

ARTICOLI

Intelligenza artificiale e polizze long-term care: tra personalizzazione del rischio ed esclusione sociale

UGO MALVAGNA

Professore Associato di Diritto dell'Economia
Università di Trento

Dialoghi di Diritto dell'Economia

Rivista diretta da

Danny Busch, Raffaele Lener, Roberto Natoli, Andrea Sacco Ginevri,
Filippo Sartori, Antonella Sciarrone Alibrandi

Direttore editoriale

Andrea Marangoni

Coordinatore editoriale

Francesco Petrosino

Direttori di area

Attività, governance e regolazione bancaria

Prof. Alberto Urbani, Prof. Diego Rossano, Prof. Francesco Ciraolo, Prof.ssa Carmela Robustella, Prof. Gian Luca Greco, Prof. Federico Riganti, Dott. Luca Lentini

Mercato dei capitali finanza strutturata

Prof. Matteo De Poli, Prof. Filippo Annunziata, Prof. Ugo Malvagna, Dott.ssa Anna Toniolo

Assicurazioni e previdenza

Prof. Paoloefisio Corrias, Prof. Michele Siri, Prof. Pierpaolo Marano, Prof. Giovanni Maria Berti De Marinis, Dott. Massimo Mazzola

Contratti di impresa, concorrenza e mercati regolati

Prof.ssa Maddalena Rabitti, Prof.ssa Michela Passalacqua, Prof.ssa Maddalena Semeraro, Prof.ssa Mariateresa Maggiolino

Diritto della crisi di impresa e dell'insolvenza

Prof. Aldo Angelo Dolmetta, Prof. Gianluca Mucciarone, Prof. Francesco Accettella, Dott. Antonio Didone, Prof. Alessio di Amato

Fiscalità finanziaria

Prof. Andrea Giovanardi, Prof. Nicola Sartori, Prof. Francesco Albertini

Istituzioni dell'economia e politiche pubbliche

Prof.ssa Michela Passalacqua, Prof. Francesco Moliterni, Prof. Giovanni Luchena, Dott.ssa Stefania Cavaliere, Dott. Lorenzo Rodio Nico

Criteri di Revisione

I contributi proposti alla Rivista per la pubblicazione sono sottoposti a una previa valutazione interna da parte della Direzione o di uno dei Direttori d'Area; il quale provvede ad assegnare il contributo a un revisore esterno alla Rivista, selezionato, *rationes materiae*, fra professori, ricercatori o assegnisti di ricerca.

La rivista adotta il procedimento di revisione tra pari a singolo cieco (*single blind peer review*) per assicurarsi che il materiale inviato rimanga strettamente confidenziale durante il procedimento di revisione.

Qualora il valutatore esprima un parere favorevole alla pubblicazione subordinato all'introduzione di modifiche, aggiunte o correzioni, la Direzione si riserva di negare la pubblicazione dell'articolo. Nel caso in cui la Direzione decida per la pubblicazione, deve verificare previamente che l'Autore abbia apportato le modifiche richieste dal Revisore.

Qualora il revisore abbia espresso un giudizio negativo, il contributo può essere rifiutato oppure inviato, su parere favorevole della maggioranza dei Direttori dell'area competente *rationes materiae*, a un nuovo revisore esterno per un ulteriore giudizio. In caso di nuovo giudizio negativo, il contributo viene senz'altro rifiutato.

Intelligenza artificiale e polizze long-term care: tra personalizzazione del rischio ed esclusione sociale*

UGO MALVAGNA

Professore Associato di Diritto dell'Economia
Università di Trento

SOMMARIO: Introduzione; 1. Il contesto normativo di riferimento; 1.1. Un quadro di sintesi del Regolamento UE 2024/1689 (AI ACT); 1.2. GDPR e governance dei dati; 1.2.1. Nozione di dato personale; 1.2.2. Principio di trasparenza e processi decisionali automatizzati (ADM); 1.2.3. (segue) Diritto alla leggibilità della decisione algoritmica; 1.2.4. (segue) Diritto all'intervento umano; 1.3. Buone pratiche per l'utilizzo di sistemi automatizzati; 1.4. Liceità del trattamento automatizzato; 2. Il rischio di discriminazione (e di conseguente esclusione sociale) derivante dall'utilizzo di sistemi di intelligenza artificiale nel diritto delle assicurazioni; 2.1. Il principio di non discriminazione nei contratti assicurativi; 2.2. L'inquadramento del principio di non discriminazione nella correttezza dei rapporti contrattuali; 2.3. L'IA e i rischi connessi alla discriminazione nel settore assicurativo; 2.4. L'adozione di buone pratiche da parte dell'assicuratore; 2.4.1. Obblighi di trasparenza e spiegabilità; 2.4.2. Il controllo sull'algoritmo; 2.4.3. Supervisione umana (*human oversight*); 2.4.4. Governo dei dati e tenuta dei registri (*data governance and record keeping*); 2.4.5. Robustezza e prestazione (*Robustness and performance*); 2.4.6. Valutazioni di impatto; 2.5. IA e processi di *product oversight and governance*; 3. *Long-term care policies*, intelligenza artificiale e *wearable devices*. Quali riscontri dal mercato?; 3.1. Utilizzo e potenziali esternalità positive dell'IA e dei *wearable data* nelle *polizze long-term care*; 3.2. Panoramica dei mercati interno, europeo ed extraeuropeo dei prodotti vita che utilizzano sistemi di IA o *wearable devices*.

* Il presente contributo si inserisce nell'ambito del PRIN denominato «*Human well-being in times of tech management of emergency*» (*principal investigator*: prof. Ugo Malvagna). Si ringraziano il prof. Francesco Petrosino, la dott.ssa Veronica Zerba, la dott.ssa Paola Dassisti e la dott.ssa Angela Aromolo de Rinaldis per il supporto nell'attività di ricerca e redazionale.

Introduzione

Il presente contributo è inteso ad esaminare il progressivo e sempre più strutturale ricorso ai sistemi di intelligenza artificiale nell'attività assicurativa, con specifico riferimento ai prodotti *long-term care*, collocando tale fenomeno nel più ampio contesto delle trasformazioni indotte dalla gestione tecnologica del rischio e delle emergenze e delle loro ricadute sul benessere umano. L'analisi muove dall'esigenza di valutare in quale misura l'integrazione dell'IA nei processi di *underwriting*, *pricing* e gestione del rapporto assicurativo – la quale consente significativi incrementi di efficienza operativa, accuratezza predittiva e personalizzazione dell'offerta – sollevi questioni rilevanti sul piano della tutela dei diritti fondamentali, della protezione dei dati personali e, in modo peculiare, della prevenzione di fenomeni discriminatori e di esclusione assicurativa a danno dei soggetti strutturalmente più vulnerabili.

Il punto di avvio dell'indagine è rappresentato dal quadro normativo europeo, profondamente ridefinito negli ultimi anni in risposta alla diffusione di sistemi algoritmici ad elevato impatto. In tale prospettiva, il Regolamento (UE) 2024/1689 (c.d. AI Act) aspira a una posizione di centralità sistematica nel quadro regolatorio, dacché introduce una disciplina orizzontale dei sistemi di intelligenza artificiale fondata su un approccio di regolazione graduata in base al rischio. L'architettura *risk-based* del Regolamento è finalizzata a coniugare l'obiettivo di promozione dell'innovazione tecnologica con l'esigenza di assicurare un livello elevato di protezione dei diritti e delle libertà fondamentali.

In questo contesto, i sistemi di IA destinati alla valutazione del rischio e alla determinazione dei prezzi nelle assicurazioni sulla vita e sanitarie sono espressamente qualificati come sistemi ad alto rischio, in ragione del loro potenziale impatto sul sostentamento, sulla salute e sulla dignità delle persone interessate. Tale qualificazione comporta l'imposizione di un complesso apparato di obblighi in capo alle imprese assicurative in specie quando esse operino in qualità di *deployer* nonché di *user*. Tali obblighi investono la progettazione e l'utilizzo dei sistemi, la gestione e la mitigazione dei rischi, la qualità, pertinenza e rappresentatività dei dati, la trasparenza dei processi decisionali, l'effettività della supervisione umana e lo svolgimento di valutazioni di impatto sui diritti fondamentali.

Accanto all'AI Act, il Regolamento (UE) 2016/679 (GDPR) continua a svolgere una funzione centrale e si pone in rapporto di complementarietà di fatto con l'AI Act. Sebbene il GDPR non contenga un'esplicita disciplina dei sistemi di intelligenza artificiale, esso fornisce un modello di *data governance* che assume particolare rilevanza quando il trattamento dei dati personali è realizzato mediante processi algoritmici e automatizzati. In tale ambito, l'articolo 22 GDPR introduce limiti significativi al ricorso a decisioni basate unicamente su trattamenti automatizzati idonei a produrre effetti giuridici o a incidere in modo analogo significativamente sulla persona e sui relativi diritti. Il presente contributo evidenzia come, nel settore assicurativo, le pratiche di profilazione algoritmica e di *pricing* automatizzato possano rientrare pienamente nell'ambito applicativo di tale

disposizione, imponendo il riconoscimento di specifiche garanzie in favore dell'interessato, tra cui il diritto all'intervento umano effettivo, alla contestazione della decisione e a un adeguato livello di intelligibilità della logica sottesa al processo decisionale.

Particolare attenzione è dedicata alla nozione di dato personale e, in specie, di dato relativo alla salute. L'interpretazione estensiva di tale nozione, adottata dalle autorità europee, conduce infatti a ricomprendere nel perimetro applicativo del GDPR anche le informazioni inferite dai sistemi di IA, ossia gli *output* algoritmici che esprimono valutazioni probabilistiche sullo stato di salute o sui rischi futuri di una persona.

Tale impostazione rafforza in modo significativo gli obblighi di trasparenza e di *accountability* e rende particolarmente problematico il ricorso a *big data* e a fonti informative eterogenee, che possono fungere da vettori di discriminazione indiretta e da *proxy* di caratteristiche protette difficilmente rilevabili *ex ante*.

Su queste basi normative si innesta la seconda parte dell'analisi, dedicata all'esame del rischio di discriminazione connesso all'impiego dell'intelligenza artificiale nell'attività assicurativa. L'automazione dei processi decisionali, soprattutto quando fondata su modelli complessi, opachi o scarsamente spiegabili, è suscettibile di amplificare *bias* preesistenti nei dataset di addestramento o di generare nuove forme di esclusione sistemica. Tali rischi non si esauriscono nella fase di determinazione del premio, ma investono l'intera catena del valore assicurativa, dalla progettazione e dal posizionamento dei prodotti, alle strategie di *marketing* e distribuzione, fino al rinnovo delle polizze, alla gestione dei sinistri e alle attività di *fraud detection*. In particolare, l'iper-personalizzazione resa possibile dall'IA può condurre a una segmentazione estrema dei mercati di rischio, con effetti potenzialmente incompatibili con la funzione sociale dell'assicurazione e con il rischio di esclusione di soggetti caratterizzati da profili considerati eccessivamente onerosi, come frequentemente avviene nel settore delle polizze *long-term care*, connotate da una marcata dimensione assistenziale.

Il contributo evidenzia come il principio di non discriminazione, pur non sempre declinato in modo puntuale con riferimento a ogni singolo fattore di rischio assicurativo, permei l'intero ordinamento europeo e nazionale e trovi fondamento tanto nelle normative antidiscriminatorie quanto nei principi di correttezza, buona fede e trasparenza che informano i rapporti contrattuali. Un utilizzo non conforme dell'IA espone pertanto le imprese assicurative non solo a violazioni della normativa applicabile, ma anche a rilevanti rischi reputazionali, di contenzioso e di condotta, con possibili ricadute sulla sana e prudente gestione e sulla sostenibilità complessiva del modello di business.

In tale prospettiva assumono rilievo centrale le buone pratiche promosse dalle autorità europee, in particolare da EIOPA, che sollecitano l'adozione di un approccio sistemico alla governance dell'intelligenza artificiale. Il contributo sottolinea l'importanza di assetti organizzativi idonei a garantire trasparenza e spiegabilità dei modelli, un controllo rigoroso sulla qualità dei dati e sugli algoritmi, forme di supervisione umana sostanziali e una chiara allocazione delle responsabilità all'interno dell'impresa. Le valutazioni

di impatto, già note nel contesto del GDPR attraverso il DPIA, risultano ulteriormente rafforzate dall'AI Act mediante l'introduzione della valutazione di impatto sui diritti fondamentali, quale strumento chiave per l'individuazione preventiva e la mitigazione dei rischi più significativi.

La parte finale del lavoro è dedicata all'analisi delle polizze long-term care e all'impiego di tecnologie emergenti, quali i sistemi di intelligenza artificiale e i wearable devices. Questi ultimi, pur non qualificabili come sistemi di IA in senso proprio, costituiscono una fonte continuativa di dati biometrici e comportamentali idonei ad alimentare modelli predittivi avanzati. Le potenziali esternalità positive di tali strumenti sono rilevanti, in termini di rafforzamento delle politiche di prevenzione, maggiore tempestività degli interventi assistenziali, promozione di comportamenti salutari e riduzione della pressione sui sistemi sanitari pubblici. Tuttavia, l'analisi empirica evidenzia come, allo stato attuale, né nel mercato italiano né in quello europeo ed extraeuropeo risultino diffusi prodotti long-term care che integrino in modo strutturale l'IA e i wearable data nella determinazione del premio e nella gestione complessiva del rischio assicurato. Le esperienze straniere esaminate si collocano prevalentemente nell'ambito di modelli service-based, orientati all'erogazione di prestazioni sanitarie e assistenziali, senza configurare ancora schemi assicurativi LTC pienamente algoritmici.

In conclusione, il contributo propone una lettura problematizzata e non riduttiva dell'impiego dell'intelligenza artificiale nel settore assicurativo. Se l'IA rappresenta una leva potenzialmente decisiva per l'innovazione del mercato e per la sostenibilità dei sistemi di welfare, il suo utilizzo, soprattutto in ambiti ad alta sensibilità sociale come le polizze long-term care, richiede una progettazione normativa, organizzativa ed etica particolarmente accurata. Solo un'integrazione responsabile delle tecnologie, fondata sulla tutela dei diritti fondamentali e sull'inclusione dei soggetti vulnerabili, appare idonea a coniugare efficienza economica, compliance regolatoria e benessere sociale.

1. Il contesto normativo di riferimento

1.1 Un quadro di sintesi del Regolamento UE 2024/1689 (AI ACT)

Un corretto approccio al tema dell'utilizzo dei sistemi IA (*Intelligenza Artificiale*) per la valutazione dei rischi e la determinazione dei prezzi per assicurazioni sulla vita e assicurazioni sanitarie destinate a persone fisiche deve anzitutto muovere dalla considerazione del quadro normativo inteso a regolare i sistemi di IA in via generale, prima di – e a prescindere da – i relativi ambiti di utilizzo pratico. Rispetto a tale normativa a vocazione generale, le peculiarità che esprime il settore assicurativo integrano specifici elementi di diversificazione, che vengono di seguito analizzati.

Nell'accostarsi alla disciplina regolatrice dei sistemi IA, va anzitutto osservato che il progressivo sviluppo di tali sistemi richiede un quadro normativo armonizzato che ne disciplini l'utilizzo e che bilanci l'esigenza di non frenare lo sviluppo tecnologico in ambito assicurativo con la necessità di mantenere un adeguato livello di tutela dei diritti della persona: interessi, questi ultimi, che possono risultare finanche confliggenti nel

concreto.

È con l'intento di ricercare questo non facile equilibrio tra tali interessi che si è strutturato il frastagliato – e talvolta lacunoso – quadro normativo europeo di riferimento, il quale si caratterizza per un reciproco rinvio – non sempre organico – tra il Regolamento sull'Intelligenza Artificiale (Regolamento UE 2024/1689, di seguito: *AI Act* o il Regolamento), che costituisce il punto focale della materia che ci occupa, e la disciplina sulla *privacy* contenuta nella *General Data Protection Regulation* (Regolamento UE 2016/679, d'ora innanzi: GDPR).

La politica regolatoria prescelta dal legislatore europeo con l'*AI Act* è incentrata, in particolare, sui profili di governo dei numerosi rischi connessi all'utilizzo di modelli basati sull'IA secondo un approccio di «*compliance by design*» che si è sostanziato nell'introduzione di nuove regole in punto di produzione e fornitura dei sistemi di IA e di doveri di condotta in capo agli utilizzatori degli stessi.

Come esplicitato nel considerando n. 58 dell'*AI ACT*, «un settore in cui l'utilizzo dei sistemi di IA merita particolare attenzione è l'accesso ad alcuni servizi e prestazioni essenziali, pubblici e privati, necessari affinché le persone possano partecipare pienamente alla vita sociale o migliorare il proprio tenore di vita, e la fruizione di tali servizi. In particolare, le persone fisiche che chiedono o ricevono prestazioni e servizi essenziali di assistenza pubblica dalle autorità pubbliche, vale a dire servizi sanitari, prestazioni di sicurezza sociale, servizi sociali che forniscono protezione in casi quali la maternità, la malattia, gli infortuni sul lavoro, la dipendenza o la vecchiaia e la perdita di occupazione e l'assistenza sociale e abitativa, sono di norma dipendenti da tali prestazioni e servizi e si trovano generalmente in una posizione vulnerabile rispetto alle autorità responsabili. I sistemi di IA, se utilizzati per determinare se tali prestazioni e servizi dovrebbero essere concessi, negati, ridotti, revocati o recuperati dalle autorità, compreso se i beneficiari hanno legittimamente diritto a tali prestazioni o servizi, possono avere un impatto significativo sul sostentamento delle persone e violare i loro diritti fondamentali, quali il diritto alla protezione sociale, alla non discriminazione, alla dignità umana o a un ricorso effettivo e dovrebbero pertanto essere classificati come sistemi ad alto rischio».

L'approccio adottato dall'*AI Act* risponde a una logica *risk-based*. In particolare, esso prevede una regolamentazione diversificata a seconda del grado di rischio che l'utilizzo di IA presenta nello specifico settore di riferimento: nessuna previsione specifica si applica nel caso di sistemi con rischio minimo o nullo; per i sistemi «a basso rischio» il legislatore «raccomanda» (ma non impone) ai fornitori e agli utilizzatori l'adozione delle medesime cautele previste per i sistemi ad alto rischio tramite appositi codici di condotta (art. 69); mentre per i sistemi «ad alto rischio» il Regolamento prevede l'adozione di un articolato sistema di gestione dei rischi (artt. 6 ss.).

Nel dettaglio, l'*AI Act* intende escludere o minimizzare i numerosi rischi connessi all'utilizzo di modelli basati sull'IA seguendo fondamentalmente tre linee direttrici: (i) la fissazione di requisiti di *design* che il prodotto di IA deve rispettare in termini di trasparenza,

di qualità dei dati e di idoneità a consentire il c.d. «intervento umano» sul funzionamento; (ii) l'individuazione di obblighi in capo ai fornitori⁰¹ dei *software*; (iii) l'introduzione di doveri di condotta (art. 26) in capo agli utilizzatori degli stessi, ossia, nel caso che ci occupa, all'impresa assicurativa che operi in qualità di *deployer* (art. 3, n. 4), ossia quella «persona fisica o giuridica, un'autorità pubblica, un'agenzia o un altro organismo che utilizza un sistema di IA sotto la propria autorità, tranne nel caso in cui il sistema di IA sia utilizzato nel corso di un'attività personale non professionale».

L'AI Act classifica «i sistemi di IA destinati a essere utilizzati per la valutazione dei rischi e la determinazione dei prezzi in relazione a persone fisiche nel caso di assicurazioni sulla vita e assicurazioni sanitarie» (allegato III, sez. 5, lett. c) quali sistemi «ad alto rischio» (art. 6, § 2), potendo gli stessi «avere un impatto significativo sul sostentamento delle persone e, se non debitamente progettati, sviluppati e utilizzati, (...) violare i loro diritti fondamentali e comportare gravi conseguenze per la vita e la salute delle persone, tra cui l'esclusione finanziaria e la discriminazione» (ancora, il considerando n. 58).

Inoltre, l'uso frequente da parte di tali sistemi di IA di *Big Data* e di informazioni ricavabili dai *social network*, connesso con alla capacità di auto-apprendimento (*machine learning*) di alcune forme di IA solleva notevoli interrogativi sul piano dell'esigenza di garantire la protezione dei dati personali e la non discriminazione.

