

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 +***

***Intelligenza artificiale e tutela dei diritti
fondamentali: il caso Clearview AI***

Giulia Giorgi

0365274

Anno accademico 2024/2025

Sommario

Abstract	3
1. Il quadro normativo europeo sull'AI	4
1.1 Il GDPR e i limiti al trattamento automatizzato nei sistemi di AI.....	4
1.2 L'AI Act: struttura, principi e classificazione dei sistemi a rischio.....	5
2. Questioni giuridiche: trasparenza, bias e responsabilità.....	6
2.1 Profilazione algoritmica e diritto alla spiegazione.....	6
2.2 Bias e discriminazioni nei sistemi automatizzati	7
2.3 Le difficoltà nell'attribuzione della responsabilità giuridica.....	8
3. Il caso Clearview AI.....	10
3.1 Come funziona il sistema e quali sono le pratiche contestate	10
3.2 Le sanzioni delle autorità europee e le violazioni accertate	12
3.3 Il caso Clearview: esempio di tensione tra innovazione tecnologica e tutele giuridiche.....	15
Conclusione	17
Bibliografia	19

Abstract

Questa tesina analizza il rapporto tra intelligenza artificiale e diritti fondamentali, con un'attenzione particolare al contesto normativo europeo. L'obiettivo è capire in che modo le regole attuali riescano a proteggere i cittadini quando vengono utilizzati sistemi di AI, sempre più presenti nella vita quotidiana. In particolare, il lavoro si concentra su due strumenti centrali: il GDPR e l'AI Act. Entrambi affrontano, da punti di vista diversi, il tema della protezione dei dati e dei rischi legati all'uso di algoritmi decisionali. Nel corso dell'elaborato vengono trattati alcuni dei problemi più rilevanti, come la trasparenza dei sistemi, la difficoltà di comprendere il funzionamento degli algoritmi, il rischio di discriminazioni e le incertezze su chi sia davvero responsabile in caso di danni. Infine viene approfondito il caso di Clearview AI, un'azienda americana che ha raccolto milioni di immagini online per addestrare un sistema di riconoscimento facciale, finendo nel mirino di diverse autorità europee per violazione della normativa sulla privacy. Il caso mostra bene quanto sia difficile far valere le regole europee contro soggetti che operano fuori dai confini dell'UE, e quanto la normativa attuale debba ancora rafforzarsi per affrontare situazioni simili. La tesina si chiude con alcune riflessioni su come il diritto potrebbe evolversi per riuscire a tenere il passo con le nuove tecnologie, garantendo allo stesso tempo innovazione e rispetto dei diritti, e su come rendere la regolazione dell'AI più vicina alla vita ordinaria.

1. Il quadro normativo europeo sull'AI

1.1 Il GDPR e i limiti al trattamento automatizzato nei sistemi di AI

L'intelligenza artificiale, intesa come la capacità di sistemi informatici di apprendere da dati e prendere decisioni in modo autonomo, pone sfide significative al diritto europeo. Uno dei primi riferimenti normativi rilevanti in questo ambito è il Regolamento (UE) 2016/679, noto come GDPR, in vigore dal 25 maggio 2018. Sebbene non sia stato pensato specificamente per l'AI, molte delle sue disposizioni si applicano direttamente ai sistemi algoritmici che trattano dati personali. In particolare, l'articolo 22 del GDPR stabilisce che l'interessato ha diritto a non essere sottoposto ad una decisione basata unicamente su un trattamento automatizzato, compresa la profilazione, che produca effetti giuridici che lo riguardano o che incida in modo analogo significativamente sulla sua persona. Questo è particolarmente rilevante nel caso di sistemi di AI che prendono decisioni su assunzioni, erogazione di crediti, accesso a servizi pubblici o valutazioni di rischio. In queste situazioni, il trattamento è consentito solo in presenza di consenso esplicito, necessità contrattuale, o autorizzazione espressa dalla legge, e deve comunque garantire il diritto dell'interessato di ottenere una spiegazione sulla logica del trattamento. Questa disposizione ha dato origine al cosiddetto "diritto alla spiegazione" (right to explanation), che non è menzionato in modo esplicito nel testo del Regolamento ma viene ricavato dagli articoli 13, 14 e 15 del GDPR, che impongono obblighi di trasparenza sul funzionamento dei processi decisionali automatizzati. Un altro punto delicato è il trattamento dei dati biometrici, che nel GDPR rientrano tra le categorie particolari di dati personali (art. 9). I dati biometrici sono dati personali ottenuti tramite un trattamento tecnico specifico, che si riferiscono alle caratteristiche fisiche, fisiologiche o comportamentali di una persona fisica o che permettono o confermano l'identificazione univoca di tale persona. Il GDPR vieta il trattamento di questi dati, salvo eccezioni tassative, tra cui il consenso esplicito o motivi di interesse pubblico rilevante. Tuttavia, nella pratica, il ricorso al consenso risulta spesso problematico quando il soggetto interessato non è consapevole di essere identificato, come accade nei sistemi di riconoscimento facciale applicati in spazi pubblici. Oltre agli articoli menzionati, il GDPR introduce il principio dell'accountability (art. 5, par. 2), secondo cui è il titolare del trattamento a dover dimostrare di aver adottato misure adeguate a garantire la conformità alla normativa. Questo principio si traduce, nei contesti AI, nella necessità di documentare accuratamente la logica dell'algoritmo, le fonti dei dati, i criteri decisionali e i sistemi di controllo umano eventualmente

