità anche successivamente all'immissione sul mercato.

Ciò include la raccolta di segnalazioni, l'analisi degli incidenti e la predisposizione di misure correttive, rafforzando la responsabilità del fornitore per tutta la vita operativa del sistema.

Un ulteriore elemento di trasparenza è rappresentato dalla banca dati pubblica dell'UE, in cui devono essere registrati tutti i sistemi ad alto rischio utilizzati da autorità pubbliche o da soggetti che agiscono per loro conto.

Per i sistemi impiegati in ambiti sensibili, come l'ordine pubblico o la gestione migratoria, è prevista invece una sezione non pubblica, accessibile solo dalle autorità competenti.

Le autorità di vigilanza del mercato svolgono un ruolo fondamentale nel garantire il rispetto delle norme, attraverso audit periodici e il monitoraggio post-commercializzazione.

In casi eccezionali, esse possono anche autorizzare temporaneamente l'uso di sistemi non ancora conformi, se ciò risponde a esigenze urgenti e giustificate.

Infine, in presenza di violazioni, le autorità nazionali hanno il diritto di accedere a tutte le informazioni necessarie per condurre indagini e valutare la legalità dell'utilizzo del sistema.

La valutazione di conformità, dunque, non è soltanto un adempimento formale, ma costituisce uno strumento giuridico e tecnico di garanzia, essenziale per promuovere un ecosistema dell'IA europeo affidabile, sicuro e rispettoso dei diritti fondamentali.

4. Implicazioni giuridiche ed etiche

4.1 Impatto sui diritti fondamentali

L'adozione di sistemi di intelligenza artificiale, in particolare quelli classificati ad alto rischio, solleva importanti questioni giuridiche ed etiche che incidono direttamente sui diritti fondamentali sanciti a livello nazionale ed europeo.

Tali diritti comprendono la tutela della privacy, la non discriminazione, la trasparenza e la sicurezza degli individui.

I sistemi di IA possono mettere a rischio la privacy a causa del trattamento massivo e spesso sensibile di dati personali e biometrici, con potenziali violazioni della riservatezza qualora non siano applicate adeguate misure di protezione e conformità al GDPR.

Inoltre, algoritmi con dati o modelli distorti possono generare risultati discriminatori, violando il principio di uguaglianza e perpetuando bias legati a genere, etnia, età o altre caratteristiche protette.

La trasparenza è un altro diritto fondamentale messo alla prova dall'"opacità" di molti algoritmi, che rende difficile per gli utenti comprendere come e perché sono state prese determinate decisioni automatizzate, limitando così la possibilità di contestazione e controllo umano.

Il rischio aumenta quando queste decisioni hanno impatti rilevanti sulla vita delle persone, come nell'ambito lavorativo, giudiziario o sanitario.

Infine, la sicurezza e l'affidabilità dei sistemi IA sono essenziali per evitare malfunzionamenti che potrebbero provocare danni fisici, economici o sociali.

Questo richiede misure di robustezza e protezione da attacchi informatici, nonché la previsione di una sorveglianza umana attiva (European Commission, 2021).

4.2 Rischi di overregulation

Sebbene l'AI Act rappresenti un avanzamento significativo verso una regolamentazione responsabile dell'intelligenza artificiale, diversi osservatori hanno espresso preoccupazioni circa il rischio di overregulation, ovvero un eccesso di vincoli normativi che potrebbe ostacolare l'innovazione tecnologica e penalizzare la competitività dell'Unione Europea.

In particolare, l'imposizione di requisiti stringenti in termini di valutazioni di conformità, gestione del rischio, trasparenza e documentazione tecnica può risultare particolarmente onerosa per le piccole e medie imprese, per le startup e per i centri di ricerca indipendenti, che spesso non dispongono delle risorse economiche e organizzative per adempiere a obblighi tanto complessi.

Il rischio è che le imprese europee più fragili decidano di spostare le attività di ricerca e sviluppo verso mercati con regolamentazioni più flessibili, compromettendo così la sovranità tecnologica dell'UE.

A tal proposito, l'associazione Digital Europe ha evidenziato come l'AI Act, nella sua versione più ampia e restrittiva, possa soffocare la crescita di attori emergenti nel settore dell'IA, dichiarando:

“For Europe to become a global digital powerhouse [...] Let's not regulate them out of existence before they get a chance to scale” (Finextra, 2023).