Per tali sistemi «ad alto rischio» il Regolamento prevede l'attuazione, documentazione e mantenimento di un articolato sistema di gestione dei rischi. A mente dell'articolo 9, tale sistema di gestione è da intendersi «come un processo iterativo continuo pianificato ed eseguito nel corso dell'intero ciclo di vita di un sistema di IA ad alto rischio, che richiede un riesame e un aggiornamento costanti e sistematici. Esso comprende le fasi seguenti: a) identificazione e analisi dei rischi noti e ragionevolmente prevedibili che il sistema di IA ad alto rischio può porre per la salute, la sicurezza e i diritti fondamentali quando il sistema di IA ad alto rischio è utilizzato conformemente alla sua finalità prevista; b) stima e valutazione dei rischi che possono emergere quando il sistema di IA ad alto rischio è usato conformemente alla sua finalità prevista e in condizioni di uso improprio ragionevolmente prevedibile; c) valutazione di altri eventuali rischi derivanti dall'analisi dei dati raccolti dal sistema di monitoraggio successivo all'immissione sul mercato di cui all'articolo 72; d) adozione di misure di gestione dei rischi opportune e mirate intese ad affrontare i rischi individuati ai sensi della lettera a)».

Nel dettaglio, ai sensi dell'art. 26 dell'AI Act, il *deployer* di sistemi di IA «ad alto rischio» – nel caso in analisi, l'impresa assicurativa – è obbligato a: (i) adottare idonee misure tecniche e organizzative per garantire l'utilizzo dei sistemi conformemente alle istruzioni

⁰¹ Ai sensi dell'art. 3, n. 3 AI Act, per «fornitore» s'intende «una persona fisica o giuridica, un'autorità pubblica, un'agenzia o un altro organismo che sviluppa un sistema di IA o un modello di IA per finalità generali o che fa sviluppare un sistema di IA o un modello di IA per finalità generali e immette tale sistema o modello sul mercato o mette in servizio il sistema di IA con il proprio nome o marchio, a titolo oneroso o gratuito».

per l'uso che li accompagnano; (ii) affidare la sorveglianza umana a persone fisiche che dispongono della competenza, della formazione, dell'autorità e del sostegno necessari; (iii) nella misura in cui il *deployer* eserciti il controllo sui dati di *input*, garantire essi siano pertinenti e sufficientemente rappresentativi alla luce della finalità prevista del sistema di IA ad alto rischio; (iv) monitorare il funzionamento del sistema di IA ad alto rischio sulla base delle istruzioni per l'uso e, qualora abbiano motivo di ritenere che l'uso del sistema di IA ad alto rischio in conformità delle istruzioni possa comportare che il sistema di IA presenti un rischio per la salute o la sicurezza o per i diritti fondamentali delle persone (di cui all'art. 79); (v) conservare i *log* generati automaticamente dal sistema di IA ad alto rischio, nella misura in cui siano sotto il loro controllo del *deployer*, per un periodo adeguato alla prevista finalità del sistema di IA ad alto rischio, di almeno sei mesi; (vi) informare i rappresentanti dei lavoratori e i lavoratori interessati che saranno soggetti all'uso del sistema di IA ad alto rischio prima di mettere in servizio o utilizzare un sistema di IA ad alto rischio sul luogo di lavoro.

A ciò si aggiunge il disposto dell'art. 4 (*Alfabetizzazione in materia di IA*) il quale dispone a monte che « i fornitori e i *deployer* dei sistemi di IA adottano misure per garantire nella misura del possibile un livello sufficiente di alfabetizzazione in materia di IA del loro personale nonché di qualsiasi altra persona che si occupa del funzionamento e dell'utilizzo dei sistemi di IA per loro conto, prendendo in considerazione le loro conoscenze tecniche, la loro esperienza, istruzione e formazione, nonché il contesto in cui i sistemi di IA devono essere utilizzati, e tenendo conto delle persone o dei gruppi di persone su cui i sistemi di IA devono essere utilizzati ».

Accanto agli obblighi di cui all'art. 26, il successivo articolo 27 impone l'obbligo di valutazione dell'impatto sui diritti fondamentali che l'uso dei sistemi di IA «ad alto rischio» può produrre.

Tale obbligo si sostanzia nei seguenti elementi: (i) una descrizione dei processi del *deployer* in cui il sistema di IA ad alto rischio sarà utilizzato in linea con la sua finalità prevista; (ii) una descrizione del periodo di tempo entro il quale ciascun sistema di IA ad alto rischio è destinato a essere utilizzato e con che frequenza; (iii) le categorie di persone fisiche e gruppi verosimilmente interessati dal suo uso nel contesto specifico; (iv) i rischi specifici di danno che possono incidere sulle categorie di persone fisiche o sui gruppi di persone individuati, tenendo conto delle informazioni trasmesse dal fornitore a norma dell'articolo 13; (v) una descrizione dell'attuazione delle misure di sorveglianza umana, secondo le istruzioni per l'uso; (vi) le misure da adottare qualora tali rischi si concretizzino, comprese le disposizioni relative alla governance interna e ai meccanismi di reclamo.

Accanto alla disciplina sopra esaminata dettata per i sistemi di IA qualificati come «ad alto rischio», il legislatore europeo si premura di fissare le pratiche di IA vietate, cui è dedicato il Capo II dell'*AI Act*.

In particolare, ai fini dell'indagine che ci occupa, ai sensi dell'art. 5 lett. c) e f) sono vieta-

te le pratiche di IA quali: «l'immissione sul mercato, la messa in servizio o l'uso di sistemi di IA per la valutazione o la classificazione delle persone fisiche o di gruppi di persone per un determinato periodo di tempo sulla base del loro comportamento sociale o di caratteristiche personali o della personalità note, inferite o previste, in cui il punteggio sociale così ottenuto comporti il verificarsi di uno o di entrambi gli scenari seguenti: i) un trattamento pregiudizievole o sfavorevole di determinate persone fisiche o di gruppi di persone in contesti sociali che non sono collegati ai contesti in cui i dati sono stati originariamente generati o raccolti; ii) un trattamento pregiudizievole o sfavorevole di determinate persone fisiche o di gruppi di persone che sia ingiustificato o sproporzionato rispetto al loro comportamento sociale o alla sua gravità» (lett. c.); nonché «l'immissione sul mercato, la messa in servizio per tale finalità specifica o l'uso di sistemi di IA per inferire le emozioni di una persona fisica nell'ambito del luogo di lavoro e degli istituti di istruzione, tranne laddove l'uso del sistema di IA sia destinato a essere messo in funzione o immesso sul mercato per motivi medici o di sicurezza» (lett. f.).

Dal dettato normativo sembrerebbe desumersi che il tratto differenziale tra il lecito e corretto utilizzo dei sistemi di IA ad alto rischio (art. 6) e i sistemi di IA vietati (art. 5) vada individuato alla luce della concreta finalità con cui gli stessi sono utilizzati, richiedendosi pertanto alla preposta autorità nazionale di vigilanza un'indagine in concreto in tal senso.

Fermo il merito di aver introdotto dei criteri per la commerciabilità e l'utilizzabilità dei sistemi di IA, sussistono una serie di problemi di coordinamento con altre discipline settoriali.

Segnatamente, vengono in rilievo il poco chiaro coordinamento con la disciplina di settore (inclusa quella delineata dal regolamento (UE) 2022/2554 (DORA) nel caso di sistemi di IA prodotti internamente dalle imprese assicurative, nonché l'assenza di regole per l'ipotesi di *outsourcing* della fase di profilazione del cliente svolta con IA che, se affidata a soggetti non vigilati, sfugge al controllo dell'IVASS. Inoltre, il Regolamento risulta apparentemente incompleto tanto nel definire le condizioni di ammissibilità di un trattamento che sia completamente automatizzato, quanto nel fornire una risposta sull'ammissibilità delle tipologie alternative di dati utilizzabili e delineare validi strumenti diretti per la tutela dei diritti fondamentali.

Al fine di colmare tali carenze, si rende necessaria una riflessione su quali fonti normative vigenti consentano di bilanciare l'esigenza di non frenare lo sviluppo e l'utilizzo delle tecnologie emergenti in ambito assicurativo con la necessità di mantenere un adeguato livello di protezione di coloro che, a vario titolo, potrebbero risultare da queste condizionati.

A tal proposito, vengono in rilievo i principi sanciti dal GDPR (in particolare l'art. 22) di

cui lo stesso AI Act lascia impregiudicata l'applicazione (considerando n. 10⁰²).

Ai sensi dell'art. 22 GDPR («*Processo decisionale automatizzato relativo alle persone fisiche, compresa la profilazione*»), risulta precluso sottoporre l'interessato a una decisione basata «unicamente» su trattamento automatizzato, compresa la profilazione, se tale decisione «produc[e] effetti giuridici che lo riguardano o che incida in modo analogo significativamente sulla sua persona». Al riguardo, il Gruppo di lavoro Articolo 29 ha chiarito che la disposizione sancisce un vero e proprio divieto generale di qualsiasi decisione basata su un trattamento interamente automatizzato di dati personali che abbia un impatto giuridico o materiale su una persona fisica⁰³.

Benché generale, tale preclusione non ha comunque carattere assoluto e, anzi, risulta fortemente limitata⁰⁴, posto che la norma prevede una serie di eccezioni in cui la decisione integralmente automatizzata è infatti consentita⁰⁵, purché rientri in una delle ipotesi previste al § 2. Tali ipotesi comprendono il caso in cui la decisione: a) sia necessaria per la conclusione o l'esecuzione di un contratto tra l'interessato e un titolare del trattamento; b) sia autorizzata dal diritto dell'Unione o dello Stato membro cui è soggetto il titolare del trattamento, che precisa altresì misure adeguate a tutela dei diritti, delle libertà e dei legittimi interessi dell'interessato; c) si basi sul consenso esplicito dell'interessato.

Al ricorrere di uno dei suddetti presupposti anche il trattamento completamente au-

02 Considerando 10 dell'AI Act "Il presente regolamento non mira a pregiudicare l'applicazione del vigente diritto dell'Unione che disciplina il trattamento dei dati personali, inclusi i compiti e i poteri delle autorità di controllo indipendenti competenti a monitorare la conformità con tali strumenti. Inoltre, lascia impregiudicati gli obblighi dei fornitori e dei *deployer* dei sistemi di IA nel loro ruolo di titolari del trattamento o responsabili del trattamento derivanti dal diritto dell'Unione o nazionale in materia di protezione dei dati personali, nella misura in cui la progettazione, lo sviluppo o l'uso di sistemi di IA comportino il trattamento di dati personali".

03 Gruppo di lavoro Articolo 29 per la protezione dei dati, *Linee guida sul processo decisionale automatizzato relativo alle persone fisiche e sulla profilazione ai fini del regolamento 2016/679*, 3 ottobre 2017, in cui si specifica che «Il termine "diritto" contenuto nella disposizione non significa che l'articolo 22, paragrafo 1, si applica soltanto se invocato attivamente dall'interessato. L'articolo 22, paragrafo 1, stabilisce un divieto generale nei confronti del processo decisionale basato unicamente sul trattamento automatizzato. Tale divieto si applica indipendentemente dal fatto che l'interessato intraprenda un'azione in merito al trattamento dei propri dati personali» (p. 21). Per l'opinione favorevole al divieto generale in dottrina, E. GIORGINI, *Profilazione automatizzata e contratto assicurativo*, cit.; G. CARAPEZZA FIGLIA, *Decisioni algoritmiche tra diritto alla spiegazione e divieto di discriminare*, in S. Orlando (a cura di), *Libertà e liceità del consenso nel trattamento dei dati personali*, Firenze, 2024; E. PELLECCIA, *Profilazione e decisioni automatizzate al tempo della black box society: qualità dei dati e leggibilità dell'algoritmo nella cornice della responsible research and innovation*, in *Nuove leggi civ. comm.*, 5, 41, 2018; A. SPANGARO, *Il concetto di profilazione tra "direttiva madre" e GDPR*, in *Giur. it.*, 2022.

04 E. GIORGINI, *Profilazione automatizzata e contratto assicurativo*, cit.

05 Inoltre, laddove il processo decisionale coinvolga categorie particolari di dati definite all'articolo 9, paragrafo 1, il titolare del trattamento deve altresì garantire di poter soddisfare il requisito del consenso esplicito (richiamato dall'art. 22, para. 4).

tomatizzato è ammesso, subordinatamente al riconoscimento da parte del titolare di talune garanzie a favore dell'interessato, al fine di non lasciarlo esposto ad eventuali errori o distorsioni del sistema automatico che possono riassumersi nella possibilità di richiedere «l'intervento umano» e la *disclosure* del procedimento.

1.2 GDPR e governance dei dati

Il riferimento al processo decisionale automatizzato consente, invero, di effettuare qualche considerazione ulteriore sulla rilevanza della disciplina tracciata dal GDPR nel presente contesto, con particolare riferimento ad alcuni principi ispiratori del governo dei dati laddove siano processati da modelli di intelligenza artificiale.

L'AI Act lascia impregiudicato quanto stabilito dal GDPR proprio perché complementare ad esso, in quanto volto a rafforzare la tutela dei dati personali nel caso in cui si faccia uso di sistemi di IA⁰⁶.

Sebbene il GDPR non contenga alcun riferimento ai sistemi di intelligenza artificiale (né a sistemi di inferenza automatica o ai *big data*)⁰⁷, il modello della *governance* dei dati personali fornisce uno standard organizzativo e di condotta che deve essere predisposto anche nella gestione dei sistemi di IA laddove venga in rilievo il trattamento di dati personali.

Nei paragrafi che seguono, saranno analizzate le disposizioni del GDPR più rilevanti in materia di IA⁰⁸, che impongono una declinazione degli obblighi del responsabile del trattamento più incisiva proprio nelle ipotesi di trattamento algoritmico dei dati, alla luce dei principi generali che regolano la materia della tutela dei dati personali (art. 5 GDPR)⁰⁹. In particolare saranno analizzati: (i) la nozione di dato personale con riguardo alle inferenze automatizzate; (ii) il contenuto degli obblighi di trasparenza nei processi decisionali automatizzati; (iii) il sistema di governance dei dati personali.

1.2.1 Nozione di dato personale

In base all'art. 4, n. 1, GDPR, costituisce dato personale «qualsiasi informazione riguardante una persona fisica identificata o identificabile». Con riguardo ai dati relativi alla salute, l'art. 9 richiede il consenso esplicito dell'interessato. Le linee guida europee, in

⁰⁶ Il Considerando 10 dell' AI Act stabilisce che il regolamento «lascia impregiudicati gli obblighi dei fornitori e dei *deployer* dei sistemi di IA nel loro ruolo di titolari del trattamento o responsabili del trattamento derivanti dal diritto dell'Unione o nazionale in materia di protezione dei dati personali, nella misura in cui la progettazione, lo sviluppo o l'uso di sistemi di IA comportino il trattamento di dati personali».

⁰⁷ Sull'insufficienza del GDPR per le attività di *Big Data analysis*, T. ZARSKY, *Incompatible: the GDPR in the age of big data*, in *Seton Hall L. Rev.*, 2017.

⁰⁸ European Parliamentary Research Service, *The impact of the General Data Protection Regulation (GDPR) on artificial intelligence*, 2020.

⁰⁹ European Data Protection Board, *Opinion 28/2024 on certain data protection aspects related to the processing of personal data in the context of AI models*.

accordo con la prevalente dottrina, convergono verso l'adozione di un'interpretazione ampia del concetto di dato personale, al fine di garantire un'estesa applicazione delle tutele predisposte del GDPR.

Per ciò che rileva in questa sede, deve essere tenuto in considerazione che i sistemi di IA possono dedurre nuove informazioni sugli interessati applicando modelli algoritmici ai loro dati personali. Il Gruppo di lavoro Articolo 29 ha chiarito che la nozione di dato personale deve essere interpretata come concetto ampio, sicché le informazioni inferite dai sistemi di IA *devono essere considerate dati personali ulteriori e diversi* rispetto ai dati originariamente raccolti presso l'interessato¹⁰. Ciò, come si vedrà, rileva nell'ipotesi di inferenza automatizzata poiché il perimetro qualificatorio di dato personale definisce i termini di applicazione degli obblighi di trasparenza (artt. 13 e 14 GDPR) e del diritto di accesso degli interessati (art. 15 GDPR).

Adottando questa lettura, l'interessato avrà diritto di accedere non soltanto ai dati di *input*, ma anche agli *output* dedotti dall'algoritmo utilizzato, siano essi *output* intermedi o finali nell'ambito del processo di trattamento. Con specifico riguardo al processo di dati nel settore assicurativo, è stato affermato che i dati desunti per condurre valutazioni o adottare decisioni devono sicuramente essere considerati come dati personali¹¹.

A tale conclusione si perviene anche tenendo in considerazione la nozione ampia generalmente attribuita ai "*dati relativi alla salute*"¹². Le informazioni relative alla salute ottenute in via indiretta occupano un'area grigia, che si estende ancora di più nell'era dei *big data* e della digitalizzazione dell'*health care*. Il Gruppo di lavoro Articolo 29 (il «Gruppo di Lavoro») ha adottato un criterio di qualificazione dinamico, che guarda non a una categorizzazione oggettiva dei dati ma alla *finalità* dell'attività di trattamento per cui sono stati raccolti. Con riguardo ai dati sanitari, si è stabilito che essi devono essere considerati relativi alla salute in ogni caso in cui «si traggono conclusioni sullo stato di salute o sui rischi per la salute di una persona (indipendentemente dal fatto che tali conclusioni siano accurate o imprecise, legittime o illegittime, o comunque adeguate o inadeguate)»¹³. Ciò che rileva dunque sono le conclusioni probabilistiche sullo stato di salute e sui rischi per la salute degli interessati che si traggono da altri dati relativi all'interessato.

Inoltre, nel trattamento di dati personali che fanno ricorso all'analisi di *big data*, cioè di quei dati *non personali* che le imprese assicuratrici combinano con i dati personali al

¹⁰ Gruppo di lavoro articolo 29, *Parere 4/2007 sul concetto di dati personali*, WP 136.

¹¹ E. GIORGINI, *Profilazione automatizzata e contratto assicurativo*, in *Ass.*, 4, 2024.

¹² Per una ricostruzione della problematica, W. SCHÄFKE-ZELL, *Revisiting the definition of health data in the age of digitalized health care*, in *Int. data protection I.*, 1, 12, 2022, dove si argomenta, a titolo esemplificativo, che può essere qualificato come dato personale relativo alla salute anche il *download* di una determinata applicazione sul telefono cellulare.

¹³ Article 29 Working Party, *Annex - health data in apps and devices, Advice paper on special categories of data ("sensitive data")*, aprile 2011.

fine di ricavare ulteriori informazioni, sussiste il rischio che dati qualificati come “non personali”, e in quanto tali sottratti all’applicazione del Regolamento, siano in realtà in grado di fornire informazioni su un determinato gruppo di individui. Tale osservazione viene in rilievo poiché dalla natura (personale o anonima) del dato dipende l’applicazione della normativa in materia di protezione dei dati personali. La corretta distinzione, da parte del titolare del trattamento, tra dato e informazione è dunque fondamentale al fine di non incorrere nella violazione della disciplina del GDPR¹⁴.

1.2.2 Principio di trasparenza e processi decisionali automatizzati (ADM)

L’articolo 5, paragrafo 1, lettera a), del GDPR richiede che i dati personali siano trattati in modo «lecito, equo e trasparente nei confronti dell’interessato»¹⁵. Il concetto di trasparenza è inoltre specificato nel considerando 58, secondo cui le informazioni rivolte all’interessato devono essere «concise, facilmente accessibili e comprensibili» e comunicate attraverso un linguaggio «chiaro e semplice», nonché dal considerando 60, in base al quale «[i] principi di trattamento corretto e trasparente implicano che l’interessato sia informato dell’esistenza del trattamento e delle sue finalità. Il titolare del trattamento dovrebbe fornire all’interessato eventuali ulteriori informazioni necessarie ad assicurare un trattamento corretto e trasparente, prendendo in considerazione le circostanze e del contesto specifici in cui i dati personali sono trattati».

La definizione del contenuto dell’obbligo di trasparenza nei casi di utilizzo di sistemi di IA pone una serie di problematiche ulteriori rispetto alla semplice *disclosure* sulle caratteristiche e le finalità del trattamento dei dati personali da parte di persone fisiche. La maggior parte dei sistemi di intelligenza artificiale non è infatti dotata di meccanismi informatici “interni”, in grado di fornire all’utente la spiegazione della logica sottesa alla produzione di un determinato *output*, generando così una “*black-box*” cognitiva¹⁶.

Tuttavia, è stato osservato che il GDPR non declina l’obbligo di trasparenza con specifico riguardo allo scopo, le capacità e le modalità operative (“*spiegabilità*”) dei sistemi automatizzati di trattamento dei dati in via generale, posto che la correttezza informativa imposta dal GDPR richiede che gli interessati non siano ingannati o fuorviati in merito

¹⁴ E. GIORGINI, *Profilazione automatizzata e contratto assicurativo*, cit.