previsti. Tuttavia, molti commentatori hanno osservato come il GDPR, pur costituendo un riferimento fondamentale, non sia sufficiente da solo a regolamentare tutte le dinamiche dell'intelligenza artificiale, ad esempio Malgieri, G. nel suo articolo “Automated Decision-Making and the GDPR: Fascinating Legal Framework and Challenges Ahead” del 2020 ha sottolineato i principali limiti del GDPR nel contesto dell'AI, ad esempio che l'art. 22 non è chiaro su cosa significhi “decisione automatizzata” e su quali siano i confini delle eccezioni o che in molti casi è difficile per i soggetti interessati capire come far valere i propri diritti (es. diritto all'opposizione o alla spiegazione), soprattutto se non hanno conoscenze tecniche. Quindi il GDPR interviene quando il sistema è già operativo, ma non contiene strumenti adeguati per valutare i rischi nella fase di progettazione o per classificare i diversi tipi di AI in base alla loro pericolosità. Per questo motivo, la Commissione Europea ha sentito il bisogno di affiancare al GDPR una normativa più specifica: il regolamento sull'intelligenza artificiale, noto come AI Act.

1.2. L'AI Act: struttura, principi e classificazione dei sistemi a rischio

L'Artificial Intelligence Act che è stato approvato in via definitiva nel 2024, anche se entrerà pienamente in vigore nei prossimi anni, rappresenta già oggi un passaggio decisivo perchè è il primo quadro normativo europeo pensato appositamente per regolamentare l'intelligenza artificiale. L'obiettivo è duplice: da un lato garantire uno sviluppo armonico e sicuro dell'AI all'interno del mercato unico europeo, e dall'altro tutelare i diritti fondamentali delle persone, evitando usi distorti o pericolosi della tecnologia. L'elemento centrale dell'AI Act è il suo impianto normativo costruito attorno ad una classificazione dei rischi. Vale a dire che il regolamento non impone le stesse regole a tutti i sistemi di intelligenza artificiale, ma distingue tra diversi livelli di rischio a seconda dell'uso previsto e del potenziale impatto sugli individui. Classificazione dei rischi:

- Rischio inaccettabile: si tratta di sistemi considerati pericolosi per i diritti e le libertà delle persone, e per questo motivo vietati in modo assoluto. Rientrano in questa categoria, ad esempio, le tecnologie che manipolano subdolamente il comportamento umano o che assegnano punteggi sociali ai cittadini, come avviene in alcuni contesti autoritari.
- Rischio elevato: include sistemi utilizzati in ambiti particolarmente delicati – sanità, istruzione, lavoro, giustizia, sicurezza pubblica – in cui un errore dell'AI può avere conseguenze gravi sulla vita delle persone. Questi sistemi non sono vietati, ma per poter essere utilizzati devono rispettare requisiti molto stringenti: deve esserci un controllo umano

significativo, i dati usati devono essere di qualità, la documentazione tecnica deve essere accurata e ogni operazione deve poter essere tracciata.

- Rischio limitato: comprende applicazioni come chatbot (programma informatico progettato per simulare una conversazione con un essere umano), assistenti virtuali o generatori di immagini e testi. In questi casi, il rischio per gli utenti è contenuto, quindi le regole sono meno severe: è sufficiente, ad esempio, informare chiaramente l'utente che sta interagendo con un sistema automatizzato.
- Rischio minimo: riguarda la maggior parte delle applicazioni di uso quotidiano dell'AI, come i sistemi di raccomandazione su piattaforme di e-commerce o streaming. In questi casi non sono previsti obblighi specifici, perché il potenziale impatto negativo è ritenuto trascurabile.

Una delle innovazioni più significative introdotte dall'AI Act è il principio della prevenzione. Mentre regolamenti come il GDPR tendono ad attivarsi dopo che si è verificata una violazione, l'AI Act cerca di anticipare i problemi, imponendo misure fin dalla fase di progettazione e sviluppo. Gli sviluppatori saranno tenuti a valutare i rischi in anticipo, ad implementare sistemi di monitoraggio continui e a garantire che gli utenti siano sempre consapevoli di interagire con una macchina e non con una persona reale. Un altro aspetto centrale del regolamento riguarda la sua portata extraterritoriale. L'AI Act, infatti, si applica non solo alle aziende europee, ma anche a quelle extra-UE, se i loro sistemi producono effetti nel territorio dell'Unione Europea. Questo è particolarmente importante in casi come quello di Clearview AI, un'azienda statunitense i cui sistemi hanno avuto effetti concreti su cittadini europei, pur essendo sviluppati al di fuori dell'UE.