Analogamente, l'organizzazione AI Chamber, nel commentare una bozza del *General Purpose AI Code of Practice*, ha parlato apertamente di un rischio di eccessiva burocrazia per gli sviluppatori indipendenti e le startup, definendo il documento un *“bureaucratic monster”* e denunciando obblighi “sproporzionati” e potenzialmente paralizzanti per il settore.

Il regolamento stesso prevede alcuni correttivi per bilanciare questi rischi, tra cui l'adozione di un approccio basato sul rischio (art. 4), che differenzia gli obblighi in base alla criticità del sistema, e la possibilità di autorizzazioni straordinarie (art. 55).

Tuttavia, sarà cruciale che le autorità competenti adottino un'applicazione flessibile del regolamento, affinché l'ambiente normativo europeo possa sostenere l'innovazione senza diventare un ostacolo alla crescita.

4.3 Posizione delle grandi piattaforme digitali

Le grandi piattaforme digitali, tra cui Google, Microsoft, Meta e altri leader del settore IA, rappresentano attori centrali nell'ecosistema tecnologico europeo e globale e sono particolarmente coinvolte dall'applicazione dell'AI Act.

Pur riconoscendo l'importanza di un quadro normativo volto a garantire sicurezza, trasparenza e tutela dei diritti fondamentali, queste aziende sottolineano la necessità di un equilibrio tra regolamentazione e innovazione.

In particolare, le grandi aziende sottolineano che l'AI Act deve bilanciare l'esigenza di regole rigorose con la flessibilità necessaria per non ostacolare lo sviluppo e l'adozione di modelli fondamentali di intelligenza artificiale.

Digital Europe (2023) evidenzia che “una regolamentazione eccessiva potrebbe limitare la capacità delle grandi piattaforme di innovare e di competere a livello globale”, e che “le normative devono garantire un approccio proporzionato per non frenare lo sviluppo tecnologico, soprattutto riguardo ai sistemi basati su modelli fondamentali”.

Nonostante le criticità sollevate, molte grandi aziende mostrano disponibilità a collaborare con le istituzioni europee per definire linee guida chiare e proporzionate, volte a costruire un ambiente di fiducia attorno all'IA e a promuovere uno sviluppo tecnologico responsabile e sostenibile (European Commission, 2024).

In sintesi, le grandi piattaforme riconoscono nell'AI Act sia un'opportunità per consolidare un modello europeo di regolamentazione etica dell'intelligenza artificiale, sia una sfida dovuta alla complessità e ai costi della conformità, sottolineando l'importanza di un approccio equilibrato che tuteli i diritti senza penalizzare l'innovazione e la competitività delle aziende leader del settore.

4.4 L'approccio europeo vs altri modelli

L'Unione Europea attraverso l'AI Act, come già ampiamente discusso all'interno di questa tesina, adotta un approccio strutturato e basato sul rischio.

I sistemi di intelligenza artificiale vengono suddivisi in diversi livelli di rischio – da “accettabile” a “inaccettabile” – con vincoli stringenti per quelli ad alto rischio, che includono obblighi di trasparenza, valutazione d'impatto, conformità certificata da enti terzi e divieti per l'utilizzo di sistemi di identificazione biometrica in tempo reale o di social scoring (Parlamento europeo, 2024; Hertie School, 2024).

Questo modello è preventivo e vincolante, con un sistema di autorità centrali (quali l'AI Office e l'AI Board) e prevede sanzioni fino al 7% del fatturato globale dell'operatore (Hertie School, 2024).

Inoltre, l'AI Act ha efficacia extraterritoriale, vincolando anche gli operatori extra-UE qualora i loro sistemi siano destinati al mercato europeo (Hertie School, 2024).

Per quanto riguarda gli Stati Uniti invece, essi privilegiano un quadro normativo decentrato e settoriale, in assenza di una legge federale specifica sull'intelligenza artificiale.

L'approccio statunitense si articola principalmente tramite:

- un Executive Order del 2023 che detta linee guida e obblighi di reporting per le agenzie federali, ma che non ha forza di legge nei confronti dei privati (Hertie School, 2024; White House, 2023);
- l'applicazione di quadri normativi preesistenti, quali le competenze della Federal Trade Commission (FTC) o della Food and Drug Administration (FDA), gestiti caso per caso;
- iniziative statali e volontarie dell'industria, come l'AI Bill of Rights e il NIST AI Risk Management Framework (Science Business, 2024; Hertie School, 2024).