¹⁵ Tale obbligo si aggiunge e specifica obblighi derivanti dalle disposizioni che regolano l’attività assicurativa, quali la Direttiva (UE) 2016/97 (IDD), che impone ai distributori di prodotti assicurativi di agire in modo onesto, imparziale e professionale, diligente e trasparente (art. 17), il Cod. ass. priv. (art. 183) e il Regolamento Ivass n. 41/2018, che impongono doveri di trasparenza in capo all’assicuratore.

¹⁶ Per il problema dal punto di vista informatico, tra gli altri, R. GUIDOTTI, A. MONREALE, S. RUGGIERI, D. PEDRESCHI, F. TURINI, F. GIANNOTTI, *Local Rule-Based Explanations of Black Box Decision Systems*, 2018, disponibile su <https://arxiv.org/pdf/1805.10820>.

al trattamento dei loro dati¹⁷. In base alla lettera delle disposizioni del GDPR, l'obbligo di fornire una spiegazione della logica sottesa al sistema di IA utilizzato viene in rilievo soltanto con riferimento ai processi decisionali interamente automatizzati, inclusa la profilazione¹⁸.

L'art. 4, para. 1, n. 4 del GDPR definisce la profilazione come «qualsiasi forma di trattamento automatizzato di dati personali consistente nell'utilizzo di tali dati personali per valutare determinati aspetti personali relativi a una persona fisica, in particolare per analizzare o prevedere aspetti riguardanti il rendimento professionale, la situazione economica, la salute, le preferenze personali, gli interessi, l'affidabilità, il comportamento, l'ubicazione o gli spostamenti di detta persona fisica».

A tal riguardo, è scontato osservare che, in effetti, l'attività profilazione è uno degli utilizzi più diffusi di sistemi IA per il trattamento dei dati personali. Nel settore assicurativo, la principale finalità della profilazione automatizzata è di minimizzare i rischi per le compagnie assicurative, attraverso una (significativamente) maggiore personalizzazione dell'offerta di polizza in base alle specifiche caratteristiche del soggetto interessato.

Come ricordato in precedenza (v. *supra*, p. 5), ai sensi dell'art. 22 GDPR «L'interessato ha il diritto di non essere sottoposto a una decisione basata unicamente sul trattamento automatizzato, compresa la profilazione, che produca effetti giuridici che lo riguardano o che incida in modo analogo significativamente sulla sua persona».

Le condizioni di liceità del ricorso al processo decisionale automatizzato relativo alle persone fisiche sono state infine declinate dalle pertinenti linee guida del Gruppo di Lavoro (le «Linee Guida»).

In primo luogo, uno dei nodi interpretativi più complessi riguarda l'individuazione di quali procedimenti possano definirsi completamente automatizzati, posto che rimangono fuori dall'ambito di applicazione i procedimenti semi-automatizzati, in cui la decisione finale è comunque assunta da una persona fisica. A tal riguardo, le Linee Guida hanno precisato che «[i]l titolare del trattamento non può eludere le disposizioni dell'articolo 22 creando coinvolgimenti umani fittizi. Ad esempio, se qualcuno applica abitualmente profili generati automaticamente a persone fisiche senza avere alcuna influenza effettiva sul risultato, si tratterà comunque di una decisione basata unicamente sul trattamento automatico. Per aversi un coinvolgimento umano, il titolare del trattamento deve garantire che qualsiasi controllo della decisione sia significativo e

¹⁷ G. COMANDÉ, *Leggibilità algoritmica e consenso al trattamento dei dati personali, note a margine di recenti provvedimenti sui dati personali*, in *Danno resp.*, 2, 2022, il quale ha evidenziato che la trasparenza richiesta dal GDPR sia «correlata, ma distinta, dall'idea di un'IA trasparente e spiegabile». Dunque, il concetto di spiegabilità dell'IA attiene propriamente alla costruzione del modello informatico del funzionamento di un sistema di IA e alla correttezza informazioni da fornire al pubblico.

¹⁸ Cfr. Art. 22 GDPR

non costituisca un semplice gesto simbolico. Il controllo dovrebbe essere effettuato da una persona che dispone dell'autorità e della competenza per modificare la decisione. Nel contesto dell'analisi, tale persona dovrebbe prendere in considerazione tutti i dati pertinenti. Nell'ambito della valutazione d'impatto sulla protezione dei dati, il titolare del trattamento dovrebbe individuare e registrare il grado di coinvolgimento umano nel processo decisionale e la fase nella quale quest'ultimo ha luogo»¹⁹. Più di recente, l'*European Data Protection Board* ha sottolineato che, per configurarsi una supervisione umana sui sistemi automatizzati idonea a qualificarli come parzialmente automatizzati sottrarli all'applicazione dei requisiti specifici stabiliti per i processi decisionali esclusivamente automatizzati, è fondamentale la formazione del personale che deve effettuare la revisione del processo decisionale²⁰.

Di maggiore interesse per l'utilizzo di profilazione automatica a base IA è la condizione che impone il carattere necessario del sistema automatizzato ai fini della conclusione o dell'esecuzione di un contratto. In tali ipotesi spetta al titolare del trattamento dimostrare la necessità del ricorso a modalità di trattamento connotate da automatizzazione, tenendo in dovuta considerazione le possibilità di adottare un metodo meno incisivo sulla riservatezza della persona fisica. Pertanto, se esistono altri mezzi efficaci e meno invasivi per il conseguimento del medesimo obiettivo, il trattamento non può essere considerato "necessario"²¹.

Quando ricorrono le ipotesi *a)* e *c)*, il paragrafo 3 richiede l'adozione di misure di salvaguardia adeguate, "almeno" per quanto riguarda il diritto dell'interessato di ottenere l'intervento umano da parte del titolare del trattamento, di esprimere la propria opinione e di contestare la decisione.

Le Linee Guida hanno chiarito quali sono gli obblighi in capo ai titolari del trattamento e ulteriori garanzie per gli interessati, tra cui: (i) chiedere agli istituti finanziari che un essere umano intervenga nella profilazione, (ii) esprimere il proprio punto di vista, (iii) contestare una decisione basata sulla profilazione; (iv) il diritto di accedere sia ai dati personali utilizzati come *input* per l'inferenza, sia ai dati personali ottenuti come risultato (finale o intermedio) dell'inferenza; (v) il diritto di rettificare le informazioni dedotte non solo quando le informazioni dedotte sono "verificabili" (i.e. la loro correttezza può essere oggettivamente determinata), ma anche quando sono il risultato di inferenze non verificabili o probabilistiche (ad esempio, la probabilità di sviluppare malattie cardiache in futuro).

¹⁹ *Linee guida sul processo decisionale automatizzato*, 23.

²⁰ EDPB-GEPD, *Parere congiunto 5/2021 sulla proposta di regolamento del Parlamento europeo e del Consiglio che stabilisce regole armonizzate sull'intelligenza artificiale (legge sull'intelligenza artificiale)*, 18 giugno 2021, 6.

²¹ Le *Linee guida sul processo decisionale automatizzato* hanno tuttavia chiarito, a titolo esemplificativo, che il processo decisionale automatizzato di cui all'articolo 22, paragrafo 1, può essere necessario anche per un trattamento precontrattuale, dove la necessità può essere anche data dal numero elevato di richieste.

Ai fini dell'individuazione delle tutele con specifico riguardo a sistemi AI, saranno analizzati (i) il diritto alla leggibilità della decisione algoritmica e (ii) il diritto all'intervento umano.

1.2.3 (segue) Diritto alla leggibilità della decisione algoritmica

Ulteriori obblighi informativi per il caso di processo decisionale automatizzato di cui all'art. 22 sono stabiliti dagli artt. 13, 14 e 15 del GDPR, da interpretarsi alla luce del considerando 71. In particolare:

(i) l'art. 13(par. 2, lett. f) e l'art. 14(par. 2, lett. g) stabiliscono il contenuto delle informazioni da fornire prima che i dati dell'interessato siano soggetti a trattamento (informazioni *ex ante*), prescrivendo che l'interessato ha il diritto di essere informato «dell'esistenza di un processo decisionale automatizzato, compresa la profilazione, di cui all'articolo 22, paragrafi 1 e 4», nonché del diritto di ricevere «informazioni significative sulla logica applicata, nonché sull'importanza e sulle conseguenze previste di tale trattamento per l'interessato»;

(ii) l'art. 15(par. 1, lett. h) definisce il contenuto delle informazioni *ex post*, cioè delle informazioni che devono essere fornite nell'ambito dell'esercizio da parte dell'interessato del diritto di accesso ai dati personali;

(iii) il considerando 71 richiede che le garanzie adeguate cui è subordinato il trattamento automatizzato devono comprendere «la specifica informazione all'interessato e il diritto di ottenere l'intervento umano, di esprimere la propria opinione, di ottenere una spiegazione della decisione conseguita dopo tale valutazione e di contestare la decisione».

Autorevole dottrina ha evidenziato come l'elenco delle informazioni dovute, previste dal GDPR, non contempla espressamente il diritto di ottenere una spiegazione *personalizzata*, che specifichi cioè i motivi per cui è stata adottata una decisione sfavorevole²². In aggiunta, non è richiesto al titolare del trattamento che la decisione sia *ragionevole*, cioè che «la misura in cui alcuni fattori di *input* influenzano la decisione sia proporzionata all'importanza causale o almeno predittiva di tali fattori rispetto agli obiettivi legittimi perseguiti»²³, sicché fattori di *input* e gli obiettivi siano accettabili dal soggetto interessato e il metodo possa essere considerato da questi affidabile.

Tuttavia, è orientamento prevalente che fra le misure di salvaguardia minime previste

22 G. COMANDÉ, *Leggibilità algoritmica e consenso al trattamento dei dati personali*, cit., Sui limiti delle tutele previste dall'art. 22 del GDPR, di recente, P. DE HERT, G. LAZCOZ, *Radical rewriting of Article 22 GDPR on machine decisions in the AI era*, 14 agosto 2024, disponibile su <https://www.europeanlawblog.eu/pub/radical-rewriting-of-article-22-gdpr-on-machine-decisions-in-the-ai-era/release/1>.

23 G. COMANDÉ, *Leggibilità algoritmica e consenso al trattamento dei dati personali*, cit.; G. COMANDÉ, G. MALGIERI, *Why a Right to Legibility of Automated Decision-Making Exists in the General Data Protection Regulation*, in *Int. Data Privacy L.*, 7, 4, 2017.

dal par. 3 dell'art. 22 debba includersi anche un «*diritto alla leggibilità*» della decisione algoritmica²⁴. Attraverso un'interpretazione sistematica degli artt. 13, 14 e 22 GDPR, nonché del considerando 71 che pone l'accento sul *contenuto* di un'inferenza o di una decisione automatizzata, può ritenersi che il titolare del trattamento debba conformarsi non soltanto a un certo standard di trasparenza ma anche a «standard di *accettabilità, pertinenza e affidabilità*»²⁵ della decisione. Le garanzie previste dall'art. 22, par. 3, GDPR implicherebbero, per la persona titolare dei dati personali, la possibilità di comprendere le logiche sottese al processo decisionale al fine di valutare la solidità metodologica, sicché il diritto di conoscere le ragioni della decisione costituisce «un pilastro del *due process algoritmico*»²⁶.

I destinatari devono essere in grado di comprendere le decisioni anche al fine di poterne contestare il contenuto e valutare l'opportunità, poiché la leggibilità della decisione si pone su un piano prodromico alla sua effettiva contestabilità²⁷. Tutto ciò è tanto più rilevante nel settore assicurativo, giacché il ricorso di sistemi di IA da parte dell'impresa assicuratrice può indurre, come si è detto, a fenomeni di discriminazione (su cui si vedano i par. 1.2.3 e ss.)

24 G. COMANDÉ, G. MALGIERI, *Why a Right to Legibility of Automated Decision-Making Exists in the General Data Protection Regulation*, cit.; R. MORTIER ET AL., *Human Data Interaction: The Human Face of the Data-Driven Society*, MIT Technology Review, 2014; G. MALGIERI, G. COMANDÉ, *Sensitive by Distance: Quasi Health Data in the Algorithmic Era*, in *Information and Communications Technology Law*, 26, 2017; G. COMANDÉ, *Leggibilità algoritmica e consenso al trattamento dei dati personali*, cit., secondo il quale «la leggibilità viene collegata alla trasparenza dell'algoritmo (di cui prima espressione operativa è certo l'informativa agli interessati), che deve essere tale da dare indicazioni comprensibili all'interessato sulle logiche con cui si inferiscono delle conclusioni e non indicazioni tecniche astruse». Per l'autore tale indirizzo è stato anche avallato dalla Corte di cassazione, che, nell'ordinanza 24-25 marzo 2021, n. 14381, ha fornito indirettamente una chiave di lettura degli artt. 13 e 14 GDPR che chiede di fornire all'interessato «informazioni significative sulla logica utilizzata, nonché l'importanza e le conseguenze previste di tale trattamento per l'interessato» per ogni tipo di trattamento automatizzato. La Corte avrebbe considerato la profilazione e i processi automatizzati come una mera esemplificazione dell'obbligo informativo, sicché gli artt. 13 e 14 GDPR assumerebbero una vera e propria portata generale, trovando applicazione anche al di fuori delle ipotesi previste dall'art. 22.

25 G. COMANDÉ, G. MALGIERI, cit.

26 G. COMANDÉ, G. MALGIERI, cit.; R. MESSINETTI, *La Privacy e il controllo dell'identità algoritmica*, in *Contr. impr./Eur.*, 2021.

27 H. FELZMANN, E. FOSCH-VILLARONGA, C. LUTZ, A. TAMÓ-LARRIEUX, *Transparency you can trust: transparency requirements for artificial intelligence between legal norms and contextual concerns*, in *Big Data & Society*, 2019; G. COMANDÉ, *Leggibilità algoritmica e consenso al trattamento dei dati personali*, cit., secondo cui «la stessa proposta disciplina sull'intelligenza artificiale richiede che le misure di sorveglianza umana (di cui l'informativa all'interessato non può non fare parte) deve permettere di «comprendere appieno le capacità e i limiti del sistema di IA ad alto rischio» (art. 14 GDPR, comma 4, lett. a) assieme ai vincoli di trasparenza specificati all'art. 13 GDPR così come la garanzia della qualità e tracciabilità dei dati e dell'addestramento (art. 10 GDPR)».

1.2.4 (segue) Diritto all'intervento umano

L'art. 22, par. 3, conferisce all'interessato, come misura di salvaguardia minima, il diritto di richiedere l'intervento umano nell'assunzione della decisione. Per i sistemi completamente automatizzati sono dunque individuabili due momenti in cui la supervisione umana può essere eseguita:

1. prima che il sistema di IA sia implementato, in conformità alle disposizioni dell'AI Act;
2. dopo che il titolare del trattamento ha assunto la decisione, in attuazione della misura di salvaguardia prevista dall'art. 22, par. 3, GDPR.

Quest'ultima attività di supervisione, che interviene in fase di revisione della decisione, è stata definita come «supervisione umana di secondo grado»²⁸, in quanto meramente correttiva (e non costitutiva) della produzione decisionale di un'IA.

L'intervento umano nel riesame della decisione deve essere effettivo, cioè effettuato da una persona che dispone dell'autorità, della competenza e della conoscenza adeguate a modificare in concreto la decisione, compiendo una valutazione approfondita di tutti i dati pertinenti, comprese eventuali informazioni aggiuntive fornite dall'interessato. In tale sede possono ritenersi valide le stesse osservazioni che il Gruppo di Lavoro ha svolto con riguardo alla qualifica di processo interamente automatizzato, ai fini dell'applicazione dell'art. 22²⁹.

Nell'ambito dei sistemi di IA, l'effettività dell'intervento umano si scontra con le reali possibilità da parte dell'operatore umano di incidere effettivamente sull'*output* decisionale: quest'ultimo, infatti, difficilmente può essere in grado di rivedere in maniera efficace una decisione automatizzata frutto di tecnologie da esso non totalmente ricostruibili nelle relative logiche di funzionamento. Ciò mostra come il diritto alla leggibilità e, dunque, all'intervento umano efficace, sia condizionato dall'intervento umano di primo grado, cioè da quello richiesto oggi nella fase della progettazione del sistema di intelligenza artificiale dell'art. 13 dell'AI Act³⁰.

Allo scopo di prevenire o ridurre al minimo i pericoli legati all'utilizzo dell'intelligenza artificiale, è previsto che i sistemi ad alto rischio siano progettati e sviluppati in modo tale da poter essere efficacemente supervisionati, durante il periodo in cui sono in uso, da persone fisiche, le quali – oltre ad essere in grado di comprendere appieno le capacità e i limiti del sistema – devono monitorarne il funzionamento (anche al fine di individuare e affrontare anomalie o disfunzioni) ed eventualmente intervenire sul sistema mede-

28 J. LAUX, *Institutionalised distrust and human oversight of artificial intelligence: towards a democratic design of AI governance under the European Union AI Act*, in *Law&Society*, 39, 2024.

29 *Linee Guida sul processo decisionale individuale automatizzato*, 31.

30 L'art. 13, par. 1, AI Act prevede che i sistemi ad alto rischio sia giustappunto progettati e sviluppati in modo tale da garantire che il loro funzionamento sia sufficientemente trasparente da consentire agli utenti di interpretare l'output del sistema e utilizzarlo adeguatamente.

simo, interpretarne correttamente gli *output* e, se necessario, decidere di discostarsi dalle indicazioni da esso fornite. Così, l'intervento umano richiesto dall'art. 22, par. 3, del GDPR non costituisce altro che la fase finale e attuativa della tutela più ampia e incisiva prevista dall'AI Act³¹.

Infine, occorre tener in considerazione il fatto per cui la diffusione di sistemi di IA e dell'attività di raccolta e di analisi dei *big data* hanno fatto sì che il processo decisionale possa risultare frammentato fra più soggetti della "catena del valore" collegata ai sistemi di IA. A tal riguardo, in una recente sentenza in materia di *credit scoring*³², la Corte di Giustizia UE ha adottato un approccio espansivo dell'ambito di applicazione dell'art. 22 GDPR. La Corte ha voluto neutralizzare i fenomeni di "segmentazione" della decisione, affermando che la parcellizzazione del trattamento automatizzato di dati personali tra più soggetti non esclude, di per sé, l'applicabilità dell'art. 22 GDPR, dovendosi piuttosto considerare il peso effettivamente attribuito alla predizione algoritmica nel processo decisionale. Le pronunce della Corte stanno progressivamente ampliando e meglio delineando l'ambito di applicazione dell'art. 22 GDPR, ben al di là della tradizionale nozione di "decisione", e indipendentemente dall'ambito in cui quel processo automatizzato viene utilizzato³³.

31 L. ENQVIS, *Human oversight' in the EU artificial intelligence act: what, when and by whom?*, in *Law, Innovation and Technology*, 15, 2, 2023, p. 512 secondo cui sebbene l'articolo 22 del GDPR preveda già un diritto limitato all'intervento umano nel processo decisionale esclusivamente automatizzato, i requisiti dell'Ala in materia di supervisione umana costituiranno una novità nel diritto europeo.

32 Corte giust. UE, sez. I, 7.12.2023, causa C-634/21, commentata da D. D'ORAZIO, *Il credit scoring e l'art. 22 del GDPR al vaglio della Corte di giustizia*, in *Nuova giur. civ. comm.*, 2, 2024; F. MATTASSOGLIO, *La Corte di giustizia europea, algoritmi e credit scoring. L'apertura del vaso di Pandora delle società che si "limitano" a elaborare gli scoring*, in *Dialoghi dir. econ.*, 1, 2025; in senso critico B. PAAL, *Article 22 GDPR: Credit Scoring Before the CJEU*, in *Global Privacy L. Rev.*, 3, 4, 2023. Al giudice europeo è stato chiesto se l'art. 22, par. 1, GPR debba essere interpretata nel senso che costituisce un «processo decisionale automatizzato relativo alle persone fisiche» ai sensi di tale disposizione, il calcolo automatizzato, da parte di una società che fornisce informazioni commerciali, di un tasso di probabilità basato su dati personali relativi a una persona e riguardante la capacità di quest'ultima di onorare impegni di pagamento in futuro.

33 D. D'ORAZIO, *Il credit scoring e l'art. 22 del GDPR al vaglio della Corte di giustizia*, cit., In particolare, la questione atteneva alla nozione di "decisione" di cui all'art. 22 GDPR. Secondo la Corte, la nozione di decisione ricomprende anche il mero «risultato del calcolo della solvibilità di una persona sotto forma di tasso di probabilità relativo alla capacità di tale persona di onorare impegni di pagamento in futuro». Di conseguenza, laddove il livello di rischio creditizio calcolato da una società terza fornitrice sia determinante ai fini della concessione di un credito da parte della banca, il calcolo svolto della società terza deve essere qualificato come decisione idonea a produrre nei confronti di un interessato «effetti giuridici che lo riguardano o che incide in modo analogo significativamente sulla sua persona», ai sensi dell'articolo 22 GDPR. Sicché l'interessato è legittimato ad esercitare il proprio diritto di accesso e di contestazione anche nei confronti di soggetti terzi comunque coinvolti nel processo decisionale.