2. Questioni giuridiche: trasparenza, bias e responsabilità

2.1 Profilazione algoritmica e diritto alla spiegazione

La profilazione algoritmica è un processo che prevede l'analisi di dati personali, solitamente raccolti in modo continuo e su larga scala, al fine di individuare o anticipare comportamenti, preferenze e tratti individuali. Le informazioni così ottenute vengono successivamente impiegate per adottare decisioni che possono avere un impatto rilevante sulla vita delle persone, ad esempio nell'ambito dell'erogazione di servizi, nella determinazione dell'affidabilità creditizia o nella selezione dei candidati per un impiego. Dal punto di vista giuridico, ciò pone un problema centrale: l'asimmetria informativa tra chi progetta o gestisce l'algoritmo e il soggetto che ne subisce gli

effetti. L'individuo profilato spesso non ha strumenti adeguati per comprendere su quali basi venga classificato o valutato, né per contestare decisioni errate, arbitrarie o penalizzanti. Da questa esigenza nasce il dibattito attorno al cosiddetto diritto alla spiegazione ovvero il diritto di conoscere in modo chiaro, accessibile e comprensibile il funzionamento dei sistemi che incidono direttamente sui propri diritti ed interessi. Le tecnologie basate sull'apprendimento automatico, soprattutto quelle più complesse come le reti neurali profonde, tendono a produrre risultati efficaci ma difficili da spiegare in modo trasparente. Si parla in questo caso di opacità algoritmica ovvero una condizione per cui le decisioni risultano affidabili dal punto di vista statistico ma incomprensibili per chi le subisce o per chi dovrebbe valutarne la legittimità.

2.2 Bias e discriminazioni nei sistemi automatizzati

Tra le principali criticità legate all'utilizzo dei sistemi automatizzati vi è la presenza di bias algoritmici, ovvero distorsioni sistematiche che possono compromettere l'imparzialità e l'equità delle decisioni generate dai modelli. Sebbene gli algoritmi siano spesso considerati strumenti oggettivi, in realtà possono riflettere e amplificare pregiudizi insiti nei dati utilizzati per il loro addestramento o derivanti da scelte progettuali umane, anche involontarie. Quindi l'origine di tali distorsioni risiede o nei dataset storici, i quali rispecchiano spesso le disuguaglianze presenti nella società, oppure nei criteri tecnici impiegati nella progettazione del modello, come la selezione delle variabili, le metriche di performance o i target di ottimizzazione, che possono introdurre bias non evidenti. Ad esempio se dovessimo insegnare ad un computer come scegliere il miglior candidato per un lavoro, se gli si dà in input come esempio solo persone assunte in passato di cui tutti erano uomini, il computer penserà che gli uomini siano automaticamente i candidati migliori. Semplicemente imita quello che ha visto nei dati, ignorando il resto. Dal punto di vista giuridico, questo fenomeno pone una sfida concreta al rispetto del principio di non discriminazione, sancito dall'articolo 21 della Carta dei diritti fondamentali dell'Unione Europea e ribadito dalla normativa UE in materia di pari trattamento (Direttive 2000/43/CE e 2000/78/CE). Le discriminazioni prodotte dagli algoritmi possono essere dirette, quando si utilizzano categorie protette in modo esplicito (età, sesso, razza, ecc...), oppure indirette, quando le decisioni si basano su variabili apparentemente neutre ma correlate a caratteristiche sensibili. Un esempio emblematico è quello del software COMPAS (Correctional Offender Management Profiling for Alternative Sanctions), utilizzato negli Stati Uniti per valutare la probabilità di recidiva dei detenuti, per cui uno studio ha dimostrato che il sistema assegnava punteggi di rischio più elevati agli imputati afroamericani

rispetto a quelli di origine caucasica, a parità di condizioni. Simili distorsioni sono state documentate anche in altri ambiti, come il credito, il reclutamento aziendale o la sorveglianza predittiva. L'individuazione di queste discriminazioni è resa più difficile dalla scarsa trasparenza dei sistemi di intelligenza artificiale, specialmente di quelli basati su modelli complessi e poco interpretabili (deep learning, reti neurali, modelli black-box). Questa opacità algoritmica limita la possibilità di effettuare controlli, esercitare diritti di contestazione e attribuire responsabilità in caso di effetti lesivi. In risposta a tali rischi, l'AI Act ha introdotto obblighi specifici per i cosiddetti sistemi ad alto rischio, tra cui rientrano quelli utilizzati nei settori del lavoro, del credito, della giustizia e dell'istruzione. Le misure previste includono: la documentazione dettagliata dei dati utilizzati, la valutazione preventiva dei rischi, la tracciabilità del processo decisionale e l'obbligo di supervisione umana. Tuttavia, l'efficacia concreta di queste disposizioni dipenderà dalla loro attuazione pratica, dalla disponibilità di competenze adeguate presso le autorità di controllo, e dalla capacità di sviluppare strumenti tecnici affidabili per il monitoraggio, la spiegabilità e l'audit dei sistemi algoritmici. In conclusione, bias e discriminazioni nei sistemi automatizzati non sono meri effetti collaterali, ma espressione di dinamiche strutturali che possono compromettere profondamente i diritti fondamentali delle persone. Far fronte a queste problematiche richiede un approccio multidisciplinare e proattivo, capace di integrare sia innovazione tecnologica che principi costituzionali che responsabilità sociale, sin dalla fase di progettazione e poi nell'uso delle tecnologie intelligenti. Per affrontare la criticità dei bias legati a sistemi di AI si dovrebbero correggere i dataset prima dell'addestramento eliminando gli squilibri tramite ribilanciamento delle classi, anonimizzazione di variabili sensibili e tecniche di etichettatura equa, o ancora applicare vincoli di equità direttamente nel modello e poi correggere le predizioni dopo che l'algoritmo ha prodotto un output, ad esempio raffinando soglie di decisione per ridurre disparità nei gruppi protetti.