Secondo la Hertie School, UE e USA convergono sull'adozione di un modello risk-based, ma divergono radicalmente nell'implementazione: l'Unione Europea si caratterizza per regole prescrittive e autorità centrali, mentre gli Stati Uniti adottano un ecosistema flessibile, ma potenzialmente frammentato e poco prevedibile per le imprese (Hertie School, 2024).

Analogamente, il think tank Brookings sottolinea come il rischio di disallineamento normativo e la creazione di barriere non tariffarie siano fattori concreti e pericolosi; pertanto, è necessaria un'armonizzazione che passi da definizioni comuni e dalla cooperazione tra le autorità di conformità (Brookings Institution, 2024).

Un ulteriore approfondimento, come evidenziato dall'articolo di The AI Innovator, sottolinea la differenza sostanziale tra i due approcci: l'AI Act è una legge vincolante, applicabile anche a operatori esterni al territorio europeo, mentre l'Executive Order statunitense resta una norma di tipo soft law, senza sanzioni e modificabile facilmente da una nuova amministrazione (Yao, 2024).

Infine, per quanto riguarda un'altra superpotenza come la Cina, il suo modello regolatorio si basa su un controllo statale fortemente centralizzato, mirato alla stabilità sociale e politica.

La Cina ha adottato normative specifiche che impongono obblighi di trasparenza e supervisione sulle raccomandazioni algoritmiche, oltre a regolamentazioni severe sui contenuti generati artificialmente e sulle licenze per sistemi critici di intelligenza artificiale.

Questo approccio privilegia la sicurezza nazionale e il controllo governativo diretto, limitando al contempo gli spazi di innovazione privata indipendente e la tutela delle libertà individuali (Hua & Zhang, 2024).

In conclusione, i modelli di regolazione dell'intelligenza artificiale adottati da Unione Europea, Stati Uniti e Cina riflettono visioni profondamente diverse del rapporto tra tecnologia, diritto e società.

L'approccio dell'UE è di tipo prescrittivo e centralizzato con un modello che punta a prevenire i rischi sistemici, tutelando i diritti fondamentali e garantendo la fiducia dei cittadini nelle tecnologie emergenti.

Gli Stati Uniti, invece, adottano un sistema decentrato e flessibile, basato su norme settoriali, linee guida volontarie e autoregolamentazione dell'industria.

Sebbene ciò favorisca l'innovazione rapida e l'adattabilità, può comportare frammentazione normativa e incertezza giuridica, soprattutto per le imprese operanti a livello internazionale.

La Cina, infine, presenta un modello statale e fortemente centralizzato, in cui l'intelligenza artificiale è regolata in funzione della sicurezza nazionale e del controllo sociale.

5. Conclusioni

5.1 Sintesi dei principali punti emersi

Alla luce di quanto emerso, può essere utile fare un breve riepilogo per tirare le fila del discorso.

L'intelligenza artificiale è ormai diventata parte integrante della nostra quotidianità, influenzando in modo crescente il modo in cui viviamo, lavoriamo e interagiamo.

Questa rapida diffusione ha inevitabilmente sollevato nuove sfide anche dal punto di vista giuridico, rendendo necessario intervenire con regole specifiche per guidarne lo sviluppo e l'utilizzo in modo sicuro e responsabile.

In Europa, la risposta più strutturata a questa esigenza è rappresentata dall'AI Act.

Questo regolamento cerca di trovare un equilibrio tra due esigenze diverse e non sempre facilmente compatibili: da un lato, favorire l'innovazione e mantenere la competitività europea nel settore dell'IA; dall'altro, proteggere i diritti fondamentali delle persone, mettendo al centro la sicurezza e la trasparenza.

Per raggiungere questo obiettivo, si è scelto di classificare i sistemi di intelligenza artificiale in base al rischio che comportano, applicando regole più o meno stringenti a seconda della pericolosità del loro utilizzo, secondo un approccio basato sulla proporzionalità.

Naturalmente, non mancano i punti critici: l'AI Act solleva ancora molti dubbi, sia per la sua complessità tecnica, sia per le difficoltà legate alla sua applicazione concreta.

Inoltre, le differenze tra gli approcci normativi dei vari Paesi rendono difficile trovare un equilibrio globale.