1.3 Buone pratiche per l'utilizzo di sistemi automatizzati

Ferme le considerazioni che precedono sui presupposti di applicazione della disciplina in materia di trattamento automatizzato, le *Linee Guida* hanno fornito un elenco (non esaustivo) di buone prassi che il titolare del trattamento dovrebbe tenere in considerazione quando prende decisioni basate unicamente sul trattamento automatizzato, compresa la profilazione (definite all'art. 22, par. 1)³⁴. Tali *best practices* consistono in:

- a) controlli regolari di garanzia della qualità dei sistemi per assicurare che le persone siano trattate in maniera equa e non siano discriminate sulla base di categorie particolari di dati personali o in altro modo;
- b) verifica degli algoritmi sviluppati da sistemi di apprendimento automatico al fine di assicurarne il corretto funzionamento, la correttezza e la non discriminatorietà dei relativi risultati;
- c) fornitura all'ispettore di tutte le informazioni necessarie su come funziona l'algoritmo o il sistema di apprendimento automatico, in quanto essenziali per garantire l'audit indipendente di "terzi" (laddove il processo decisionale basato sulla profilazione abbia un impatto elevato sulle persone fisiche);
- d) ottenimento di garanzie contrattuali che sono stati effettuati *audit* e *test* e che l'algoritmo è conforme alle norme concordate, quanto meno per quanto riguarda gli algoritmi di terzi;
- e) predisposizione di misure specifiche per la minimizzazione dei dati al fine di prevedere periodi di conservazione chiari per i profili e per tutti i dati personali utilizzati durante la creazione o l'applicazione dei profili;
- f) utilizzo di tecniche di anonimizzazione o pseudonimizzazione nel contesto della profilazione;
- g) adozione di strumenti efficienti per consentire all'interessato di esprimere il proprio punto di vista e contestare la decisione;
- h) adozione di un meccanismo per l'intervento umano in determinati casi, ad esempio fornendo un collegamento a una procedura di ricorso nel momento in cui la decisione automatizzata viene trasmessa all'interessato, con termini concordati per il riesame e la designazione di un punto di contatto per qualsiasi domanda.

Infine, il titolare del trattamento può altresì valutare opzioni quali:

- a) l'applicazione di meccanismi di certificazione per i trattamenti;
- b) l'adozione di codici di condotta per la verifica dei processi che comportano appren-

³⁴ Gruppo di lavoro articolo 29, *Linee guida sul processo decisionale automatizzato relativo alle persone fisiche e sulla profilazione ai fini del regolamento 2016/679*, 3 ottobre 2017.

dimento automatico;

c) la costituzione di comitati di revisione etica per valutare i potenziali danni e benefici per la società di particolari applicazioni per la profilazione.

1.4 Liceità del trattamento automatizzato

A completamento della ricostruzione del quadro normativo rilevante ai fini qui in considerazione, vanno considerate le disposizioni speciali dedicate ai dati relativi alla salute quali *dati particolari*.

In merito, rileva la previsione dell'art. 9 del GDPR, che subordina il trattamento al requisito per cui l'interessato abbia prestato «il proprio consenso esplicito al trattamento di tali dati personali per una o più finalità specifiche».

Le condizioni del consenso ai fini della liceità del trattamento sono state chiarite dal Gruppo di Lavoro³⁵. Nell'ambito dei trattamenti automatizzati mediante IA, si pongono una serie di problematiche circa la possibilità di realizzare i requisiti di specificità, granularità e libertà del consenso³⁶. Il considerando 43 affronta la questione della granularità, stabilendo che «[s]i presume che il consenso non sia stato liberamente espresso se non è possibile esprimere un consenso separato a distinti trattamenti di dati personali, nonostante sia appropriato nel singolo caso, o se l'esecuzione di un contratto, compresa la prestazione di un servizio, è subordinata al consenso sebbene esso non sia necessario per tale esecuzione». Ciò presenta almeno due ordini di implicazioni per l'applicazione dell'IA. In primo luogo, all'interessato non deve essere richiesto un consenso *congiunto per tipi di trattamento* basati sull'IA essenzialmente *diversi*. In secondo luogo, il consenso alla profilazione deve essere prestato separatamente dall'accesso al servizio.

Ai fini della liceità del trattamento dei dati personali³⁷ è essenziale che la raccolta iniziale dei dati sia compatibile con finalità dichiarate dal titolare del trattamento, sulla base dei criteri previsti dall'art. 6.

L'art. 13, par. 1, lett. c), e l'art. 14, par. 1, lett. c), impongono ai responsabili del trattamento di fornire informazioni riguardanti «le finalità del trattamento cui sono destinati i dati personali e la base giuridica del trattamento». Il principio di *accountability* impone dunque al titolare del trattamento di fornire, per la validità del consenso, la prova che l'accesso e il trattamento contestati siano riconducibili alle finalità per le quali è stato validamente richiesto – e validamente ottenuto – un adeguato consenso, unitariamente

³⁵ Gruppo di lavoro articolo 29, *Linee guida sul consenso ai sensi del regolamento 2016/679*, WP 259, 28 novembre 2017.

³⁶ S. THOBANI, *La libertà del consenso al trattamento dei dati personali e lo sfruttamento economico dei diritti della personalità*, in *Eur. e dir. priv.*, 2, 2016, 532.

³⁷ In generale si veda D. POLETTI, *Le condizioni di liceità del trattamento dei dati personali*, in *Giur. it.*, 12, 2019.

alla prova circa la significatività delle informazioni sulla logica utilizzata nonché sull'importanza e le conseguenze previste di tale trattamento per l'interessato³⁸.

La limitazione delle finalità trova un problema con riguardo all'uso delle tecnologie di intelligenza artificiale e dei *big data*, in quanto consentono il riutilizzo dei dati personali per nuove finalità diverse da quelle per cui i dati sono stati originariamente raccolti. Di talché, per stabilire se il reimpiego dei dati sia legittimo, è necessario determinare se il nuovo scopo sia compatibile con quello per cui i dati sono stati originariamente raccolti. Secondo il Gruppo di Lavoro³⁹, i criteri rilevanti sono: (a) la distanza tra la nuova finalità e quella originaria; (b) l'allineamento della nuova finalità con le aspettative degli interessati, la natura dei dati e il loro impatto sugli interessi degli interessati; e (c) le misure di salvaguardia adottate dal responsabile del trattamento per garantire un trattamento equo e prevenire impatti indebiti.

2. Il rischio di discriminazione (e di conseguente esclusione sociale) derivante dall'utilizzo di sistemi di intelligenza artificiale nel diritto delle assicurazioni

2.1 Il principio di non discriminazione nei contratti assicurativi

Tracciato il contesto normativo di riferimento, l'attenzione deve ora essere concentrata sul rischio di discriminazione – e conseguente esclusione sociale – dei potenziali assicurandi. Questi ultimi, a fronte della sempre maggiore diffusione di modelli algoritmici complessi e funzionanti grazie a dinamiche operative difficilmente comprensibili dal fattore umano⁴⁰, potrebbero infatti non essere ricompresi all'interno dei *target markets* individuati *ex ante* dalle *policies* produttive e distributive dell'assicuratore.

L'utilizzo dell'IA nell'impresa assicurativa espone il mercato a potenziali conseguenze negative, tra cui il rischio di adottare delle forme di discriminazione che possono potenzialmente penalizzare le categorie più vulnerabili. Il tema può essere esaminato da varie angolazioni: da un lato, quello dell'utilizzo dell'IA in una prospettiva migliorativa del benessere sociale⁴¹ (principio, questo, che difficilmente potrebbe dirsi rispettato ove, per effetto dell'utilizzo dell'IA, venissero aggravate le situazioni di disagio esistenti); dall'altro, l'esigenza di garantire all'assicuratore la libertà di esercizio dell'impresa, in

³⁸ Il problema pratico è evidenziato da G. COMANDÉ, *Leggibilità algoritmica e consenso al trattamento dei dati personali*, cit.

³⁹ Gruppo di lavoro Articolo 29, *Parere 03/2013 sulla limitazione delle finalità* (WP 203).

⁴⁰ Ci si riferisce ai modelli di *machine based*, *machine learning* e *deep machine learning*, di cui, più approfonditamente, si rimanda *infra*, 31.

⁴¹ Proposta questa che ha trovato appoggio in varie sedi. In proposito si veda la Dichiarazione di Montreal sullo Sviluppo Responsabile dell'Intelligenza Artificiale, 2018, disponibile presso <https://montrealdeclaration-responsibleai.com/>, ma nello stesso senso si veda: L. FLORIDI, J. COWLS, *A unified framework of Five Principle for AI in society*, 2 luglio 2019, on *Harvard Data Science Review*, available at <https://hdsr.mitpress.mit.edu/>; si veda inoltre D. CAPONE, *La governance dell'Artificial Intelligence nel settore assicurativo tra principi etici, responsabilità del board e cultura aziendale*, in Quaderni Ivass, n. 16/2021.

particolare attraverso l'integrazione, nei propri processi interni, di strumenti tecnologici che consentano significativi risparmi in termini di costi operativi.

È necessario quindi capire quali siano i doveri di non discriminazione a cui è sottoposta l'impresa assicurativa, quali rischi ponga l'implementazione dell'IA, e infine, quali siano le *best practices* che è opportuno l'assicuratore adotti al fine di gestire al meglio tali rischi.

In questa prospettiva, è opportuno soffermarsi sulla portata del principio di non discriminazione all'interno dei contratti assicurativi. Il dibattito sulla possibilità di riconoscere al principio di non discriminazione efficacia diretta anche nei rapporti contrattuali tra privati – al di fuori, dunque, della sfera del diritto pubblico da cui origina – ha conosciuto rinnovata vitalità negli ultimi anni. L'espansione dell'ambito applicativo tipico del principio, operata del legislatore, limitatamente a specifici settori del diritto privato e ad alcune determinate forme di discriminazione, ha consentito un parziale superamento della tradizionale tesi secondo cui l'autonomia privata risulterebbe prevalente rispetto all'interesse della controparte a non essere esclusa dalla possibilità di soddisfare il proprio interesse mediante contratto – o, comunque, a non accedervi a condizioni deteriori – per ragioni discriminatorie, salvo negli specifici casi individuati dal legislatore⁴². Una parte della dottrina è infatti arrivata a sostenere che si debba attribuire al principio una portata generale: che cioè, anche al di fuori degli ambiti di applicazione tipica, questo possa avere forza cogente, tale da garantire al soggetto leso da comportamenti discriminatori il diritto a un rimedio effettivo (reale o risarcitorio)⁴³.

Alla luce delle diverse posizioni appena esposte, l'analisi non potrà che partire dall'individuazione delle norme di non discriminazione applicabili al settore assicurativo. In questo contesto, di particolare rilievo sono le disposizioni sul divieto di discriminazione per ragioni di etnia, razza o religione, per ragioni di genere o disabilità.

Il divieto di discriminazione per ragioni legate all'etnia e alla razza è desumibile dall'articolo 43 lett. b) del t.u. immigrazione (d. lgs. 286/1998) in cui rientra esplicitamente il

⁴² *Ex multis*, P. RESCIGNO, *Il principio di eguaglianza nel diritto privato: (a proposito di un libro tedesco)* in *Rivista trimestrale di diritto e procedura civile*, 1959, 1515 ss.

⁴³ In particolare, in senso di ampliamento rispetto alle ipotesi tipiche, si veda: E. NAVARRETTA, *Principio di uguaglianza, principio di non discriminazione e contratto* in *Rivista di diritto civile*, 2014, 548 ss; straordinariamente rilevante, per quanto qui di interesse, la posizione di R. SACCO, voce «Parità di trattamento», nel *Digesto IV, Disc. priv., sez. civ.*, 2012, su banche dati Wolters Kluwers, da leggersi qui congiuntamente a ID., voce «Giustizia contrattuale», nel *Digesto IV, Disc. priv., sez. civ.*, 2012, su banche dati Wolters Kluwers, in cui l'A. sottolinea che sebbene non sia riscontrabile nell'ordinamento una norma generale che garantisca la parità di trattamento, comunque "(...)verosimilmente opera nel nostro ordinamento un divieto di trattare il cliente in modo discriminatorio (...), e questo divieto è espresso se la lesione è operata nell'area del lavoro, o a danno di rifugiati e profughi, o è correlata a trattamenti deteriori praticati in ragione del sesso o della razza o della cultura del contraente (...)". L'A. procedimentalizza la regola nel corretto e libero accesso al mercato, consentendo dunque la difformità di trattamento ove questa trovi una giustificazione economica nei processi di impresa.

comportamento di chi pratici «condizioni più svantaggiose o si rifiuti di fornire beni o servizi offerti al pubblico ad uno straniero soltanto a causa della sua condizione di straniero o di appartenente ad una determinata razza, religione, etnia o nazionalità». Data tuttavia la rilevanza anche assistenziale del contratto *long-term care* qui considerato, potrebbero in tale sede rilevare anche le «condizioni più svantaggiose» o il rifiuto «di fornire l'accesso all'occupazione, all'alloggio, all'istruzione, alla formazione e ai servizi sociali e socio-assistenziali allo straniero regolarmente soggiornante in Italia soltanto in ragione della sua condizione di straniero o di appartenente ad una determinata razza, religione, etnia o nazionalità» di cui all'art 43 lett. c) del t.u. immigrazione (d. lgs. 286/1998). Il medesimo principio è stato poi statuito anche nella dir. CE/2000/43, che nel sancire il divieto di discriminazione anche indiretta in relazione alla razza o all'origine etnica, lo ritiene applicabile, ai sensi dell'art. 3 lett. h), «all'accesso a beni e servizi a disposizione del pubblico e alla loro fornitura», e, ai sensi della lett. f), «alle prestazioni sociali». Questa direttiva è stata recepita in Italia con d. lgs. 215/2003.

Il divieto di discriminazione di genere è stato introdotto con la dir. 2004/113, recepita in Italia con d. lgs. 196/2007, in cui si vieta ogni discriminazione, diretta o indiretta, fondata sul sesso. La direttiva aveva una specifica disposizione dedicata all'utilizzo del genere quale fattore di rischio nei servizi assicurativi e lasciava agli Stati Membri la possibilità di consentire differenze nei premi e nelle indennità ove il genere fosse un fattore determinante per il rischio. Sulla mancanza di una limitazione temporale a questa deroga è intervenuta nel 2011 la sentenza della CGUE *Test-Achats*⁴⁴, in cui la Corte, ritenendo che la mancanza di un limite temporale infici l'obiettivo di introdurre polizze assicurative *unisex*, ha deciso che la norma era invalida, seppure con effetto *ex nunc* a partire dal 21 dicembre 2012. A chiarimento e parziale contemperamento degli effetti della medesima sentenza è intervenuta la Commissione⁴⁵, la quale ha ritenuto che il genere potesse essere utilizzato: i) per il calcolo di rischio aggregato nel *pool* assicurativo; ii) per il calcolo del rischio nel contratto di riassicurazione; iii) nelle campagne pubblicitarie e di *marketing*; iv) per ragioni di *underwriting* delle polizze salute e vita, ove rilevino differenze fisiologiche tra uomo e donna. Infine, esclude la rilevanza delle discriminazioni indirette. Il principio ha trovato attuazione con Regolamento ISVAP n. 30/2009.

⁴⁴ C-236/09 – *Association Belge des Consommateurs Test-Achats and Others*.

⁴⁵ Commissione, 2012/C 11/01, *Guidelines on the application of Council Directive 2004/113/EC to insurance, in the light of the judgment of the Court of Justice of the European Union in Case C-236/09 (Test-Achats)*, 13 gennaio 2012. In realtà la CGUE in questa occasione ha dato un giudizio basato sulla coerenza interna della direttiva. In tale senso si veda: M. L. REGO, *Insurance Segmentation as Unfair Discrimination: What to Expect in the Wake of Test-Achats Proceedings of the 16th Annual Conference of the Insurance Law Association of Serbia*. Insurance law, governance and transparency: basics of the legal certainty, AIDA Serbia 2015, 382 ss available at SSRN: <https://ssrn.com/abstract=4325562> Sulla sentenza si veda anche P. MARANO, *Sex Discrimination in Private Insurance Contracts and the EU Law in Challenges in Harmonisation of the Serbian Insurance Law with the European (EU) Insurance Law*, 2012 Available at SSRN: <https://ssrn.com/abstract=2081983>.

Il divieto di discriminazione di persone affette da disabilità è sancito dalla l. n. 67/2006. Specificamente per quanto attiene al settore assicurativo, tuttavia, si deve ricordare che l'Italia è firmataria della Convenzione ONU Sui Diritti delle Persone con Disabilità, in cui all'art 25 lett. e) si vieta, nelle polizze assicurative del ramo vita, la discriminazione in ragione dell'invalidità, nell'evidente prospettiva di garantire l'accesso del disabile alla miglior tutela del proprio diritto alla salute⁴⁶; la Convenzione è stata ratificata con l. n. 18/2009. In data 23 luglio 2012, l'ISVAP ha emesso una lettera al mercato in materia di assicurazioni contro la malattia, ove l'Autorità Indipendente in adeguamento alla convenzione delle Nazioni Unite, e alla relativa disciplina di recepimento, vietava l'utilizzo di alcune clausole che restringevano l'accesso alle polizze vita alle persone disabili⁴⁷.

Da questa breve disamina si può notare come le disposizioni sul divieto di discriminazione di genere siano le uniche esplicite rispetto alla considerazione del fattore discriminatorio quale fattore di rischio rilevante ai fini della definizione di premi. Cionondimeno, benché il divieto non sia altrettanto esplicito in riferimento ai fattori di rischio concernenti etnia/razza/disabilità, si ritiene che anch'essi non possano essere utilizzati quali fattori di *pricing* giacché, altrimenti, il divieto di discriminazione nell'accesso ai servizi (anche) assicurativi sarebbe svuotato di contenuto. I comportamenti discriminatori hanno quindi ricadute di tipo legale (la violazione del principio di discriminazione) e determinano quindi un illecito presidiato dalla tutela reale o risarcitoria individuata dalle norme rilevanti, esponendo conseguentemente l'assicuratore a un importante *litigation risk*, nonché a rischi reputazionali e di condotta.

Quest'ultima considerazione sembra per contro valere anche in relazione ai comportamenti discriminatori che, pur non ricadendo in divieti tipici, siano comunque sanzionati dall'articolo 21 Carta Fondamentale dei diritti dell'Unione Europea. Nonostante dalla citata norma non possa sicuramente trarsi un divieto di utilizzo del fattore attuariale rilevante, è ipotizzabile un onere di motivazione rafforzato atto a giustificare l'utilizzo in termini di causalità e corretta identificazione del *pool* di rischio⁴⁸.

⁴⁶ In proposito: D. SANTOVITO, C. BOSCO, A. POLOTTI DI ZUMAGLIA, *Assicurazioni malattia e vita, la nuova frontiera posta dall'art. 25 della Convenzione delle Nazioni Unite sui diritti delle persone con disabilità*, in *Assicurazioni*, 2021, 193 ss.

⁴⁷ ISVAP, Lettera al mercato 23 luglio 2012, disponibile su https://www.ivass.it/normativa/nazionale/secondaria-ivass/lettere/2012/Im-07-23/LETTERA_AL_MERCATO.pdf

⁴⁸ M. L. REGO, *Statistics as a basis for discrimination in the insurance business in Law, Probability and Risk* 2015, 119 ss; anche in M. L. REGO, *Insurance Segmentation as Unfair Discrimination: What to Expect in the Wake of Test-Achats*, cit, 385 ss. Pur non in materia di diritto assicurativo, in favore di un generale onere di giustificazione verso comportamenti discriminatori sembra esprimersi anche D. MAFFEIS, *Libertà contrattuale e divieto di discriminazione*, in *Rivista trimestrale di diritto e procedura civile*, 2008.

2.2 L'inquadramento del principio di non discriminazione nella correttezza dei rapporti contrattuali

Il principio di non discriminazione rileva nei rapporti contrattuali ai fini della trasparenza e della correttezza dell'intermediario verso l'utente dei prodotti assicurativi⁴⁹. Esso infatti conforma la correttezza dell'agire imprenditoriale dell'assicuratore da un lato, nell'individuazione dei fattori di rischio rilevanti per l'esercizio dell'impresa e nella conseguente predisposizione di *pool* adeguati, e dall'altro, nel calcolo dei premi assicurativi individuali. In tale senso si colloca nell'organizzazione e governance dell'impresa, e contestualmente incide sui profili di tutela del cliente.