2.3 Le difficoltà nell'attribuzione della responsabilità giuridica

Ad oggi quando un algoritmo prende decisioni che possono causare danni, capire chi ne è davvero responsabile non è affatto semplice, ed è diventato uno dei punti critici più urgenti con cui il diritto deve fare i conti. La complessità dei modelli algoritmici, la difficoltà di ricostruire il nesso causale tra input e output, e la pluralità di soggetti coinvolti nel ciclo di vita di un sistema di intelligenza artificiale rendono problematica l'applicazione degli schemi classici della responsabilità civile e amministrativa. Nell'ambito delle decisioni tradizionali, il soggetto che commette un illecito può

essere identificato con relativa chiarezza; al contrario, nei sistemi algoritmici la responsabilità si diluisce tra sviluppatori, fornitori, utenti finali, titolari del trattamento dei dati e soggetti pubblici o privati che adottano il sistema. Una delle principali difficoltà risiede nella complessa struttura decisionale tipica degli algoritmi di apprendimento automatico, spesso non completamente prevedibile neppure da chi ha progettato il sistema. Questo fenomeno, comunemente definito come “opacità algoritmica”, rende ostica la ricostruzione del processo logico che ha condotto ad una determinata decisione e, di conseguenza, ostacola l’accertamento del nesso causale tra un comportamento umano e l’evento dannoso. Non a caso, la letteratura giuridica recente insiste sulla necessità di sviluppare modelli di responsabilità che tengano conto della natura distribuita e probabilistica delle decisioni automatizzate, superando l’idea che ogni responsabilità debba necessariamente derivare da un comportamento umano volontario o consapevole. La giurisprudenza europea ha cominciato ad affrontare la questione, come dimostra il caso SyRI (Systeem Risico Indicatie) nei Paesi Bassi. Questo sistema, utilizzato per identificare potenziali frodi nel settore del welfare, elaborava grandi quantità di dati personali e attribuiva punteggi di rischio ai cittadini sulla base di variabili socio-demografiche. Il sistema operava però in modo opaco e non offriva agli interessati la possibilità di conoscere i criteri decisionali o di contestarne gli esiti. Nel 2020, il Tribunale dell’Aia ha dichiarato l’illegittimità del sistema, affermando che SyRI violava l’articolo 8 della Convenzione europea dei diritti dell’uomo, in quanto interferiva in modo sproporzionato con il diritto alla vita privata. La Corte ha inoltre evidenziato che l’assenza di trasparenza e controllo umano nella valutazione del rischio comprometteva il diritto ad una tutela effettiva, sancito anche dalla Carta dei diritti fondamentali dell’Unione Europea. Questo caso dimostra come le difficoltà nell’attribuzione della responsabilità non siano solo un problema di diritto civile, ma coinvolgano anche il diritto costituzionale e quello sovranazionale. Il punto centrale non è solo intercettare il colpevole, ma garantire che le persone colpite da decisioni algoritmiche lesive possano esercitare i propri diritti, accedere ad un ricorso effettivo e ottenere una spiegazione comprensibile. In questo senso, l’assenza di responsabilità chiaramente definita rischia di tradursi in un vuoto di tutela giuridica. A livello normativo, l’Unione Europea ha tentato di colmare alcune di queste lacune con il Regolamento sull’intelligenza artificiale di cui precedentemente si è discusso, che introduce obblighi di trasparenza, tracciabilità e supervisione umana per i sistemi ad alto rischio. Tuttavia, il regolamento non istituisce una nuova categoria autonoma di responsabilità. L’attribuzione del danno rimane affidata al diritto nazionale, con l’ulteriore difficoltà di armonizzare norme di natura tecnica, informatica e giuridica. Per rispondere a questa esigenza, nel 2022 la Commissione

Europea ha presentato una proposta di direttiva sulla responsabilità per danno da intelligenza artificiale (COM/2022/496), che prevede una presunzione di nesso causale in favore della parte lesa nei casi in cui venga accertata una violazione degli obblighi previsti per i fornitori o gli utilizzatori del sistema. Si tratta di un meccanismo innovativo che mira a ridurre l'onere probatorio, riequilibrando l'asimmetria informativa tra chi utilizza il sistema e chi ne subisce gli effetti. Tuttavia la presenza di molti attori coinvolti nel processo decisionale continua a rendere difficile stabilire chi debba rispondere in caso di danni. In molti casi, il danno non deriva da un errore evidente, ma da una serie di micro-decisioni progettuali, di addestramento e di implementazione che, nel loro complesso, producono effetti discriminatori o ingiusti. In questo contesto si discute dell'opportunità di introdurre forme di responsabilità oggettiva sul modello della direttiva 85/374/CEE sulla responsabilità da prodotto difettoso, oppure di istituire obblighi assicurativi per i sistemi automatizzati che operano in settori ad alto impatto, come la sanità, il credito o l'amministrazione pubblica. Un'altra proposta è quella di attribuire ai sistemi di intelligenza artificiale una forma limitata di personalità giuridica, simile a quella delle persone giuridiche o delle società. Tuttavia, il Parlamento Europeo ha già espresso una netta opposizione a questa ipotesi, in una risoluzione del 2017, sottolineando che ciò rischierebbe di deresponsabilizzare i soggetti umani e di compromettere i diritti delle persone danneggiate. In definitiva, l'attribuzione della responsabilità giuridica nei sistemi automatizzati non può essere risolta con l'applicazione meccanica delle categorie esistenti. Richiede un ripensamento profondo del rapporto tra tecnologia, diritto e tutela dei diritti fondamentali. È necessario riconoscere che l'innovazione tecnologica deve procedere di pari passo con l'evoluzione delle garanzie giuridiche, affinché nessun danno causato da un algoritmo resti privo di un responsabile e nessun individuo resti privo di tutela.