Quel che è certo è che l'approccio europeo si distingue nettamente da quello di altri grandi potenze mondiali come gli Stati Uniti, più orientati all'autoregolamentazione e alla flessibilità, e la Cina, che invece punta su un controllo statale molto forte.

Capire quale modello risulterà più efficace nel lungo periodo è tutt'altro che semplice, perché entrano in gioco tantissimi fattori diversi – economici, politici, culturali e sociali – che rendono il confronto complesso e ancora aperto.

5.2 Riflessioni sul futuro dell'IA e del diritto

Il rapporto tra intelligenza artificiale e diritto diventerà sempre più importante negli anni a venire.

L'IA non è più una tecnologia di nicchia, confinata ai laboratori o a qualche settore specialistico: oggi è parte integrante della vita di tutti i giorni, influenza l'economia e le decisioni sia pubbliche che private.

Di conseguenza, il diritto si trova di fronte a una sfida complessa, perché l'innovazione tecnologica corre molto più veloce rispetto ai tempi della legislazione tradizionale.

Da un lato, sarà fondamentale continuare a garantire che lo sviluppo dell'IA avvenga nel pieno rispetto dei diritti fondamentali, della trasparenza e della responsabilità.

Questo non significa solo aggiornare le leggi, ma anche riflettere con attenzione su principi giuridici che potrebbero dover essere reinterpretati o addirittura ripensati, per far fronte alle nuove sfide che questa tecnologia porta con sé.

Dall'altro lato, però, la regolamentazione non deve diventare un freno all'innovazione.

Il diritto dovrà essere flessibile, proporzionato e capace di adattarsi ai cambiamenti rapidi di questo settore, senza però perdere di vista la sua funzione principale: quella di proteggere e guidare la società.

Per questo, il ruolo delle autorità di controllo e della comunità scientifica sarà fondamentale per mantenere un equilibrio tra progresso tecnologico e interesse pubblico.

Infine, non bisogna dimenticare l'importanza di coinvolgere attivamente la società civile e soprattutto le nuove generazioni in questo dibattito.

Educazione digitale e consapevolezza etica saranno strumenti fondamentali per costruire un futuro in cui la tecnologia sia davvero al servizio delle persone, e non il contrario.

Damiano Pace

Bibliografia

- AI Chamber. (2025, 3 Aprile). *AI Chamber Issues New Open Letter Warning Against Overregulation in EU AI Code*.
- BicoccaNews. (2024). *Classificazione e categorie di rischio nell'AI Act*. Università degli Studi di Milano-Bicocca.
- Brookings Institution. Meltzer, J. P., & Kroll, J. A. (2024, Gennaio). *The EU and US Diverge on AI Regulation: A Transatlantic Comparison and Steps to Alignment*.
- Commissione Europea. (2021, 21 aprile). *Domande e risposte: Proposta di regolamento sull'intelligenza artificiale*.
- Digital Europe. (2023). *Position paper on the EU Artificial Intelligence Act*.
- European Commission. (2021). *Proposal for a Regulation laying down harmonised rules on artificial intelligence (Artificial Intelligence Act) – Impact Assessment*.
- European Commission. (2024). *Communication on the Implementation of the AI Act*.
- Finextra. (2023, 24 Novembre). *AI Act: Industry warns EU against overregulation*.
- Hertie School. (2024). *AI Governance: EU and US converge on risk-based approach amid stark differences*.
- Hua, X., & Zhang, Y. (2024). *AI Governance in China: State Control and Algorithmic Regulation*. *Journal of Asian Public Policy*, 17(2), 134–150.
- negg Group. (2025, 15 gennaio). *Pratiche IA vietate dall'AI Act*.
- Parlamento europeo. (2024, 13 marzo). *Regolamento sull'intelligenza artificiale (AI Act)*.
- Parlamento Europeo e Consiglio dell'Unione Europea. (2024). *Regolamento (UE) 2024/1689 del 13 giugno 2024: Norme armonizzate sull'intelligenza artificiale*. *Gazzetta ufficiale dell'Unione europea*, L 1689/1.
- Science Business. (2024). *US AI regulation: voluntary frameworks and sectoral approaches*.
- Unione Europea. (2007). *Trattato sull'Unione Europea (TUE)*, art. 5, par. 4.
- White House. (2023, 30 Ottobre). *Executive Order on the Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence*.
- Yao, D. (2024). *EU's AI Act vs. U.S. Approach on Regulating AI*. *The AI Innovator*.