Pertanto, oltre ai principi generali posti dal codice civile (artt. 1175, 1337, e 1375 c.c.), nei contratti assicurativi, in ragione delle specificità del relativo mercato e dei prodotti offerti, gli assicuratori sono sottoposti a doveri di correttezza e buona fede che conformano le fasi di pubblicità, *design*, stipula ed esecuzione del contratto dettate dalle normative specialistiche sul tema. Sotto questo profilo, particolare attenzione meritano gli artt. 17, 20, e 25 IDD, così come trasposti nell'ordinamento nazionale.

L'art. 17 IDD, trasposto nel nostro ordinamento all'art 183 del codice assicurazioni private (d'ora in avanti c.a.p.), impone all'intermediario assicurativo di agire «*honestly, fairly, and professionally in accordance with the best interest of their clients*» e, in relazione alla fase pubblicitaria, richiede una comunicazione «*clear and not misleading*». In quest'ambito, la rilevanza dei principi discriminatori può derivare in particolare dallo sfruttamento di pubblicità targettizzate di prodotti assicurativi. Si noti che tali pratiche, allo stato, non sono sanzionate dal legislatore, e che, al contrario, le Linee Guida della Commissione in materia di discriminazione di genere le consentono esplicitamente. Tuttavia, è limitato l'utilizzo di pubblicità che possano sfruttare vulnerabilità (derivanti, per esempio, da una marginalizzazione già esistente) per la collocazione di prodotti assicurativi a premi maggiori⁵⁰.

L'art 20 IDD attiene specificamente all'informazione precontrattuale: il criterio a cui questa si conforma, al netto degli obblighi specifici in materia di documentazione che deve essere fornita al cliente, è quello di fornire un'indicazione comprensibile e oggettiva, che consenta al cliente di assumere una decisione informata.

Infine, l'art. 25 IDD regola i processi di *product oversight and governance*, che individuano anche le demografiche di rischio, e sono fondamentali ai fini del *pricing* assicurativo.

In tale prospettiva, integrando una violazione delle sopra citate norme, i comportamenti discriminatori che pregiudichino i clienti dell'assicuratore in base a caratteristiche personali integrano un rischio di condotta.

⁴⁹ Si parla in questo senso di «*actuarial fairness*»: E. W. (JED) FREES, F. HUANG *The Discriminating (Pricing) Actuary*, in *North American Actuarial Journal*, 2023, 2 ss.

⁵⁰ In materia di assicurazioni non vita, una preoccupazione di questo tipo è espressa in EIOPA, *Supervisory statement on differential pricing practices in non-life insurance lines of business*, cit.

Fatto salvo quanto detto in materia di non discriminazione e in alcuni interventi regolatori dell'Autorità di Vigilanza (riferiti a specifici profili di rischio⁵¹), non ci sono *de iure condito* ulteriori disposizioni specifiche che impongano all'assicuratore doveri di inclusione finanziaria, in ragione della necessità di rendere accessibile il prodotto. Tuttavia, non si può non rilevare come anche nell'impresa assicurativa i processi normativi in atto abbiano portato all'integrazione dei fattori ESG nei rischi rilevanti per la sana e prudente gestione, che rendono ipotizzabile la rilevanza prudenziale di spinte volte a promuovere l'inclusione finanziaria.

2.3 L'IA e i rischi connessi alla discriminazione nel settore assicurativo

È un dato acquisito che l'integrazione dell'IA nei processi operativi possa incrementare il rischio di incorrere in discriminazioni non consentite. Questo dipende direttamente dalle modalità di funzionamento delle IA⁵² basate sull'individuazione di correlazioni statistiche tra i dati forniti come *input*: a essere fattore di discriminazione potrebbe essere – in astratto – direttamente il design dell'algoritmo, ma, più frequentemente, si tratta di *dataset* utilizzati nella fase di *training* non adeguatamente «puliti», e che riflettono *bias* e pregiudizi, in quanto le variabili ritenute determinanti ricorrono maggiormente in alcuni gruppi marginalizzati, oppure perché altri gruppi risultano sottorappresentati nei *dataset*⁵³.

A seconda di come l'IA venga implementata nel processo produttivo dell'assicuratore, ne possono derivare comportamenti discriminatori in molteplici fasi.

a) Al momento del *pricing* del contratto assicurativo, è possibile che l'IA discrimini direttamente (fondando la propria decisione su fattori di rischio vietati in quanto discriminatori) o indirettamente (traendo correlazioni da variabili apparentemente neutre, ma che restituiscono il potere predittivo di una variabile legalmente vietata in quanto discriminatoria: c.d. *proxy discrimination*)⁵⁴. Il fenomeno della discriminazione indiretta ha finora trovato scarsa attenzione da parte dell'ordinamento, in virtù della sua più limitata

51 In particolare, in materia di rischi sistemici EIOPA, *Supervisory statement on exclusions in insurance products related to risks arising from systemic events*, disponibile al link [Supervisory statement on exclusions in insurance products related to risks arising from systemic events](https://www.eiopa.europa.eu/sites/default/files/2021-03/Supervisory_statement_on_exclusions_in_insurance_products_related_to_risks_arising_from_systemic_events.pdf).

52 EIOPA, *Report on Artificial intelligence governance principles: towards ethical and trustworthy artificial intelligence in the European insurance sector*, 2021.

53 J. SMITH, *AI Bias in Life & Annuities insurance*, 5 marzo 2024, in <https://www.insurancethoughtleadership.com/>; più in generale, si veda G. CARAPEZZA FIGLIA, *Decisioni algoritmiche tra diritto alla spiegazione e divieto di discriminare*, in *Persona e mercato*, 2023, 640.

54 Sui rischi di discriminazione indiretta si vedano: E. W. (JED) FREES, F. HUANG, op. cit., 15 ss, A: PRINCE, D. SCHWARCZ, *Proxy Discrimination in the Age of Artificial Intelligence and Big Data* in *Iowa Law Review*, 2020, 1257, Available at SSRN: <https://ssrn.com/abstract=3347959>. Quest'ultimo in particolare sottolinea che l'utilizzo dell'IA potrebbe incrementare esponenzialmente le discriminazioni indirette giacché il funzionamento di tali tecnologie si basa sulla rilevazione di correlazione tra variabili fra loro non necessariamente causalmente legate. Specificamente sull'uso di *credit scoring* assicurativo, si veda P. MANES, *Credit scoring assicurativo, machine learning e profilo di rischio: nuove prospettive*, in *Contratto e impresa*, 2021, 469 ss.

incidenza nei processi decisionali non integralmente automatizzati. Tuttavia, il tema è destinato ad acquisire sempre maggiore importanza con la progressiva implementazione dell'IA con riferimento alla politica di *underwriting* dell'impresa assicurativa. Si pensi al funzionamento della stessa per correlazioni statistiche, di talché potrà rintracciare la presenza di variabili dipendenti dall'appartenenza a gruppi protetti e trarne conseguenze sulla probabilità di concretizzazione del rischio. Ancora, è possibile che vengano individuate correlazioni tra caratteristiche personali (il cui uso attuariale non è a prescindere vietato, ma necessita di giustificazione) e rischi, così che il diverso esito del processo di *pricing* non rifletta un reale rapporto di causalità.

b) In fase pubblicitaria, l'IA potrebbe individuare *pattern* di consumo non correlati al rischio, che consentono all'assicuratore di sfruttare una maggiore elasticità di prezzo (c.d. *differential pricing*), con l'effetto di aggravare potenziali situazioni di vulnerabilità⁵⁵.

c) Nella fase di *design* del contratto, potrebbe determinare un *tailoring* eccessivo del profilo di rischio, per la grande disponibilità di dati legati all'individuo: in tale situazione il soggetto con profilo di rischio troppo alto può essere escluso dal contratto assicurativo. Tale aspetto risulta problematico specialmente alla luce del valore sociale riconosciuto alle polizze *long-term care*, nella prospettiva dell'isolamento e dell'aggravamento della marginalizzazione di categorie vulnerabili.

Allo stesso modo, l'uso di algoritmi potrebbe determinare l'esclusione o l'aumento sproporzionato dei costi per soggetti che non forniscono dati sufficienti quantitativamente e qualitativamente, con conseguente aggravamento di situazioni di vulnerabilità sociale.

d) In fase di rinnovo si potrebbe verificare un aumento dei costi dovuto alla maggiore elasticità di prezzo⁵⁶; in alcuni casi questa pratica potrebbe essere sanzionata quale pratica commerciale scorretta⁵⁷.

e) A seguito del concretizzarsi del rischio assicurato, nella fase di *claims management*, l'assicuratore potrebbe formulare offerte deteriori in relazione all'appartenenza a categorie protette, sulla base di ipotesi formulate dall'IA sulla maggiore o minore probabilità che il cliente agisca in giudizio. Nella stessa fase, l'implementazione dell'IA quale strumento di *fraud detection* potrebbe rivelarsi discriminatorio ove portasse a rilevare una maggiore frequenza di frodi in relazione a caratteristiche personali protette.

⁵⁵ Si veda il già citato in EIOPA, *Supervisory statement on differential pricing practices in non-life insurance lines of business*; evidenzia il problema anche EIOPA, *Report on Artificial intelligence governance principles: towards ethical and trustworthy artificial intelligence in the European insurance sector*, 32 ss.

⁵⁶ V. anche International Association of Insurance Supervisor, *Draft application paper on the supervision of Artificial Intelligence*, November 2024, 26.

⁵⁷ EIOPA, *Supervisory Statement on differential pricing practices in non-life insurance lines of business*, EIOPA-BoS-23/076 22 February 2023.

I comportamenti discriminatori sopra individuati⁵⁸ pongono non solo il problema di garantire il rispetto del divieto di discriminazione posto dall'ordinamento, ma, in un senso più ampio, anche di come implementare nuovi strumenti tecnologici all'interno dell'attività di impresa in modo etico e orientato al benessere della comunità nel suo complesso. Il tema risulta particolarmente problematico in relazione a polizze, quali quelle di *long-term care*, che hanno una causa assistenziale per i soggetti assicurati. In una prospettiva etica e di *Corporate Social Responsibility*, la necessità di garantire un'inclusione finanziaria realmente accessibile richiede che l'implementazione dell'intelligenza artificiale avvenga con peculiari cautele.

Si tenga conto che il tema rileva anche per l'integrazione dei principi ESG nell'impresa assicurativa. L'adozione di strategie di impresa che determinano un deterioramento del fattore S può rilevare per la sana e prudente gestione nella misura in cui questo determini dei rischi finanziari per l'assicuratore, per esempio sotto forma di *reputation* o *litigation risk*. Il tema va oltre l'identificazione del confine tra pratiche ammesse e pratiche illecite, e incide sulle modalità attraverso cui deve essere pensata l'integrazione tecnologica nell'industria.

Tanto viene ribadito esplicitamente all'interno dei considerando dell'*AI Act*⁵⁹, laddove si statuisce che le imprese sono invitate ad adottare le buone pratiche e i codici di condotta coerenti con un'IA etica e affidabile, secondo le linee guida individuati dall'High Level Expert Group nel 2019⁶⁰, tra cui spicca anche il principio di diversità, non discriminazione ed equità (considerando 27).

⁵⁸ Va ricordato che l'utilizzo dell'IA o, più precisamente, il fenomeno della c.d. *Big Data Analysis* (BDA) – cioè di quell'insieme di tecniche, che possono avvalersi anche dell'utilizzo dell'intelligenza artificiale, per estrarre informazioni, pattern e tendenze significative da enormi quantità di dati – è stato da tempo oggetto di una certa attenzione da parte di EIOPA e, soprattutto, con riferimento ad alcuni rami, tra tutti, il più significativo, si è rivelato quello corrispondente al ramo 10 danni. In questo senso, dunque, appare quanto più interessante porre mente al documento di EIOPA, *Big data analytics in motor and health insurance: a thematic review*, 2019. Con questo atto di *soft law* il regolatore europeo aveva, invero, anticipato alcuni degli sviluppi dell'IA nel settore assicurativo in seguito considerati sempre da altre fonti EIOPA e, in generale, dagli operatori del settore. Le principali aree di utilizzo della BDA individuate in questo atto da EIOPA erano infatti: *pricing* e *underwriting*; vendita e distribuzione; servizi successivi alla vendita e, in questo senso, va detto che i BDA forniscono una serie di vantaggi anche nella gestione del rapporto contrattuale, come lo sviluppo di comunicazioni semplici e rapide con i clienti, la funzionalizzazione dei reclami; *rilevamento delle frodi*, che può essere implementato anche grazie all'utilizzo di tecniche *ad hoc* come il *claims scoring*.

⁵⁹ Si noti fin da subito che una parte della dottrina ha evidenziato, tuttavia, come l'*AI Act* abbia dato un apporto modesto in tema di non discriminazione. In tale senso: R. DE CARIA, *L'AI Act e il divieto di discriminazioni*, in *Media Laws*, 30 marzo 2022, available at <https://www.medialaws.eu/lai-act-e-il-divieto-di-discriminazioni/>.

⁶⁰ Artificial Intelligence High Level Expert Group, *Orientamenti etici per un'IA affidabile*, 8 aprile 2019.

Si noti in questa prospettiva che l'AI Act ha proibito gli usi dell'IA considerati più pericolosi in quanto idonei a sfruttare situazioni di vulnerabilità o a perpetrare direttamente discriminazioni⁶¹. In particolare, per quanto qui di interesse, l'art 5 AI Act proibisce l'uso di un sistema di AI che sfrutta le vulnerabilità di una persona, dovute all'età, alla disabilità, o a una specifica situazione sociale ed economica, con l'obiettivo o l'effetto di distorcerne il comportamento e con la ragionevole probabilità di creare un danno significativo (art 5, par. 1, lett. b)⁶².

Inoltre, si proibisce l'uso di IA che pratici *social scoring*, ossia, dia valutazioni delle persone basate su caratteristiche personali o di personalità o comportamenti sociali, ove questo determini un effetto pregiudizievole in contesti *non collegati a quello di origine del comportamento, o comunque un effetto pregiudizievole ingiustificato o sproporzionato* (art. 5, par. 1, lett. c).

Infine, il divieto colpisce anche l'uso di IA di categorizzazione biometrica per trarre deduzioni o inferenze in merito a razza, opinioni politiche, appartenenza sindacale, convinzioni religiose o filosofiche, vita sessuale o orientamento sessuale (salvo per scopi di etichettatura o filtraggio di dati acquisiti legalmente) (art 5, par. 1, lett. g).

Il *profiling* assicurativo condotto con l'IA rientra nella definizione di *social scoring*⁶³. Infatti, in un caso⁶⁴ in cui un ente creditizio si era servito di un algoritmo per la valutazione del merito creditizio dei propri clienti, la CGUE ha ritenuto di trovarsi di fronte a una forma di profilazione, in quanto trattavasi di previsione sulla capacità futura di adempiere a un'obbligazione alla luce dell'appartenenza a una determinata classe di persone con caratteristiche omogenee. In particolare, è considerato *social scoring* coperto dal divieto di cui all'art. 5, par. 1, lett. c) dell'AI Act l'utilizzo di IA per la definizione del premio o la decisione di stipulare un'assicurazione sulla vita, ove i dati raccolti siano di spesa e finanziari, non correlati al rischio sulla vita⁶⁵. Come anticipato, il divieto di cui all'art 5 par. 1 lett. c) sarà tuttavia applicabile solo nel caso in cui tutte le condizioni richieste dalla norma siano soddisfatte, e, significativamente, ove derivino dalla valutazione effetti pregiudizievoli in contesto diverso da quello in cui è stata posta in essere l'azione valutata, o comunque si tratti di un effetto pregiudizievole sproporzionato o ingiustificato. Se ne trae, da un lato, che non sono considerate forme di *social scoring* vietate le ipotesi di utilizzo dell'IA da parte dell'assicuratore per la valutazione del rischio a cui

⁶¹ Si vedano sul punto i consideranda 28 e ss. dell'AI Act

⁶² La rilevanza della norma in rapporto al principio di non discriminazione è esplicitamente riconosciuta nelle, European Commission, *Guidelines on prohibited artificial intelligence practices established by Regulation (EU) 2024/1689 (AI Act)* 52, C(2025) 884, 4 febbraio 2025, 52.

⁶³ European Commission, *Guidelines on prohibited artificial intelligence practices established by Regulation*, cit., 53.

⁶⁴ CGUE 7 December 2023, SCHUFA Holding (Scoring), C-634/21,

⁶⁵ European Commission, *Guidelines on prohibited artificial intelligence practices established by Regulation*, cit. 60.

è esposto l'assicurato a scopo di *underwriting*, a condizione che l'algoritmo consideri dati rilevanti per il rischio assicurato e da esso derivi un pregiudizio proporzionato al comportamento considerato⁶⁶. Il divieto di *social scoring* ex art. 5 par. 1 lett. c) ha infine l'effetto di imporre all'assicuratore un rafforzamento del rapporto causale tra i fattori di rischio considerati e l'esito della decisione algoritmica, richiedendo sia un collegamento di contestualità tra il comportamento o la caratteristica personale e il rischio assicurato, sia la proporzionalità tra la gravità del comportamento o della caratteristica e la valutazione algoritmica. Lo scopo è anche quello di minimizzare i potenziali effetti discriminatori del *social scoring*, specialmente ove le caratteristiche personali siano afferenti all'appartenenza a gruppi marginalizzati, e pertanto siano potenzialmente discriminatorie, in via diretta o indiretta⁶⁷. Il divieto di cui all'art. 5 lett. c) può coprire anche le ipotesi di *unfair differential pricing*, condotto attraverso condotte miranti a sfruttare caratteristiche personali diverse dalle mere preferenze di consumo⁶⁸.

L'AI Act vieta esplicitamente quegli usi dell'IA correlati a un elevato rischio di discriminazione. Fermi restando tali divieti, lo sviluppo di un'IA affidabile ed etica pone il problema di elaborare ulteriori buone pratiche per garantire il rispetto della diversità e l'uguaglianza nell'esercizio dell'attività assicurativa: tanto dovrà avvenire sulla base di un principio di proporzionalità che tenga conto, attraverso le sopra citate valutazioni di impatto: i) da un lato, da come l'IA venga integrata nella *value chain* assicurativa; ii) dall'altro, del tipo di IA effettivamente utilizzata.

2.4 L'adozione di buone pratiche da parte dell'assicuratore

2.4.1 Obblighi di trasparenza e spiegabilità

Nel contesto in discorso, lo High Level Expert Group ed EIOPA hanno definito una serie di buone pratiche rilevanti nell'ambito assicurativo. Tali buone pratiche costituiscono requisiti a sé stanti per l'implementazione dell'IA nei processi produttivi, assumendo rilevanza quali condotte strumentali al contrasto dei fenomeni discriminatori sopra descritti⁶⁹.

⁶⁶ Esplicitamente, European Commission, *Guidelines on prohibited artificial intelligence practices established by Regulation*, cit., 61.

⁶⁷ Tuttavia, è stato sottolineato che il legislatore Europeo ha predisposto, così facendo, una norma non coerente con le attuali prassi del settore assicurativo, in cui si raccolgono enormi quantità di dati, sostanzialmente decontestualizzati, e poi, a seguito di una valutazione di correlazione statistica *machine-based*, si seleziona i fattori di rischio ritenuti rilevanti: D. MINTY, *The social scoring ban – confusion for insurers*, in <https://www.ethicsandinsurance.info>, 19 febbraio 2025.

⁶⁸ European Commission, *Guidelines on prohibited artificial intelligence practices established by Regulation*, cit., 63.

⁶⁹ Riflessioni non dissimili d'altronde si trovano anche nel *Guidance Resource: Artificial intelligence and discrimination in insurance pricing and underwriting*, predisposta dalla Australia Human Rights Commission e dal Actuaries Institute nel 2022, e disponibile presso <https://actuaries.asn.au/Library/Opinion/DataScienceAI/2022/2022GuidanceResourceAI.pdf>

Un primo gruppo di buone pratiche si riconduce agli obblighi di informazione e trasparenza, la cui correlazione con il rischio di discriminazione è di immediata evidenza: in effetti, il principio di non discriminazione può essere letto come una specificazione degli obblighi di correttezza e buona fede. La *disclosure*, quale strumento primario di attuazione di questi obblighi, diventa quindi essenziale per un uso legittimo dell'IA. D'altronde, gli artt. 13 e 14 del GDPR già richiedono il rilascio una specifica informativa per l'utilizzo dei dati degli utenti nei processi decisionali automatizzati, che ha contenuto rafforzato in caso di decisione totalmente automatizzata. Quest'ultimo requisito è stato letto da una parte della dottrina come un vero e proprio diritto alla spiegazione della logica sottesa all'algoritmo, che ha finalità di tutela contro l'utilizzo dei dati per scopi discriminatori⁷⁰. La scelta sembra infine confermata anche dallo stesso *AI Act*, che nell'allegato IV, a precisazione dei contenuti dell'informativa ex art. 11⁷¹, richiede che vengano date informazioni anche sui possibili impatti discriminatori derivanti dall'integrazione di strumenti algoritmici ad alto rischio nei processi produttivi.