3. Il caso Clearview AI

3.1 Come funziona il sistema e quali sono le pratiche contestate

Clearview AI è una società statunitense fondata nel 2017 da Hoan Ton-That, programmatore australiano, e Richard Schwartz, ex consulente politico. L'azienda è rapidamente finita al centro del dibattito pubblico per il suo ruolo altamente controverso nel campo del riconoscimento facciale, in particolar modo per aver sviluppato una tecnologia capace di identificare volti umani con estrema rapidità e precisione, sfruttando un'enorme banca dati di immagini raccolte da internet. Il suo software è stato presentato come uno strumento investigativo al servizio delle forze dell'ordine, ma

le modalità con cui è stato progettato, distribuito e utilizzato hanno sollevato numerose critiche sia sotto il profilo tecnico che giuridico. Il sistema di Clearview si basa sul web scraping, ovvero una tecnica informatica che consente di “catturare” e archiviare automaticamente contenuti presenti su siti pubblici, inclusi social media, portali di notizie, blog e database accessibili senza autenticazione. Secondo le dichiarazioni della stessa società, il software ha raccolto oltre 30 miliardi di immagini da fonti pubblicamente disponibili, molte delle quali contenenti volti identificabili. Tali immagini vengono elaborate attraverso un algoritmo di deep learning, in grado di estrarre vettori numerici (embedding) rappresentativi delle caratteristiche facciali di ogni soggetto. Questo perché gli algoritmi non confrontano direttamente le immagini, ma confrontano questi vettori e più i vettori sono simili, maggiore è la probabilità che i due volti appartengano alla stessa persona. Quindi quando l’utente carica un’immagine da analizzare, il sistema confronta il volto con gli embedding presenti nel database, restituendo risultati visivi con i volti simili e i link alle fonti da cui le immagini sono state prelevate. Questa modalità operativa ha reso Clearview uno strumento estremamente potente per l’identificazione di individui sconosciuti, specialmente in contesti di sicurezza pubblica e investigazioni criminali. Secondo vari report, il software sarebbe stato adottato da migliaia di agenzie di polizia negli Stati Uniti. Tuttavia, in poco tempo sono emerse criticità profonde, legate sia alla legittimità del trattamento dei dati, sia all’assenza di controlli e limiti nell’utilizzo. La prima grande questione riguarda la raccolta non consensuale delle immagini. Clearview non richiede l’autorizzazione degli interessati, né fornisce informative sul trattamento, sulle finalità, sulla conservazione o sui diritti esercitabili. L’azienda si è difesa sostenendo che le immagini raccolte sono già “pubbliche”, in quanto presenti su internet, ma questa posizione si scontra con il principio cardine del Regolamento (UE) 2016/679 (GDPR) secondo cui ogni trattamento di dati personali deve fondarsi su una base giuridica valida e deve rispettare i principi di liceità, trasparenza e minimizzazione. L’articolo 9 del GDPR, che disciplina il trattamento di categorie particolari di dati, incluso il dato biometrico, stabilisce che tale trattamento è essenzialmente vietato, salvo specifiche eccezioni, tra cui il consenso esplicito o un rilevante interesse pubblico previsto dal diritto dell’Unione o degli Stati membri. Il ricorso al concetto di “dato manifestamente reso pubblico” come giustificazione appare dunque forzato, soprattutto considerando che la pubblicazione di una fotografia su un social network non implica automaticamente la rinuncia alla tutela dei propri dati biometrici. In secondo luogo, Clearview non ha previsto alcun meccanismo di accesso, rettifica o cancellazione dei dati da parte degli interessati, impedendo di fatto l’esercizio dei diritti riconosciuti dal GDPR agli articoli 15-22. Gli utenti non

possono sapere se le loro immagini siano presenti nel database, né possono opporsi al trattamento. Inoltre, non sono state definite politiche chiare sulla conservazione dei dati né sulle misure di sicurezza informatica adottate per proteggerli. Questo vuoto normativo ha generato un rischio concreto di violazione sistematica della privacy, amplificato dal fatto che il sistema può essere utilizzato anche da soggetti privati, con finalità non sempre lecite o controllate. Un altro punto debole del sistema riguarda l'ambiguità della destinazione d'uso. Sebbene Clearview si sia presentata come fornitore di tecnologie per la pubblica sicurezza, in più occasioni l'azienda ha offerto accesso al proprio software ad imprese private, agenzie investigative, banche e persino individui in cerca di servizi di sorveglianza o monitoraggio. Secondo un'inchiesta del New York Times, in alcune fasi iniziali dello sviluppo la società avrebbe concesso accessi di prova ad amici, finanziatori o aziende interessate, senza alcuna regolamentazione esterna. Questo solleva dubbi circa il rispetto del principio di limitazione delle finalità, anch'esso previsto dal GDPR, e del principio di accountability, che richiede al titolare del trattamento di dimostrare la conformità del proprio operato ai principi di protezione dei dati. Emergono inoltre criticità che riguardano il funzionamento tecnico del sistema in quanto i modelli di riconoscimento facciale, anche quando costruiti su dataset ampi, possono generare errori significativi nei confronti di soggetti appartenenti a minoranze etniche o persone con tratti non standardizzati, a causa della rappresentazione non bilanciata nei dati di addestramento. Diversi studi, tra cui quelli del MIT Media Lab e dell'Electronic Frontier Foundation, hanno dimostrato che gli algoritmi di riconoscimento facciale hanno tassi di errore molto più elevati per le persone di colore e asiatiche rispetto ai soggetti caucasici. Questi squilibri non sono solo problemi tecnici, ma hanno un impatto diretto sui diritti delle persone e danno luogo a forme di discriminazione. Alla luce di quanto detto, Clearview AI rappresenta un caso emblematico di come l'uso di tecnologie avanzate, in assenza di regole chiare e strumenti di controllo efficaci, possa compromettere diritti fondamentali come la protezione dei dati personali, la non discriminazione e la libertà individuale. L'uso massiccio di dati, la portata internazionale del trattamento e la scarsa trasparenza hanno spinto varie autorità per la privacy a prendere provvedimenti, avviando indagini e imponendo sanzioni.

3.2 Le sanzioni delle autorità europee e le violazioni accertate

L'espansione globale di Clearview AI ha attirato fin da subito l'attenzione delle autorità europee incaricate della tutela dei dati personali, in particolare per il modo con cui l'azienda ha costruito il proprio database, come detto mediante il web scraping massivo di fotografie e contenuti visivi

disponibili online. La logica sottostante a questo tipo di raccolta, apparentemente basata sul fatto che le immagini siano “pubblicamente accessibili”, si è rivelata incompatibile con i principi fondamentali sanciti dal GDPR. La portata delle violazioni è stata tale da innescare una risposta coordinata a livello sovranazionale per cui diverse autorità per la protezione dei dati dei Paesi membri hanno aperto procedimenti distinti ma convergenti, arrivando ad imporre sanzioni economiche rilevanti e ad ordinare la cessazione delle attività di trattamento dei dati da parte dell’azienda sul territorio europeo. Uno dei casi più significativi, per entità della sanzione e chiarezza argomentativa, è quello del Garante per la protezione dei dati personali italiano, che nel marzo 2022 ha concluso un’istruttoria formale contro Clearview AI con l’imposizione di una sanzione amministrativa pecuniaria pari a 20 milioni di euro, l’importo massimo previsto dall’articolo 83 del GDPR per le violazioni più gravi. Il provvedimento del Garante italiano ha avuto un valore non solo sanzionatorio, ma anche rappresentativo, perché ha ribadito che il semplice fatto di non operare fisicamente in un determinato territorio non solleva un’azienda dagli obblighi derivanti dalla normativa europea se tratta dati di soggetti residenti nell’UE. Clearview, pur non avendo filiali o stabilimenti in Italia, è stata ritenuta responsabile in quanto raccoglieva e analizzava informazioni (immagini biometriche) di cittadini italiani, attraverso il proprio algoritmo. Contestualmente anche altre autorità europee si sono espresse in maniera analoga come l’Information Commissioner’s Office (ICO) del Regno Unito che ha emesso nel novembre 2022 una multa di 7,5 milioni di sterline, accusando Clearview di aver trattato in modo illecito dati personali di milioni di cittadini britannici senza alcuna base giuridica. La sanzione è stata accompagnata da un’ingiunzione di cancellazione di tutti i dati relativi ai soggetti del Regno Unito, e da un divieto futuro di raccolta dati su cittadini britannici senza il rispetto delle regole locali. A questo si aggiunge un’analoga sanzione della CNIL francese, autorità garante francese, che ha ordinato a Clearview di cessare il trattamento dei dati di soggetti francesi e di cancellarli entro un termine prestabilito, ritenendo che l’azienda avesse violato gravemente i principi di liceità e trasparenza. Simili misure sono state prese in Austria, Grecia e Germania, dove le autorità competenti hanno aperto istruttorie formali e avviato procedimenti amministrativi. Nonostante l’assenza, nella maggior parte dei casi, di una sede giuridica europea per Clearview, le autorità hanno fatto leva sull’articolo 3 del GDPR, che estende l’applicazione del regolamento anche a soggetti extra-UE che trattano dati di cittadini europei nel contesto della fornitura di beni o servizi, o del monitoraggio del comportamento online. È proprio quest’ultimo il criterio attivato nel caso di Clearview, la cui tecnologia di tracciamento ed identificazione attraverso il riconoscimento facciale

è stata ritenuta una forma avanzata e sistematica di sorveglianza comportamentale. Per quanto concerne le violazioni maggiormente criticate a Clearview rientrano:

- Violazione del principio di liceità, correttezza e trasparenza (art. 5, par. 1, lett. a): Clearview non ha fornito alcuna informazione ai soggetti interessati circa l'uso dei loro dati personali, impedendo loro di esercitare consapevolmente i propri diritti.
- Mancanza di una base giuridica per il trattamento (art. 6): l'azienda ha trattato dati personali senza il consenso degli interessati né altro fondamento giuridico valido.
- Trattamento illecito di dati biometrici (art. 9): secondo le autorità, le immagini facciali raccolte costituiscono dati biometrici utilizzati per l'identificazione univoca di una persona fisica, il cui trattamento è vietato salvo specifiche eccezioni, che nel caso in esame non sussistevano.
- Mancata osservanza del diritto all'informazione (art. 14): i dati venivano raccolti da fonti pubbliche senza che i soggetti fossero informati, né veniva indicata l'identità del titolare del trattamento, le finalità del trattamento o i diritti esercitabili.
- Ostacolo all'esercizio dei diritti (artt. 15–22): le autorità hanno riscontrato l'impossibilità pratica, per gli interessati, di accedere ai propri dati o di richiederne la cancellazione, configurando una grave lesione dei diritti fondamentali sanciti dal GDPR.