Si deve tenere conto del fatto che il contenuto dell'informativa prospettata da EIOPA⁷² è direttamente rapportato agli scopi di garantire un'IA etica, di tal che tale informativa deve essere tenuta ben distinta dagli obblighi informativi generali appena rappresentati. Inoltre, solo la prima si rivolge anche alle funzioni interne di controllo, nonché agli organi di controllo e supervisione e alle autorità di Vigilanza. Dunque, EIOPA individua una serie di informazioni che devono essere fornite all'utente del servizio assicurativo, al consiglio di amministrazione dell'impresa assicurativa, nonché alla funzione di *audit* della stessa e all'autorità di supervisione. Alcune di queste informazioni sembrano in effetti finalizzate a consentire l'esercizio di un controllo contro possibili fenomeni discriminatori e comprendono: (i) quali dati sono immessi nell'algoritmo; (ii) come è integrata l'IA nel processo produttivo; (iii) quali dati sono utilizzati per il *pricing*; e (iv) l'individuazione dei fattori più influenti attraverso ragionamenti controfattuali⁷³. Tuttavia, per evitare fenomeni di discriminazione indiretta, questi strumenti potrebbero non essere sufficienti, data la difficoltà di identificare la correlazione tra una determinata variabile e l'appartenenza a una categoria protetta. In dottrina si è pertanto suggerito di integrare gli obblighi di trasparenza con la *disclosure* degli effetti che l'utilizzo di sistemi di IA può produrre sulle categorie più vulnerabili e maggiormente esposte a rischio di discriminazione⁷⁴.

⁷⁰ G. CARAPEZZA FIGLIA, *op.cit.*, 643 ss.

⁷¹ Sui sistemi ad alto rischio nel Regolamento, si veda *supra*, pp. 2-4.

⁷² EIOPA; Report on Artificial intelligence governance principles: towards ethical and trustworthy artificial intelligence in the European insurance sector, cit., 43.

⁷³ EIOPA; Report on Artificial intelligence governance principles: towards ethical and trustworthy artificial intelligence in the European insurance sector, cit., 43.

⁷⁴ A: PRINCE, D. SCHWARCZ, *op. cit.*, 1313.

Con particolare riferimento ad alcuni sistemi di IA, quali i modelli di *deep learning* – e in particolare architetture di reti neurali (c.d. *neural network*) – che precludono un’effettiva comprensione del funzionamento dell’algoritmo e della logica sottostante a un determinato *output*, EIOPA prescrive le seguenti misure che l’impresa assicuratrice può adottare per garantire la trasparenza nei confronti degli interessati⁷⁵:

a) L’adozione di spiegazioni adattate alle specificità del caso di utilizzo dell’IA e agli *stakeholders* destinatari.⁷⁶ A parità di condizioni, maggiore è l’impatto di un caso di utilizzo dell’IA, maggiori sono le misure di trasparenza e spiegabilità che l’impresa di assicurazione dovrebbe adottare. Ad esempio, nella determinazione dei prezzi e nella sottoscrizione, dovrebbe essere chiaramente spiegato ai consumatori perché un determinato rischio non è ritenuto accettabile o quali sono i principali fattori di *rating* che influenzano il premio⁷⁷.

b) L’utilizzo di modelli di IA spiegabili, che si rende particolarmente necessario nei casi di utilizzo dell’IA ad alto impatto⁷⁸. In questo senso, EIOPA individua due principali opzioni in materia di *explainable AI*. La prima consiste nell’utilizzo di modelli intrinsecamente interpretabili, quali gli alberi decisionali o la regressione lineare, che l’Autorità raccomanda di privilegiare quando l’esigenza di trasparenza e comprensibilità dell’algoritmo prevale su quella di massimizzare l’accuratezza. La seconda opzione, invece, prevede l’impiego di sistemi *black-box* supportati da strumenti di spiegabilità supplementari – quali LIME (*Local Interpretable Model-agnostic Explanations*) e SHAP (*SHapley*

⁷⁵ L’obbligo di trasparenza è previsto da una serie di disposizioni della normativa europea. Ai sensi dell’articolo 20, par. 1, IDD, i distributori assicurativi devono fornire al cliente informazioni obiettive sul prodotto assicurativo in forma comprensibile per consentirgli di prendere una decisione informata. Come è stato visto, il principio di trasparenza è sancito anche dall’art. 5 GDPR, ulteriormente sviluppato dagli articoli 13 e 14 GDPR, che impongono alle imprese di informare i consumatori in modo tempestivo, adeguato e trasparente sulle modalità di trattamento dei loro dati personali. Il GDPR afferma esplicitamente che i consumatori devono essere informati anche dell’esistenza di processi decisionali automatizzati (spesso sostenuti dalla tecnologia AI) e di fornire loro informazioni “significative” sulla logica coinvolta, nonché sul significato e sulle conseguenze previste di tale trattamento.

⁷⁶ EIOPA; Report on Artificial intelligence governance principles: towards ethical and trustworthy artificial intelligence in the European insurance sector, cit., 41-42.

⁷⁷ EIOPA; Report on Artificial intelligence governance principles: towards ethical and trustworthy artificial intelligence in the European insurance sector, cit., 43-44.

⁷⁸ Nonostante infatti EIOPA suggerisca di superare le criticità connesse alla *black box* mediante programmi di *explainable AI* (che cercano di replicare il funzionamento dell’*AI Machine Learning*), attualmente non può ritenersi raggiunto un grado di certezza significativo tra le spiegazioni fornite e il software utilizzato. Sulla funzionalità e l’efficienza di strumenti di *explainable AI*. Si vedano sul punto J. PEARL, *The limitations of opaque learning machines*, in *Possible minds*, 25 ways of looking at AI, J. Brockman (a cura di), New York, 2019; in Z. BEDNARZ, K. MANWARING, *Keeping the (good) faith: implications of emerging technologies for consumer insurance contracts*, in *Sidney Law Review*, 2021, 455 ss.

Additive exPlanations)⁷⁹ – per garantire una comprensione adeguata del funzionamento del modello e del relativo *output*.

c) Spiegazioni significative e di facile comprensione per aiutare gli *stakeholder* a prendere decisioni informate.

A fronte di un obbligo di trasparenza e spiegabilità che, in aderenza con quanto espresso da EIOPA, si rivela quanto meno corposo, appare necessario comprendere quali informazioni debbano essere fornite all'utente assicurativo medio. Questo è vero specialmente in relazione alle polizze *long-term care* qui considerate, dal momento che gli specifici rischi assicurati porteranno l'assicuratore a rapportarsi con utenti persone fisiche e consumatori, con un grado di conoscenza del funzionamento dell'IA ridotta.

2.4.2 Il controllo sull'algoritmo

Un altro strumento utile al contrasto ai fenomeni discriminatori è l'intervento diretto sull'algoritmo. In questo senso, le regole di non discriminazione vengono direttamente incorporate nel funzionamento dell'IA⁸⁰. Questo risultato può essere ottenuto in diversi modi, e l'assicuratore dovrà scegliere quale soluzione adottare secondo un principio di proporzionalità rispetto al rischio posto dall'utilizzo concreto dell'IA.

In particolare, uno strumento efficace consiste nell'intervenire direttamente sui *dataset* usati per l'addestramento, convalida e prova del modello. Sul punto, l'art 10 dell'AI Act stabilisce esplicitamente che nei sistemi IA ad alto rischio tali *dataset* devono essere sottoposti a pratiche adeguate di *governance* e gestione dei dati, e in particolare devono essere sottoposti a un esame atto ad evitare che l'algoritmo abbia effetti discriminatori. Più in generale, un uso dell'AI etico richiede di garantire che i dati siano accurati, completi e rappresentativi per il fine a cui il sistema è preposto⁸¹; si devono rimuovere le caratteristiche personali vietate in quanto discriminatorie e le *proxy* evidenti⁸². Inoltre, i dati utilizzati per il *training* devono essere raccolti nel rispetto delle norme applicabili⁸³.

In una prospettiva più generale, è opportuno verificare le deduzioni algoritmiche, al fine di garantire che le differenze nelle condizioni di contratto dipendano da fattori discretivi sul piano causa – effetto. Resta possibile un utilizzo di fattori causali definiti *ex ante*

79 Per una spiegazione dei vari sistemi si veda R. GUIDOTTI ET AL., *A Survey of Methods for Explaining Black Box Models*, ACM Computing Surveys, 5, 51, 2018, disponibile su <https://doi.org/10.1145/3236009>.

80 IAIS, *Draft application paper on the supervision of Artificial Intelligence*, cit., 27.

81 IAIS, *Draft application paper on the supervision of Artificial Intelligence*, cit., 28; EIOPA, *Report on Artificial intelligence governance principles: towards ethical and trustworthy artificial intelligence in the European insurance sector*, cit., 57.

82 EIOPA, *Report on Artificial intelligence governance principles: towards ethical and trustworthy artificial intelligence in the European insurance sector*, cit., 40.

83 IAIS, *Draft application paper on the supervision of Artificial Intelligence*, cit., 30.

quali variabili di controllo sugli *outcomes* algoritmici⁸⁴.

Infine, a fronte del rilievo di eventuali pregiudizi insiti nell'algoritmi è raccomandabile intervenire direttamente sull'algoritmo⁸⁵. Specificamente in materia di discriminazione indiretta si può ipotizzare un'integrazione di variabili protette all'interno dell'algoritmo, in modo da fare sì che le altre variabili vengano considerate dall'IA per il loro valore predittivo diretto e non in quanto *proxy* delle prime, per poi rimuovere gli effetti dell'integrazione delle variabili protette⁸⁶.

2.4.3 Supervisione umana (*human oversight*)

Lo *human oversight* si afferma in principio come un generale obbligo di controllo sul buon funzionamento del modello di IA, affinché siano ridotti al minimo i rischi per i diritti fondamentali⁸⁷. Tale principio rappresenta un presidio di tutela, sia con riferimento al plesso normativo del GDPR che dell'AI Act, benché con sfumature e significati differenti.

Se nel GDPR esso ha infatti una funzione di tutela nella sola ipotesi del processo automatizzato in un'ottica di tutela del consumatore (v. *supra* par. 1.2.4.)⁸⁸, nel governo dell'IA esso assume una portata più ampia, in quanto comporta un vero e proprio sistema di controllo che, nella catena produttiva assicurativa, investe l'intera struttura dell'impresa assicuratrice dalla fase del disegno del prodotto a quelle successive alla distribuzione.

Con riferimento alla gestione del rischio di discriminazione derivante dall'utilizzo dell'IA, il principio di *human oversight* si concretizza nella predisposizione di un organigramma idoneo al controllo dell'algoritmo. In considerazione della rilevanza prudenziale del tema, la presenza di soggetti competenti nella corretta gestione dei modelli di intelli-

⁸⁴ IAIS, *Draft application paper on the supervision of Artificial Intelligence*, cit., 29.

⁸⁵ Per esempio, sul tema si vedano gli studi di M. LINDHOLM, R. RICHMAN, A. TSANAKAS AND M.V. WÜTHRICH, *Discrimination free insurance pricing*, in *Austin Bulletin*, 2022, 55 ss.

⁸⁶ A: PRINCE, D. SCHWARCZ, *op. cit.*, 1313 ss.

⁸⁷ Sul punto, tra gli altri, R. KOULU, *Proceduralizing Control and Discretion: Human Oversight in Artificial Intelligence Policy*, in *Maastricht Journal of European and Comparative Law*, 27, 2020; *Id.*, *Human Control over Automation: EU Policy and AI Ethics*, in *European Journal of Legal Studies*, 1, 9, 41, 2020; G. LAZCOZ, P. DE HERT, *Humans in the GDPR and AIA Governance of Automated and Algorithmic Systems. Essential Pre-Requisites against Abdicating Responsibilities*, Brussels Privacy Hub Working Paper, 8, 2022; J. LAUX, *Institutionalised distrust and human oversight of artificial intelligence: towards a democratic design of AI governance under the European Union AI Act*, in *AI & Society*, 39, 2024.

⁸⁸ Come è stato esposto nel paragrafo 1.2.4., quando le decisioni dell'assicurazione sono assunte mediante sistemi completamente automatizzati, sono individuabili due momenti in cui la supervisione umana può essere eseguita: 1) prima che il sistema di IA sia implementato, in conformità alle disposizioni dell'AI Act e del Report; 2) dopo che il titolare del trattamento ha assunto la decisione, in attuazione della misura di salvaguardia prevista dall'art. 22, par. 3, GDPR.

genza artificiale può intendersi quale presidio obbligatorio di *governance*⁸⁹.

In questa prospettiva e coerentemente con un principio di proporzionalità della misura adottata, rispetto alla valutazione di impatto degli strumenti di AI, il controllo del personale incaricato a che non vi siano comportamenti discriminatori potrebbe comprendere⁹⁰, laddove rilevanti, l'enucleazione di figure formalmente incaricate di condurre in via sistematica una review dei *dataset* in modo da rimuovere i *biases* che dovessero essere rilevati. Ancora, può essere opportuno in talune circostanze limitare l'autonomia operativa dell'algoritmo per favorire un controllo sugli esiti, e nella fase di operatività dell'algoritmo prevedere figure dedicate allo scrutinio dei processi, alla gestione degli incidenti, nonché alla revisione degli *output*⁹¹. In questo scenario, importanza centrale spetta altresì al tema di AI literacy, di cui all'art. 4 dell'AI Act.

Le funzioni poc'anzi descritte possono essere affidate a uffici interni previsti dalla normativa assicurativa o, diversamente, a figure professionali precedentemente non contemplate nella struttura organizzativa d'impresa.

Nel primo caso, una particolare attenzione al rischio di discriminazione dovrebbe essere prestata dalla funzione di *compliance* (o *legal department*), da quella di *risk management* e dalla funzione attuariale; in ottica di supervisione, un ruolo preminente potrebbe essere assunto dalla funzione di *audit*. Nel secondo, invece, potrebbero essere create figure professionali *ad hoc*, con un *expertise* mirato alla comprensione del funzionamento dell'IA e delle possibili criticità che solleva nella prospettiva dell'etica e della correttezza nei rapporti con la clientela⁹².

In termini più generali, l'assicuratore dovrebbe adottare strumenti soluzione delle controversie con il cliente, in cui l'utente che ritenga di essere stato illegittimamente discriminato a causa dell'utilizzo di IA da parte dell'assicuratore, possa far valere le proprie ragioni. Questo diritto è connesso all'effettiva comprensione da parte dell'utente del funzionamento dell'IA e della sua implementazione nei processi produttivi, renden-

89 EIOPA, *Consultation paper on opinion on artificial intelligence governance and risk management*, EIOPA-BoS-25-007, 10 febbraio 2025, laddove l'inserimento di professionalità altamente specializzate nella gestione dei modelli di IA potrebbe perfino essere integrata nell'ufficio di *risk management*, integrando anche il rischio di non discriminazione tra quelli potenzialmente realizzabili durante l'esercizio dell'attività assicurativa.

90 EIOPA; *Report on Artificial intelligence governance principles: towards ethical and trustworthy artificial intelligence in the European insurance sector*, cit., 49 ss.

91 Specialmente in relazione all'*machine learning* un controllo sugli esiti può essere fondamentale per controllare che non vi siano eventuali discriminazioni indirette: v. anche R. J. LEHMAN, *Why Big Data will force insurance companies to think hard about race*, in *Insurance Journal*, 27 marzo 2018, <https://www.insurancejournal.com/blogs/right-street/2018/03/27/484530.htm>

92 La rilevanza del cd. "*human in the loop*" è sottolineata anche nel disegno di legge modello predisposto dalla NAIC: si veda NAIC Model Bulletin, *Use of Artificial Intelligence system by insurers*, 12 febbraio 2023; in questo documento è significativa la distinzione tra la gestione dell'IA (con relativa ripartizione dei compiti e delle responsabilità nell'organigramma) e il momento di controllo, da affidarsi alla funzione di *compliance* e controllo del rischio.

do particolarmente importante la trasparenza verso l'utente⁹³.

2.4.4 Governo dei dati e tenuta dei registri (*data governance and record keeping*)

Ai fini dell'implementazione di un'IA etica, EIOPA sottolinea l'importanza dei processi che garantiscono un'elevata qualità dei dati utilizzati⁹⁴.

A questi fini rileva sia il momento *ex ante*, che il momento *ex post*.

Quanto al primo, esso investe la fase di raccolta dei dati che – come si è detto in precedenza – prevede di identificare e selezionare dati conformi alle specifiche finalità di utilizzo dell'AI e, in fase di preparazione, di procedere a un primo studio che permetta anche una valutazione qualitativa dei dati. In seguito, si procederà alla gestione delle problematiche qualitative e, in particolare, di eventuali *bias*, non solo in riferimento ad attributi protetti ma anche ad eventuali *proxies*.

Quanto al secondo, appare necessaria una valutazione *ex post* dell'*output* dell'IA per evitare esiti discriminatori. Queste regole sono da applicarsi, secondo EIOPA, tanto ai dati forniti da parte del cliente, quanto a quelli di fonte esterna, che poi saranno sempre più frequenti con la diffusione di *IoT*. Si aggiunge che però in quest'ultimo caso potrebbero trovare applicazione le norme del *Data Governance Act* oppure del *Data Act*.

Tra le buone pratiche suggerite per il miglioramento qualitativo dei *dataset*, si è proposto anche un aumento quantitativo delle informazioni: questo può funzionare soprattutto quando i dati attualmente a disposizione dell'assicuratore siano insufficienti, incompleti, o comunque viziati in ragione di discriminazioni che la comunità marginalizzata abbia vissuto⁹⁵.

Un'ulteriore tecnica a disposizione dell'assicuratore consiste nel *record keeping*, con cui l'impresa tiene traccia delle informazioni concernenti l'implementazione dell'IA (le finalità strategiche sottese, i processi di gestione dei dati, i codici, la performance e sicurezza dell'algoritmo, e la valutazione di impatto etico) consentendone una maggiore controllabilità, così da parte degli uffici interni come pure da parte dell'Autorità di Vigilanza.

La tenuta dei registri costituisce un obbligo organizzativo d'impresa sancito tanto dal

⁹³ IAIS, *Draft application paper on the supervision of Artificial Intelligence*, cit., 30.

⁹⁴ EIOPA; *Report on Artificial intelligence governance principles: towards ethical and trustworthy artificial intelligence in the European insurance sector*, cit., 56 ss. In particolare EIOPA sottolinea che la *governance* dei dati debba essere solida (*sound*) e prudente (*prudent*). Sul punto si veda anche D. ACHILLE, *I «Registri delle attività di trattamento» tra responsabilizzazione e presupposto di liceità del trattamento dei dati personali*, in *European Journal of Privacy and Technologies*, 2, 2024. Sulla *governance* dei dati si era già ragionato in anticipo, v. *supra*, par. 1.2.

⁹⁵ In tale senso si veda anche Australia Human Rights Commission, Actuaries Institute, *Guidance Resource: Artificial intelligence and discrimination in insurance pricing and underwriting*, cit.

GDPR, all'art. 30⁹⁶ quanto dagli artt. 12 e 26 dell'AI Act. Sulla scorta dell'articolo 12, il legislatore europeo stabilisce che sistemi ad alto rischio devono tenere automaticamente traccia del periodo di utilizzo, dei dataset utilizzati, dei dati di *input* per cui l'IA ha riscontrato corrispondenza.

Mentre, in base all'articolo 16, si prevede che i *deployer* conservino i *log* – cioè i registri di attività – generati automaticamente dal sistema di IA ad alto rischio sotto il loro controllo per un periodo adeguato alle finalità previste.

Appare evidente come la finalità sottesa a questi articoli non risulti soltanto nell'individuazione di situazioni di rischio per i diritti fondamentali. Semmai, e soprattutto, nel monitoraggio e nella tracciabilità dei processi di gestione dei dati e delle metodologie di modellazione.

2.4.5 Robustezza e prestazione (*Robustness and performance*)

Il rafforzamento della componente tecnica e algoritmica è compresa tra le *best practices* indicate da EIOPA⁹⁷ e in questa visuale devono essere considerate, dunque, la *performance* e la *robustness* di un sistema IA. Ciò significa che robustezza e prestazione devono essere continuamente monitorate, adattate e convalidate, tenendo conto delle limitazioni del modello e delle sue applicazioni.

Le imprese assicurative devono fare un uso accurato dei dati, gestire correttamente le infrastrutture *IT* e avere piani di emergenza pronti in caso di fallimento dei sistemi IA.