Le autorità non hanno visto queste violazioni come semplici eccezioni o errori occasionali, ma come parte di un sistema organizzato e continuo, aggravato dal fatto che il servizio operava su scala globale e con tecnologie molto intrusive. È anche per questo che le sanzioni sono state così pesanti: non solo per punire quello che è stato fatto, ma per il potenziale impatto che un uso simile dei dati potrebbe avere su milioni di persone. Un aspetto centrale nella vicenda riguarda il principio di extraterritorialità del GDPR, che consente l'applicazione delle norme europee anche a soggetti esterni all'UE qualora le loro attività coinvolgano direttamente cittadini europei. Clearview ha tentato, in più occasioni, di contestare la competenza giurisdizionale delle autorità europee, sostenendo di non essere soggetta al regolamento perché priva di sedi operative nel continente. Tuttavia, le autorità garanti hanno ribadito che il semplice trattamento dei dati dei cittadini europei, indipendentemente dal luogo fisico in cui ciò avviene, impone il rispetto pieno del GDPR. Le autorità hanno fatto esplicito riferimento all'art. 3, par. 2, lett. b, che impone l'applicazione del regolamento ad ogni attività di monitoraggio sistematico dei comportamenti dei soggetti interessati che si trovano nel territorio dell'Unione. Nonostante le sanzioni, Clearview ha finora dimostrato un

atteggiamento elusivo, non uniformandosi agli ordini di cancellazione dei dati e non implementando meccanismi efficaci per garantire l'esercizio dei diritti da parte degli interessati. Questo atteggiamento ha sollevato ulteriori interrogativi sulla capacità effettiva delle istituzioni europee di far rispettare i propri provvedimenti a soggetti extra-UE, evidenziando la necessità di meccanismi di cooperazione internazionale più efficaci. Il caso Clearview si è imposto come un precedente giurisprudenziale e politico di grande rilevanza, non solo per la sua dimensione tecnica e legale, ma anche per le implicazioni etiche che solleva. In un contesto in cui le tecnologie di sorveglianza si fanno sempre più integrate nella vita quotidiana, la risposta delle autorità europee ha rappresentato un tentativo chiaro di riaffermare la centralità dei diritti fondamentali della persona, anche davanti a innovazioni tecnologiche potenzialmente utili ma eccessivamente invasive. Infine questo caso ha dimostrato la volontà dell'Unione Europea di riaffermare con forza che valori come la dignità umana, la trasparenza e il controllo sui dati personali non sono in alcun modo negoziabili.

3.3 Il caso Clearview: esempio di tensione tra innovazione tecnologica e tutele giuridiche

Il caso Clearview AI non è rilevante soltanto per le sue specificità tecniche o per le reazioni che ha suscitato, ma anche e soprattutto perché esprime la crescente tensione tra la spinta innovativa delle tecnologie digitali e la necessità di garantire diritti fondamentali all'interno di ordinamenti democratici. Infatti questa vicenda solleva interrogativi più ampi sul rapporto tra ciò che la tecnologia è in grado di fare e ciò che dovrebbe essere autorizzata a fare, alla luce dei vincoli giuridici, etici e sociali che regolano la convivenza civile. Ciò che rende il caso particolarmente significativo è la profonda asimmetria tra potere tecnologico e capacità di controllo normativo in quanto gli strumenti di intelligenza artificiale e machine learning si evolvono ad un ritmo accelerato, al contrario, il diritto procede con tempi decisamente più lenti e non riesce ad offrire risposte tempestive. Non si tratta semplicemente di un ritardo burocratico, ma di un disallineamento strutturale: le logiche dell'innovazione si basano su velocità, sperimentazione e adattamento continuo invece quelle del diritto su prudenza, legittimazione e tutela dell'interesse generale. Questo scarto genera un vuoto regolativo che attori tecnologici privati, spesso transnazionali, tendono ad occupare in modo unilaterale, definendo di fatto le condizioni di utilizzo di strumenti che incidono profondamente sulla vita delle persone. Clearview AI, in questo senso, è un caso esemplare di governance tecnologica assente o posticipata, in cui l'uso di un sistema ad alto impatto sociale ha preceduto ogni confronto pubblico, riflessione collettiva o valutazione dei rischi. La

tensione che emerge non è dunque tra tecnologia e legge in senso astratto, ma tra modelli opposti di gestione dell'innovazione perché da un lato vi è una visione orientata al profitto e alla sperimentazione libera, dove l'interesse economico guida l'evoluzione degli strumenti senza un confronto preliminare con i valori costituzionali, mentre dall'altro si afferma un modello europeo che mira a regolare l'innovazione sulla base del rischio e del rispetto dei diritti, promuovendo un'idea di progresso che non sia solo tecnico, ma anche civile e sostenibile. In questa prospettiva, Clearview AI diventa un “test” per le istituzioni democratiche, chiamate non solo a reagire a posteriori con sanzioni, ma a sviluppare meccanismi di prevenzione, valutazione d'impatto e controllo preventivo delle tecnologie. Il passaggio da una logica reattiva ad una logica proattiva è oggi più che mai necessario, soprattutto quando si ha a che fare con strumenti che operano in modo opaco, automatizzato e potenzialmente permanente nella vita delle persone. Un altro elemento centrale della tensione riguarda il ruolo dell'opinione pubblica e della consapevolezza collettiva. L'innovazione tecnologica spesso avanza senza che i cittadini ne siano pienamente informati o messi in condizione di esprimere un consenso libero e consapevole. La mancanza di trasparenza non è solo un problema tecnico, ma un nodo politico e culturale, che impedisce il formarsi di un dibattito pubblico informato sul tipo di società digitale che si vuole costruire. In tale ottica, casi come quello di Clearview non sollevano solo dubbi sulla legalità di una tecnologia in particolare, ma aprono interrogativi più profondi sulla legittimità delle decisioni prese senza un reale controllo democratico. In ultima analisi, il caso Clearview non va letto come un'eccezione o una deviazione isolata, ma come una manifestazione concreta delle sfide che ogni ordinamento giuridico dovrà affrontare nel confrontarsi con l'intelligenza artificiale e i suoi usi più indiscreti.

Conclusione

In questa tesina si è analizzato il quadro normativo europeo, i principali nodi giuridici e il caso di Clearview AI, e indubbiamente si evince che il rapporto tra intelligenza artificiale e diritti fondamentali è ad oggi uno dei punti più delicati e urgenti del diritto digitale. Le regole ci sono, ma non sempre riescono ad affrontare con efficacia la complessità delle tecnologie attuali, infatti il GDPR e l'AI Act rappresentano tentativi significativi di rispondere alle sfide poste da un uso sempre più diffuso e sofisticato dell'AI, ma da soli non bastano a colmare tutte le lacune. Alcune criticità, come la difficoltà nell'attribuire la responsabilità, l'opacità degli algoritmi e la mancanza di trasparenza nei sistemi automatizzati, richiedono strumenti nuovi, più agili e più vicini alla realtà delle persone che queste tecnologie le subiscono, spesso senza saperlo. Per rimediare a queste mancanze si potrebbe pensare di elaborare una Carta dei diritti digitali per l'intelligenza artificiale che sia uno strumento chiaro e accessibile e che includa il diritto a sapere se si è sottoposti ad una decisione automatizzata, il diritto a capirne le logiche o il diritto a dire no. Un secondo punto riguarda la trasparenza verso i cittadini perché oggi molti sistemi di intelligenza artificiale operano in modo nascosto o comunque poco comprensibile per chi ne è coinvolto e invece sarebbe importante introdurre l'obbligo di fornire informazioni chiare e sintetiche ogni volta che viene utilizzato un sistema automatizzato per valutare, selezionare o profilare una persona. Le informazioni dovrebbero includere quali dati vengono usati, per quali scopi, con quali margini di errore, e chi è responsabile del funzionamento del sistema; il tutto in un linguaggio semplice e non tecnico di modo che sia comprensibile per chiunque. Un altro elemento che probabilmente è fondamentale è la partecipazione democratica poiché attualmente le decisioni sull'uso dell'intelligenza artificiale vengono prese quasi sempre da soggetti tecnici, politici o economici, eppure le conseguenze riguardano tutti. Per questo si dovrebbe proporre la creazione di comitati etici territoriali o settoriali (per esempio nella scuola, nella sanità, nel lavoro), composti anche da cittadini, studenti, associazioni e realtà locali, non per sostituire le autorità competenti ma per portare un punto di vista diverso, più vicino alla realtà. Un altro punto importante è il controllo continuo sui sistemi automatizzati dato che questi sistemi apprendono e si aggiornano nel corso del tempo, serve quindi un monitoraggio periodico in grado di verificare che le decisioni restino eque, trasparenti e rispettose dei diritti. A tutto questo si aggiunge la questione della formazione che spesso viene trascurata, invece è importante introdurre corsi formativi di base sull'intelligenza artificiale già nelle scuole superiori e in tutti i corsi universitari, anche quelli non scientifici. In

conclusione, la regolazione dell'intelligenza artificiale rimane, ad oggi, una questione ancora aperta che impone al diritto di misurarsi con sfide nuove, spesso senza precedenti, ma inevitabili per garantire tutele reali nel presente e nel futuro.

Bibliografia

Regolamento (UE) 2016/679 (GDPR).

Artificial Intelligence Act, COM(2021) 206 final – Regolamento UE sull'intelligenza artificiale (2024).

Carta dei diritti fondamentali dell'Unione Europea, 2000/C 364/01.

Direttiva 2000/43/CE del Consiglio del 29 giugno 2000.

Direttiva 2000/78/CE del Consiglio del 27 novembre 2000.

Proposta di Direttiva sulla responsabilità da intelligenza artificiale, COM(2022) 496 final, 28 settembre 2022.

Malgieri, G. (2020). *Automated Decision-Making and the GDPR: Fascinating Legal Framework and Challenges Ahead*. *Computer Law & Security Review*, 36, 105567.

Hill, K. (2020). *The Secretive Company That Might End Privacy as We Know It*. *The New York Times*, 18 gennaio.

Buolamwini, J., & Gebru, T. (2018). *Gender Shades: Intersectional Accuracy Disparities in Commercial Gender Classification*. *Proceedings of Machine Learning Research*, 81, 1–15.

Garante per la protezione dei dati personali (Italia). Provvedimento contro Clearview AI del 10 marzo 2022.