In questa prospettiva, il *Report* EIOPA sui principi di *governance* dell'intelligenza artificiale traccia alcuni requisiti a cui l'impresa deve conformarsi e che appaiono centrali per un uso efficiente e non discriminatorio dell'IA, ossia:

a) *Impatto dell'uso dell'IA sulla robustezza e performance*: la robustezza e le misure di *performance* richieste per un sistema di intelligenza artificiale dipendono dall'uso specifico cui è destinato e dal danno potenziale che potrebbe derivare. Maggiore è il rischio di danno, più rigorosi devono essere i requisiti di *performance* e monitoraggio. Al fine di determinare questo standard sono essenziali le politiche di *governance* adottate, come la supervisione umana.

b) *Definizione dell'obiettivo dell'IA*: le imprese assicurative devono chiarire l'obiettivo che intendono raggiungere con il sistema IA in modo da stabilire metriche di *performance* appropriate. Per applicazioni ad alto impatto, è fondamentale che le compagnie possano spiegare il processo decisionale del modello, anche se ciò potrebbe ridurre le prestazioni.

⁹⁶ Tale disposizione prevede la tenuta di un registro in cui devono essere conservate le informazioni relative al titolare del trattamento, le finalità, il tipo di dati analizzati e i soggetti interessati.

⁹⁷ EIOPA, *Report on Artificial intelligence governance principles: towards ethical and trustworthy artificial intelligence in the European insurance sector*, cit., 63 ss.

c) *Monitoraggio e valutazione continuativa delle performance dell'IA, per verificarne la precisione e per comprendere la natura degli errori*: le compagnie assicurative devono essere consapevoli delle limitazioni dei modelli IA e implementare restrizioni d'uso per evitare che vengano impiegati in circostanze che possano compromettere la *performance*. Quando vi sono incertezze significative sul rischio è accettabile adottare modelli meno performanti; tuttavia, le imprese devono fornire giustificazioni sul perché un determinato livello di *performance* sia accettabile, soprattutto in presenza di incertezze o mancanza di dati. La mancanza di dati rilevanti può, infatti, compromettere la *performance* del modello e in questi casi le compagnie devono fornire spiegazioni chiare su come i dati insufficienti influenzino le previsioni. Per evitare conseguenze negative, è importante che la giustificazione sia documentata per ogni variazione dei prezzi e che sia possibile spiegare ogni differenza sostanziale tra periodi di tempo. Se i dati cambiano in modo significativo, l'impatto di questi cambiamenti deve essere quantificato e divulgato. Le compagnie devono inoltre pianificare delle opzioni di emergenza nel caso in cui un modello di *pricing* non funzioni correttamente, per esempio se il modello risultasse discriminatorio o non performante.

d) *Adattamento delle metriche di performance*: le metriche di *performance* devono essere personalizzate in base all'uso specifico dell'IA. Ad esempio, in un sistema di rilevamento delle frodi, l'obiettivo potrebbe essere massimizzare la precisione. È necessario monitorare anche l'impatto della scelta della metrica sui consumatori vulnerabili

e) *Gestione dei dati per assicurare la performance dell'IA*: affinché sia predisposta una gestione solida dei dati, essi devono essere accurati, completi, imparziali e rappresentativi per il compito dell'IA. Ogni modello IA deve essere sviluppato con una valutazione della qualità dei dati e con un aggiornamento delle procedure di gestione.

f) *Calibrazione e validazione del modello*: i sistemi IA dovrebbero essere rivalutati, ricalibrati e validati periodicamente, soprattutto in caso di cambiamenti significativi nei dati di *input* o nell'ambiente legale ed economico. La validazione deve essere documentata e, per i sistemi IA ad alto impatto, deve includere analisi di scenario e *stress test*.

g) *Sistemi IT resilienti per supportare l'AI e sistemi di recupero*⁹⁸: le compagnie assicurative devono aggiornare le loro infrastrutture IT per far fronte alle esigenze computazio-

⁹⁸ A tal proposito rileva anche il Regolamento (UE) 2022/2554 (DORA), seppur non faccia espresso riferimento a sistemi di IA o a particolari sistemi rafforzati per la protezione di dati personali sensibili. Il DORA sancisce (i) il principio di proporzionalità, per cui le entità finanziarie attuano le norme tenendo conto delle loro dimensioni e del loro profilo di rischio complessivo, nonché della natura, della portata e della complessità dei loro servizi, delle loro attività e della loro operatività (art. 4); (ii) un quadro per la gestione dei rischi informatici solido, esaustivo e adeguatamente documentato, che consenta alle entità finanziarie di affrontare i rischi informatici in maniera rapida, efficiente ed esaustiva, assicurando un elevato livello di resilienza operativa digitale (art. 6); (iii) le disposizioni su diritti e obblighi dell'entità finanziaria e del fornitore terzo di servizi TIC, incluse le disposizioni in materia di disponibilità, autenticità, integrità e riservatezza in relazione alla protezione dei dati, compresi i dati personali (art. 30).

nali necessarie per implementare soluzioni IA. Ciò include la protezione contro attacchi informatici che potrebbero compromettere il modello di IA utilizzato. Per le applicazioni ad alto impatto, le compagnie devono sviluppare piani di recupero nel caso in cui l'IA non funzioni come previsto o subisca un attacco informatico.

h) *Sistemi IA esternalizzati*: quando le compagnie assicurative esternalizzano soluzioni IA da fornitori terzi, questi devono rispettare gli stessi requisiti di robustezza e *performance*, assicurandosi che i fornitori offrano AI di alta qualità e comprendano le limitazioni dei modelli IA. Inoltre le imprese devono essere preparate ai rischi di concentrazione, derivanti dalla dipendenza di un singolo fornitore e sviluppare le relative strategie di *exit*.

2.4.6 Valutazioni di impatto

A conclusione dell'analisi delle misure di contrasto alla discriminazione e all'esclusione finanziaria di natura assicurativa deve essere, infine, ricordato un ulteriore strumento che viene individuato dall'AI Act, in parte anche dal GDPR e, infine, è supportato anche da EIOPA. In verità, non fa specificamente parte di quelle *best practices* sostenute direttamente dal regolatore europeo, ma costituisce un supporto specificamente stabilito dalla normativa summenzionata e che potrebbe dunque trovare spazio tra gli strumenti di contrasto alle discriminazioni e di sostegno a un uso *fair* dell'IA.

Si tratta della valutazione di impatto nel caso di utilizzo dell'IA che, ben prima dell'entrata in vigore del Regolamento, era stata assunta dall'Autorità europea come mezzo che consentisse alla *governance* d'impresa di affrontare in termini proporzionati i rischi a cui sia l'assicurato che l'impresa sono esposti a causa dell'utilizzo dell'IA nella catena di valore assicurativa⁹⁹.

Ai fini della valutazione d'impatto, l'impresa di assicurazione dovrebbe pertanto procedere a classificare la probabilità (*likelihood*) e la gravità (*severity*) del danno derivante da uno specifico utilizzo dell'IA, definendo, ad esempio, tre livelli di impatto potenziale dell'AI (alto/medio/basso)¹⁰⁰.

a) Per quanto concerne il livello di gravità del rischio cui è esposto il cliente, il *Report EIOPA on Artificial Intelligence* del 2021, fornisce alcuni possibili parametri utili, quali: (i) il numero di consumatori coinvolti; (ii) la tipologia di consumatori, considerando maggiormente quelli vulnerabili; (iii) il grado di autonomia umana; (iv) il rischio di discriminazione. Con riguardo, invece, all'impatto sull'impresa assicuratrice, devono essere valutati, ad esempio: (i) la continuità aziendale; (ii) l'impatto finanziario; (iii) i rischi legali,

99 L'atto con cui EIOPA ha per la prima volta introdotto la valutazione d'impatto è il *Report on Artificial Intelligence Governance Principles: Towards Ethical And Trustworthy Artificial Intelligence In The European Insurance Sector. A report from EIOPA's Consultative Expert Group on Digital Ethics in insurance*.

100 Le imprese di assicurazione possono definire un numero maggiore di livelli se lo ritengono opportuno.

specie nell'ipotesi in cui l'utilizzo possa risultare in una illegittima discriminazione fra i consumatori; (iv) l'impatto reputazionale¹⁰¹.

b) Sono invece considerati fattori di probabilità di concretizzazione del danno: (i) il ricorso a valutazioni o attribuzioni di punteggi, compresa la profilazione e la previsione; (ii) il ricorso a processi decisionali automatizzati con effetti legali o materiali significativi; (iii) l'adozione di sistemi di monitoraggio sistematico; (iv) il grado di complessità del modello/combinazione di insiemi di dati; (v) l'uso innovativo o l'applicazione di una soluzione tecnologica o organizzativa nuova per l'impresa; (vi) il tipo e la quantità di dati utilizzati; (vii) l'esistenza di meccanismi di esternalizzazione dei dati e del trattamento tramite IA.

Per ragioni di completezza, va ricordato che la valutazione d'impatto promossa da EIO-PA era già stata predisposta attraverso le linee guida del Gruppo di Lavoro in materia di valutazione di impatto sulla protezione dei dati (DPIA)¹⁰², nella misura in cui anche il DPIA mirava ad affrontare rischi simili a quelli derivanti dalle applicazioni di IA.

In seguito, la normativa primaria dell'Unione ha seguito un approccio analogo, come confermato dall'articolo 27 dell'AI Act che richiede al *deployer* di procedere alla valutazione d'impatto sui diritti fondamentali (FRIA) per i sistemi di IA ad alto rischio. Il par. 4 dello stesso articolo chiarisce che, qualora tale obbligo sia stato già rispettato mediante la valutazione d'impatto sulla protezione dei dati effettuata mediante il DPIA a norma dell'articolo 35 del GDPR¹⁰³, non è necessario procedere ad ulteriori valutazioni. Le misure *governance* dei dati previste dal GDPR hanno dunque diretta rilevanza anche per la *governance* dell'AI, che ne costituisce un'integrazione.

¹⁰¹ L'impatto sulle imprese di assicurazione si basa fundamentalmente sui rischi che le imprese di assicurazione valutano nell'ambito del *Own Risk and Solvency Assessment* (ORSA) ai sensi degli artt. 44 e 45 della direttiva 2009/138/CE (Solvency II).

¹⁰² Gruppo di lavoro Articolo 29, *Linee guida in materia di valutazione d'impatto sulla protezione dei dati e determinazione della possibilità che il trattamento "possa presentare un rischio elevato" ai fini del regolamento (UE) 2016/679*, 4 aprile 2017.

¹⁰³ L'articolo 35 GDPR impone che «quando un tipo di trattamento, allorché prevede in particolare l'uso di nuove tecnologie, considerati la natura, l'oggetto, il contesto e le finalità del trattamento, può presentare un rischio elevato per i diritti e le libertà delle persone fisiche, il titolare del trattamento effettua, prima di procedere al trattamento, una valutazione dell'impatto dei trattamenti previsti sulla protezione dei dati personali». Al paragrafo successivo vengono individuati una serie di casi che il legislatore presume «a rischio elevato per i diritti e le libertà» e per i quali è necessaria la valutazione, quali: «a) una valutazione sistematica e globale di aspetti personali relativi a persone fisiche, basata su un trattamento automatizzato, compresa la profilazione, e sulla quale si fondano decisioni che hanno effetti giuridici o incidono in modo analogo significativamente su dette persone fisiche»; e «b) il trattamento, su larga scala, di categorie particolari di dati personali di cui all'articolo 9, paragrafo 1».

Le linee guida sul DPIA hanno specificato quali dovrebbero essere considerate le ipotesi di trattamento di dati ad alto rischio, includendovi:

- a) la valutazione o assegnazione di un punteggio – inclusiva di profilazione e previsione – in particolare in considerazione di aspetti riguardanti il rendimento professionale, la situazione economica, la salute, le preferenze o gli interessi personali, l'affidabilità o il comportamento, l'ubicazione o gli spostamenti dell'interessato (ad esempio, un ente finanziario che esamini i suoi clienti rispetto a una banca dati di riferimento in materia di crediti);
- b) ogni processo decisionale automatizzato che produca effetti giuridici o incida significativamente in modo analogo;
- c) il monitoraggio sistematico di comportamenti degli interessati;
- d) il trattamento di dati sensibili o dati aventi carattere altamente personale;
- e) il trattamento di dati su larga scala, da individuare tenendo conto, in particolare, del numero di soggetti interessati, del volume dei dati e/o le diverse tipologie di dati oggetto di trattamento, della durata o persistenza dell'attività di trattamento; della portata geografica dell'attività di trattamento;
- f) la creazione di corrispondenze o combinazione di insiemi di dati;
- g) il trattamento di dati relativi a soggetti interessati vulnerabili, in quanto tale trattamento costituisce motivo dell'aumento dello squilibrio di potere tra gli interessati e il titolare del trattamento;
- h) l'uso innovativo o applicazione di nuove soluzioni tecnologiche od organizzative;
- i) il trattamento che impedisce agli interessati di esercitare un diritto o di avvalersi di un servizio o di un contratto, come nel caso dei trattamenti che mirano a consentire, modificare o rifiutare l'accesso degli interessati a un servizio oppure la stipula di un contratto.

2.5 IA e processi di *product oversight and governance*

Le buone pratiche in materia di implementazione dell'IA coinvolgono anche i processi di *product oversight and governance*. Specificamente in quest'ultima prospettiva rileva il *differential pricing*, con cui l'assicuratore differenzia il premio del soggetto assicurato non perché questi rientra in una diversa classe di rischio, ma in ragione di caratteristiche personali che, determinando il comportamento di consumo dell'assicurato, gli consentono di sfruttare a proprio favore una maggiore elasticità di prezzo; la maggiore incidenza di queste condotte si riscontra in fase di rinnovamento della polizza.

Va rimarcata la rilevanza di queste condotte non solo ai fini della correttezza dei rapporti tra assicuratore e assicurato ma anche ai fini della disciplina in materia di pratiche commerciali scorrette (Dir. 29/2005/CE), laddove vi sia concretamente il rischio che

possano sviare il comportamento del consumatore.

In materia di *differential pricing*, EIOPA ha emanato un *supervisory statement*¹⁰⁴ finalizzato a fornire indicazioni alle autorità indipendenti di settore e agli operatori su come verrà esercitata la vigilanza. L'ambito applicativo è circoscritto ai prodotti assicurativi «non-vita», rispetto a cui il fenomeno ha più ovvia incidenza in quanto essi richiedono un periodico rinnovamento. Conseguentemente, tale aspetto interessa il prodotto assicurativo *Long Term Care* solo nella misura in cui esso assuma la forma di polizza di danno. Il *differential pricing* può essere particolarmente dannoso, in quanto idoneo ad aggravare situazioni di vulnerabilità ed esclusione finanziaria e sociale. Pertanto, nell'ottica di limitarne gli effetti pregiudizievoli, EIOPA, nella propria funzione di promotore della convergenza tra le Autorità di supervisione, propone un rafforzamento dei processi di *product oversight and governance*.

Nel proprio intervento regolatorio, l'Autorità di Vigilanza individua pratiche non coerenti con nessun *target market*, quali ad esempio pratiche di «*price walking*» che determinano un aumento costante del premio al momento del rinnovo, non legato all'aumento del rischio.

Rispetto ad altre pratiche, invece, si richiede l'adozione di accorgimenti operativi. Specialmente in riferimento alla fase di approvazione del prodotto, è richiesta una chiara distribuzione delle responsabilità che preveda il coinvolgimento di soggetti in posizioni anche apicali. In senso più sostanziale, inoltre, EIOPA richiede la predisposizione di soglie per le differenze di prezzo tra soggetti aventi un profilo di rischio simile; infine, viene prospettato un rafforzamento degli obblighi di trasparenza e informazione. In caso di utilizzo di *software*, si fissa l'obbligo di illustrare come questa impatti sul prodotto offerto all'utente. Inoltre, per il particolare caso dell'uso di strumenti algoritmici, il *supervisory statement* richiede di fornire al cliente spiegazioni adeguate dell'impatto di questi algoritmi sul *pricing* e l'implementazione dei dovuti processi di *governance* (in questo caso, *human in the loop*).

Inoltre, EIOPA richiede che nel praticare sconti ad alcuni utenti l'assicuratore tenga in considerazione il rischio di prodotto, la propria strategia e modello di *business* nonché il *target market*, rispetto a cui si tratta di garantire la coerenza del prodotto con i relativi interessi e caratteristiche.

Nella fase di *product testing* il prodotto dovrebbe inoltre essere testato contro impatti avversi verso il *target market*, conseguentemente evitando di immetterlo nel mercato nel caso in cui emergano possibili effetti dannosi. Successivamente, l'assicuratore è onerato dal monitoraggio a che questi prodotti non abbiano un impatto negativo, compito che potrà assolvere attraverso l'uso di metriche e predisponendo adeguati rimedi per i casi già verificatisi. I processi così descritti devono essere adeguatamente docu-

104 EIOPA, *Supervisory Statement on differential pricing practices in non-life insurance lines of business*, cit.

mentati. La documentazione relativa al prodotto dovrà essere fornita al distributore, per consentirgli di agire nel migliore interesse dell'utente.

3. Long-term care policies, intelligenza artificiale e wearable devices. Quali riscontri dal mercato?

3.1 Utilizzo e potenziali esternalità positive dell'IA e dei wearable data nelle polizze long-term care

A conclusione della presente analisi, occorre opportuno svolgere una panoramica dello stato attuale del mercato, europeo ed extraeuropeo, in relazione alla presenza di polizze contro il rischio di non autosufficienza funzionanti attraverso modelli di intelligenza artificiale.

In questa prospettiva, appare opportuna una notazione preliminare volta a evidenziare come l'impiego dell'intelligenza artificiale si collochi in naturale continuità con la funzione primaria dell'attività assicurativa, tradizionalmente imperniata sulla gestione del rischio. In particolare, la fase di underwriting presuppone che l'assicuratore raccolga e analizzi il maggior numero possibile di informazioni relative al potenziale assicurato, al fine di ricondurlo a uno specifico mercato di riferimento. Più precisamente, tale attività è funzionale all'inserimento del soggetto all'interno di un determinato pool di rischi omogenei, composto da clienti che, condividendo caratteristiche analoghe, risultano riconducibili alla medesima classe di rischio.

In questo contesto, i sistemi di intelligenza artificiale consentono all'assicuratore di misurare e valutare il rischio con maggiore efficacia, permettendo una più puntuale associazione tra il profilo dell'assicurato e il rischio assunto. Non sorprende, dunque, che negli ultimi anni le imprese assicurative si siano progressivamente dotate di modelli di calcolo computazionale fondati su intelligenze artificiali avanzate, quali i sistemi machine based, di machine learning e di deep machine learning. Tali strumenti tecnologici permettono infatti di elaborare, in tempi significativamente ridotti e con elevati livelli di accuratezza, volumi di dati sempre più ampi e complessi, incidendo profondamente sulle modalità di valutazione e gestione del rischio assicurativo¹⁰⁵.

Al netto delle problematiche che l'utilizzo delle nuove tecnologie porta con sé – relative al governo e corretto utilizzo dell'IA, nonché al trattamento di dati sensibili e alla concreta possibilità che alcuni soggetti vulnerabili rimangano esclusi dal servizio as-

¹⁰⁵ La maggiore accuratezza nei processi di calcolo da parte delle intelligenze artificiali potenziate, è possibile proprio grazie ad una processazione di una grande quantità di dati e, con riferimento ai sistemi che si basano sui cc.dd. alternative data, dall'utilizzo da parte dei sistemi AI di metodologie di tipo induttivo. Questi dati non riguardano esclusivamente aspetti strettamente legati alle caratteristiche personali del plausibile assicurato – ad esempio, l'età, la condizione fisica, la provenienza territoriale, le abitudini di vita –, ma possono essere desunti da altre fonti, come le ricerche effettuate *on-line* da parte del cliente, le abitudini di spese desunte dallo studio dei movimenti bancari, l'attività sui social media.

sicurativo¹⁰⁶ – non se ne possono negare le evidenti ricadute positive. Esse consistono, in particolare, in un *risk assesment* puntualmente definito e ciò, di conseguenza, si concretizza in una politica di *pricing* coerente con il profilo dell'assicurato e in linea con l'operatività dell'impresa assicurativa.

In tale contesto, l'impiego dell'intelligenza artificiale a fini di innovazione e di miglioramento della qualità dei prodotti di *long-term care* assume particolare rilievo, soprattutto con riferimento a quei prodotti che operano mediante sistemi di IA capaci di elaborare dati raccolti attraverso l'utilizzo di *wearable data*. Si tratta, in particolare, di applicazioni integrate in *smartwatches* o *smartphones*, nonché di dispositivi dotati di tecnologia GPS, che vengono messi a disposizione dell'assicurato direttamente dall'impresa assicurativa¹⁰⁷.

Dal punto di vista qualificatorio, va peraltro precisato che i «dati indossabili» non rientrano all'interno della definizione di «sistema di IA» fornita dall'art. 3 co. 1 n. 1) del Regolamento¹⁰⁸. Piuttosto, essi sono *fonti di dati*, per lo più biometrici o ambientali, che possono alimentare dei sistemi di IA che analizzano comportamenti, realizzano determinati ragionamenti – come predire malattie, anticipare certe abitudini, delineare con certezza il rischio di avveramento di un sinistro – e prendono, infine, decisioni in maniera automatica.

Nella filiera produttiva assicurativa, l'utilizzo dei *wearable data* si colloca, non solo, ma principalmente, in una fase del processo produttivo successiva alla valutazione del rischio e al conseguente *pricing*. Una volta concluso il contratto essi consentono di mo-

¹⁰⁶ Quanto al rischio di discriminazione e di esclusione sociale si rimanda *supra* al par. 2 del report.

¹⁰⁷ Una prima definizione, più che di *wearable data*, di *wearable technology* era stata fornita dieci anni fa dal *Business Innovation Observatory* della Commissione Europea, in cui si specificava che «Wearable technology refers to clothing and personal accessories that incorporate advanced computer and electronic technologies. The products that contain wearable technology are called wearables. Applications of wearable technologies include wearable cameras, smart clothing, wearable apps platforms, smart glasses, activity trackers, smart watches, as well as health and happiness wearables». Esistono per vero anche versioni più «invasive» di *wearable data*, come *microchips* sottocutanei, addirittura dei veri e propri *smart tatoos*. Per una panoramica completa di queste nuove frontiere dei dispositivi indossabili si leggano, *ex multis*, A. OMETOV, *A Survey on Wearable Technology: History, State of Art and Current Challenges*, *Computer Networks* 193 (2021) 108074; A. PIRRERA, D. GIANANTI, *Smart Tatoos Sensors 2.0: A Ten Years Progress Report through a Narrative Review*, *Bioengineering* 11(4) 376.

¹⁰⁸ La disposizione citata considera un sistema di IA come: «un sistema automatizzato progettato per funzionare con livelli di autonomia variabili e che può presentare adattabilità dopo la diffusione e che, per obiettivi espliciti o impliciti, deduce dall'*input* che riceve come generare *output* quali previsioni, contenuti, raccomandazioni o decisioni che possono influenzare ambienti fisici o virtuali».

nitorare le abitudini di vita, i comportamenti¹⁰⁹, nonché i parametri biologici dell'assicurato. In seguito, una volta che l'evento integrante il rischio assicurato, tempestivamente rilevato dai dispositivi, dovesse realizzarsi, la compagnia esegue la propria controprestazione¹¹⁰.

Le esternalità positive legate all'assunzione dei *wearable data* sono di non poco momento¹¹¹.

La prima è legata strettamente all'antecedente politica di *pricing* della polizza, giacché tali presidi tecnologici intercettano i livelli vitali a cui il rischio è collegato e, laddove dovessero mutare, l'assicurazione potrebbe modulare diversamente il prezzo contrattuale, pur sempre in una prospettiva *client oriented*, ispirata a un alto grado di personalizzazione del prodotto.

Secondariamente, le modalità di funzionamento della polizza così strutturata facilitano il flusso comunicativo cliente-impresa. In concreto, l'obbligo di avviso previsto ai sensi dell'art. 1913 c.c. si realizza istantaneamente, poiché il *trigger event* viene comunicato in tempo reale all'assicuratore tramite l'applicazione connessa al supporto tecnologico in dotazione al cliente. Non è dunque necessario che l'assicurato denunci la propria non autosufficienza per mezzo di raccomandata a.r. o PEC, come si può, invece, attualmente desumere dallo studio dei prodotti ramo danni e vita – comprensivi, questi ultimi, anche delle *LTC policies* – attualmente offerti sul mercato.

109 Non a caso gli autori, come pure gli operatori del settore assicurativo, ragionano, con riferimento all'utilizzo degli *wearable data*, di una nuova prospettiva operativa del tipo assicurativo, definita *behaviour-based insurance*. A tal proposito si vedano M. TANNINEN, *Contested technology: social scientific perspectives of behaviour-based insurance*, in *Big Data and Society*, 2020, 1 ss.; M. V. BEKKUM, F. ZUIDERVEEN, T. HESKES, *AI Insurance, discrimination and unfair differentiation. An overview and research agenda*, in *Law, Innovation and Technology*, 2025, 9-10. Di sicuro interesse sul tema, si legga anche il contributo di EIOPA, *Big Data Analytics in Motor and Health Insurance*, 2019 in part. 19, in cui il regolatore fa riferimento a degli *usage-based insurance products*. Vi è chi, poi, con riferimento a tali strumenti che permettono un monitoraggio continuo dei comportamenti degli assicurati, ha ragionato, in una prospettiva di determinazione precisa del premio, di *behaviour-based personalisation*, come nel caso di G. MEYERS, I. VAN HOYWEGHEN, *Enacting actuarial fairness: from fair discrimination to behaviour based fairness*, in *Science as a Culture*, 27(4) 2017, 413 ss.

110 Come si vedrà *infra*, par. 3.2, dall'analisi dei prodotti distribuiti sui mercati assicurativi europei ed extraeuropei, la prestazione fornita dall'assicurazione non consta esclusivamente nel pagamento di una rendita periodica nei confronti di soggetti che, per patologie pregresse o per l'età avanzata, non sono più autosufficienti. L'impresa può infatti disporre a favore dell'assicurato servizi di cura e assistenza domiciliare, di consulenza personalizzata da parte di un medico specialista, di supporto agli spostamenti dalla propria abitazione a centri di cura, di piani di cura personalizzati disponibili *on-line* o accessibili tramite applicazione sul proprio device personale.

111 Per una breve panoramica dell'impatto che i *wearable data* hanno determinato sull'attività e sul contratto assicurativi, nella prospettiva di analisi dell'impresa, si ponga attenzione al report di MUNICH RE, *The future is now: wearable for insurance risk assessment*, 2018.

Il costante monitoraggio delle *health best practices*¹¹² permette inoltre di sostenere politiche di prezzo premiali a favore dei sottoscrittori. In sostanza, se il cliente mantiene uno stile di vita sano – continuamente sorvegliato¹¹³ –, preserva il proprio stato di salute e diminuisce sensibilmente il rischio di trovarsi entro un breve lasso di tempo in una grave situazione di non autosufficienza. Le condotte virtuose dell'assicurato consentono di beneficiare dunque di sensibili sconti sul premio.

Quanto precede consente altresì di mettere in luce come le ricadute positive derivanti dall'apporto tecnologico dell'intelligenza artificiale assumano una portata sistemica. Una politica produttiva che integri, all'interno del contratto assicurativo, soluzioni di intelligenza artificiale associate all'utilizzo dei wearable data è infatti idonea ad abilitare un approccio prevalentemente preventivo alla gestione del rischio. Come già anticipato, tali modelli contrattuali non si limitano alla copertura dell'evento dannoso, ma incentivano l'adozione di buone pratiche, stimolando gli assicurati al mantenimento di corretti stili di vita; in questa prospettiva, il ricorso a cliniche o strutture ospedaliere viene considerato quale *extrema ratio* nel percorso di assistenza al paziente-assicurato.

I servizi che queste polizze possono offrire – dal più ampio ricorso alla telemedicina, alla consulenza personalizzata, fino alla predisposizione di programmi di cura individualizzati – consentono, in tal modo, ai beneficiari di ricevere assistenza sanitaria in ambito domiciliare, contribuendo contestualmente a ridurre il carico operativo gravante sulle strutture ospedaliere.

Infine, alla luce del progressivo invecchiamento della popolazione¹¹⁴ e del conseguente incremento dei costi che tale fenomeno comporta per la sostenibilità finanziaria del sistema sanitario nazionale, l'elaborazione di soluzioni contrattuali di questo tipo potrebbe configurarsi come uno strumento efficiente ed economicamente sostenibile, idoneo a favorire un godimento più ampio ed effettivo del diritto alla salute.

3.2 Panoramica dei mercati interno, europeo ed extraeuropeo dei prodotti vita che utilizzano sistemi di IA o wearable devices

Venendo ai dati della distribuzione di polizze LTC nel mercato interno, da una prima analisi si è potuto constatare che i principali operatori assicurativi (*manufacturers*), così come i distributori, non hanno ancora realizzato né commercializzato questo tipo di

¹¹² Sul rapporto tra *weareable devices*, monitoraggio della salute degli assicurati e conseguenti incentivi al mantenimento di uno stile di vita adeguato, nuovamente vedasi EIOPA, *ult. op.*, 19-20.

¹¹³ MUNICH RE, *Physical activity data from wearables*, pubblicato il 24/02/2025, disponibile sul sito <https://www.munichre.com/us-life/en/insights/future-of-risk/Next-gen-data-sources-life-underwriting-physical-activity-data-wearables.html>.

¹¹⁴ Sulla relazione tra l'invecchiamento della popolazione e lo sviluppo di soluzioni assicurative *long-term care* – benché nella loro forma "tradizionale", dunque, non alimentata dall'utilizzo dell'IA o dei *wearable data*, si consiglia la lettura del *report* elaborato dal *think tank* Bruegel, a cura di S. HOUGAARD JENSEN, N. RUER, *Long-term policies in practice: a European perspective*, *Working paper*, 21/2024, in part. 4-5.

prodotti. In sintesi, non sono presenti sul mercato italiano né soluzioni contrattuali vita (ramo IV, contro il rischio di non autosufficienza che siano garantite mediante contratti di lunga durata, non rescindibili, per il rischio di invalidità grave dovuta a malattia o a infortunio o a longevità), né danni (art. 2 co. 3, ramo II, c.a.p., Malattia: prestazioni forfetarie; indennità temporanee; forme miste) la cui operatività preveda l'utilizzo di sistemi di intelligenza artificiale e/o di *wearable data*. Le opzioni negoziali LTC attualmente esistenti, di contro, funzionano attraverso sistemi "tradizionali" di *risk assesment* e di denuncia del sinistro¹¹⁵.

Pertanto è parso necessario, nell'intento di trovare un orientamento più saldo e utile per "costruire" un contratto che inglobi i modelli tecnologici di cui si discute, volgere lo sguardo alle soluzioni contrattuali proposte nei mercati europei ed extraeuropei.

In tale prospettiva ricognitiva, va fatto riferimento al *report* dell'Organizzazione per la cooperazione e lo sviluppo economico (OECD), pubblicato nel 2024, dal titolo *Digital tools for health and wellness in insurance*¹¹⁶. Con questo documento l'OECD propone una precisa visuale di come nell'assicurazione vita, e poi anche in quella propriamente sanitaria (cui si ascrivono le polizze LTC, vi sia una sempre maggior tendenza all'utilizzo di modelli predittivi di IA, nonché di supporti tecnologici «indossabili».

L'analisi delle soluzioni contrattuali illustrate dall'OECD consente di rilevare come, nella totalità dei casi esaminati, le prestazioni poste a carico degli assicuratori non si traducano prevalentemente nell'erogazione periodica di una somma di denaro a titolo di rendita al verificarsi dell'evento costitutivo del rischio assicurato, bensì nell'offerta articolata di servizi. In particolare, le imprese assicurative propongono, da un lato, servizi di natura strettamente sanitaria – quali consulenze personalizzate, programmi di alimentazione o di attività fisica predisposti da professionisti qualificati, nonché strumenti di diagnosi a distanza – e, dall'altro lato, servizi che, pur non qualificandosi come prestazioni mediche in senso stretto, risultano funzionali a garantire agli assicurati adeguate condizioni di autonomia e autosostentamento. A titolo esemplificativo, rientrano in tale seconda categoria i servizi di care-giving, volti ad agevolare la mobilità e l'indipendenza dell'assistito nei contesti della vita quotidiana.

115 Dall'analisi dei DIP delle polizze LTC distribuite sul mercato interno, non è emerso l'utilizzo dei supporti tecnologici poc'anzi rammentati. I prodotti LTC analizzati – in particolare, si tratta di Allianz Longevity Care; AXA Assicurazioni Più autonomia; Per una lunga vita di Generali; XME Protezione Intesa Sanpaolo; Per Me Domani di Gruppo Itas assicurazioni; PostaPersona Sempre presente di Poste Italiane; Unipol Autonomia Costante; Fianco a fianco LTC di Vittoria Assicurazioni; Zurich Group LTC – sono assicurazioni sulla vita di ramo IV che prevedono, al verificarsi dell'evento dedotto nel contratto ed espressione del rischio assicurato, la corresponsione di una somma di denaro nei termini di una rendita periodica a favore dell'assicurato.

116 OECD, *Digital tools for health and wellness in insurance*, 2024, reperibile al link https://www.oecd.org/en/publications/digital-tools-for-health-and-wellness-in-insurance_d3764184-en.html.

Di seguito, si illustrano alcune delle proposte fornite dai principali *manufacturers* a livello continentale e mondiale.

1. Prudential PLC – Pulse: La compagnia inglese Prudential PLC ha elaborato un prodotto assicurativo che si avvale del supporto di un'applicazione sviluppata dall' *health data provider* Babylon Health, nonché insieme ad altre imprese come Doctoroncall – società specializzata nella consulenza medica *online* – e grazie al software Fictrac che consente il monitoraggio e lo studio del corpo attraverso una telecamera che visualizza prospetticamente il movimento umano. L'applicazione offre una serie di funzionalità, consistenti in una valutazione dello stato di salute psico-fisica dell'assicurato e in una visita anatomica in 3D del corpo. L'applicazione fornisce, tramite *chatbot*, una consulenza su come affrontare i rischi sanitari riscontrati e specifica, inoltre, se sia necessario affidarsi alle cure di un medico a fronte di determinati sintomi. Non può essere definito propriamente come polizza LTC in senso stretto, ma il sottostante tecnologico che ne garantisce l'operatività non può essere trascurato. Sul sito della compagnia non sono reperibili documenti informativi, ad eccezione di un documento pubblicitario che descrive il funzionamento in concreto del contratto¹¹⁷.

2. Zurich Australia – LiveWell: come nel caso di Prudential PLC, anche in questa ipotesi sarebbe più corretto ritenere il contratto in questione come un'assicurazione sulla vita, e segnatamente un'assicurazione caso malattia, in cui non può passare inosservata l'integrazione di elementi tecnologici. In particolare, la struttura contrattuale richiede l'ausilio di un'applicazione che consente il caricamento dei dati relativi all'attività fisica dell'utente, appresi da dispositivi indossabili. La stessa applicazione fornisce consulenze e consigli personalizzati su un corretto stile di vita e sulle abitudini alimentari da seguire per prevenire il sopraggiungere delle malattie. Sul sito della compagnia non sono reperibili documenti informativi, ad esclusione di un documento pubblicitario che descrive l'operatività della polizza¹¹⁸.

3. Nan Shan Life Insurance Co. Ltd: questa compagnia cinese ha sviluppato un prodotto che prevede la collaborazione con Home Angel – piattaforma di *matchmaking* per l'assistenza domiciliare e ospedaliera – e che consente di fornire ai pazienti assicurati un servizio di *care-giving* via app. Quanto ai prodotti LTC di Nanshan, i documenti informativi reperiti sul sito sono in lingua cinese. Pertanto, nonostante il tentativo di contatto con l'ufficio clienti, è risultato impossibile comprendere come un (eventuale) modello di AI e l'utilizzo di *smart devices* siano stati concretamente inseriti all'interno della polizza.

4. Ping An – Ping An Family doctor: la compagnia cinese ha elaborato un prodotto che

¹¹⁷ Il documento è disponibile al seguente indirizzo: <https://www.prudentialplc.com/~media/Files/P/Prudential-V13/presentations/2020/pulse-by-prudential-investor-presentation-nov-2020.pdf>.

¹¹⁸ Il documento è disponibile al seguente indirizzo: <https://edge.sitecorecloud.io/zurichinsur0b40-zwpaustali7d46-prod0760-b10c/media/project/zurich-headless/shared/content/dam/au-documents/livewell/livewell-brochure.pdf>.

prevede una prestazione di assistenza medica agli anziani e un servizio di *conciergerie* personalizzata che faciliti tali soggetti vulnerabili nelle azioni quotidiane. Con questo prodotto l'impresa fornisce una serie di servizi consulenziali e di assistenza, con un *team* medico dedicato per ogni paziente. Per i pazienti con malattie croniche, l'impresa ha pensato a un servizio che combina il monitoraggio quotidiano, tramite *IoT*, con comunicazioni volte ad avvertire circa i rischi legati allo stato di salute e con la fornitura di un programma per gestire al meglio lo stile di vita dell'assicurato.

Come per il prodotto realizzato da Nan Shan Life Insurance Co. Ltd, in attesa di un riscontro dalla compagnia, non sono disponibili documenti informativi in lingua inglese che consentano di comprendere l'inserimento di modelli di IA e di *smart devices* nella struttura contrattuale.

A seguito di queste considerazioni sui dati della prassi commerciale relativa alla distribuzione di prodotti che utilizzano l'IA e/o dispositivi intelligenti che assumono (ed eventualmente rielaborano) i dati biometrici relativi agli assicurati, uno spazio deve essere riservato ad alcune conclusioni finali.

In primo luogo, il riscontro fattuale dei dati attinenti ai prodotti distribuiti ci consente di affermare l'assenza sui mercati interno, europeo ed extra-europeo, di polizze LTC che utilizzino i supporti tecnologici di cui abbiamo discusso finora. Più precisamente, nell'odierno contesto produttivo il coordinamento tra sistemi di IA e *wearable data* non è parte integrante dei prodotti vita ramo IV ex art. 2 c.a.p.. Attualmente, cioè, non esistono soluzioni negoziali in cui l'IA sia utilizzata per la determinazione del premio né per la gestione integrata del servizio¹¹⁹.

In una prospettiva futura, non si esclude che possano essere proposte delle coperture che integrino l'utilizzo dei *devices* indossabili come fonte primaria e continuativa di

¹¹⁹ Cioè non esistono contratti in cui l'IA consente la presa in carico automatizzata dell'evento rappresentativo del rischio, una valutazione clinica dello stato di salute del cliente e la successiva fase di liquidazione della rendita periodica.

raccolta dei dati e come sistema di monitoraggio del comportamento dell'assicurato¹²⁰. Al contempo l'AI verrebbe assunta come strumento di «lavorazione» e interpretazione delle informazioni, utile a determinare – tanto a monte quanto durante l'esecuzione del programma negoziale – il giusto prezzo contrattuale¹²¹.

¹²⁰ Una puntualizzazione appare necessaria in questa sede e si ritiene utile collocarla in nota. Nella panoramica delle assicurazioni – sia vita che danni – contro il rischio di non autosufficienza non sono presenti prodotti che intercettino il rischio assicurato – l'evento costitutivo del rischio o il cui livello, una volta raggiunto, rappresenti la soglia minima di rischio – attraverso *wearable data*. Al contrario, però, esistono alcuni prodotti in cui le imprese adoperano degli strumenti – i cc.dd. oracoli, che possono essere degli enti terzi certificati, ma anche dei *devices* che rilevano le soglie minime del *trigger event* – che assumono e rielaborano i dati. Si pensi ai casi del ramo 16 danni, laddove vengono offerte delle polizze contro i danni da ritardo del volo, come nel caso dei prodotti *Revo ParametricXFlight Delay* (di cui si riporta il link al DIP: <https://travel.b2b.i4t.it/download/documento?uid=7Pk1dWUTeOPtjOHYtAh5Au0lJQ2dSJB8>) e *Revo Travel Care-Medico Bagaglio* (di cui il link al DIP: <https://travel.b2b.i4t.it/download/documento?uid=RonluEeh-hXpqkrPzX2Mn998xuJL-WaN>). In un certo senso, potrebbero rientrare in questo tipo di prodotti anche le polizze ramo 3 danni contro i furti in abitazione, laddove le compagnie prevedano l'acquisto da parte dell'assicurato di un sensore collegato a un'applicazione che consente di individuare all'istante l'evento oggetto di copertura, come nel caso della polizza *Homix Smart Protection* di NET Insurance (di cui si riporta il DIP: https://www.netinsurance.it/wp-content/uploads/2021/07/SET_INFORMATIVO_HOMIX_SMART_PROTECTION_336.pdf).

¹²¹ Con riferimento all'utilizzo di modelli di IA, la loro diffusione nella catena produttiva del settore assicurativo si è dimostrata, al momento, di un certo rilievo solo nel momento del contatto con i potenziali clienti. In particolare in quelle attività di *front end* attinenti all'accesso ai servizi sui siti *internet* delle compagnie, come la navigazione guidata sui portali *on-line* e l'assistenza digitale per mezzo di *chatbot*.