



Una legge sull'intelligenza artificiale che includa la disabilità

Di Kave Noori, responsabile delle politiche in
materia di intelligenza artificiale del FES |
Ottobre 2025

Una guida per monitorare l'attuazione nel tuo paese

Gli sforzi di advocacy e capacity building di EDF
sull'Intelligenza Artificiale sono possibili solo grazie alle
generose sovvenzioni che abbiamo ricevuto dal Fondo

European
**Artificial Intelligence
& Society Fund**

Informazioni sul Forum europeo sulla disabilità	1
Scopo di questo toolkit	1
Informazioni su questa versione del toolkit	1
Riconoscimenti	2
Sintesi	2
1. Panoramica della legge sull'IA	5
1.1 Introduzione alla legge sull'IA	5
1.2 Che cos'è un regolamento? Come viene attuato?	5
1.3 La base giuridica per l'uso dei sistemi di IA	6
1.4 Impatto: cosa fa l'AI Act per le persone con disabilità?	6
1.5 La definizione giuridica di un sistema di IA	8
2. Comprendere le categorie di rischio	8
2.1 Panoramica delle categorie di rischio	9
2.2 Rischio inaccettabile	12
3. IA ad alto rischio	25
3.1 Banca dati di IA ad alto rischio	26
3.2 Requisito di accessibilità per l'IA ad alto rischio	26
3.3 Pregiudizio	27
3.4 Il diritto di ottenere una spiegazione	29
3.5 Categorie di IA ad alto rischio	29
4. Rischio di trasparenza	43
5. Rischio minimo	44
6. IA per uso generale (GPAI)	45
6.1 Protezioni dalle minacce alla salute, alla sicurezza e alla democrazia umana	46
6.2 Orientamenti sull'ambito di applicazione degli obblighi per i modelli di IA per uso generale	47
6.3 Codice di condotta GPAI	48
6.4 Codici di condotta volontari	49
6.5 Raccomandazioni per i membri del FES	49
7. Diritti fondamentali, sandbox e test	50
7.1 Il diritto di reclamo	50
7.2 Segnalazione di incidenti gravi	51

7.3	Sandbox normativi	52
7.4	Valutazioni d'impatto sui diritti fondamentali (FRIA).....	53
7.5	Test in condizioni reali.....	55
8.	Misure supplementari: «Incentivazione»	55
9.	Struttura di governance.....	57
9.1	Autorità di vigilanza del mercato	58
9.2	Autorità che tutelano i diritti fondamentali	61
9.3	Sanzioni per la violazione degli obblighi dell'AI Act.....	62
10.	Monitoraggio della legge sull'IA nel tuo paese	63
10.1	Importanza del coinvolgimento delle persone con disabilità....	63
10.2	Tempistica dettagliata dell'attuazione dell'AI Act:	64
10.3	Cosa bisogna fare? (Misure da adottare per l'attuazione a livello nazionale).....	65
10.4	Raccomandazioni ai membri del FES	65
11.	Valutazione delle carenze della legge sull'IA	68
11.1	La scelta di disegnare la legge in base ai rischi invece che ai diritti	68
11.2	Accessibilità nell'intelligenza artificiale GPAI e rischio minimo .	69
12.	Glossario	70
13.	Crediti del documento	77

Nota: un glossario dei termini chiave è incluso alla fine di questo toolkit (vedi pagina 70). Se incontri parole o concetti sconosciuti, fai riferimento al glossario per spiegazioni chiare.

Informazioni sul Forum europeo sulla disabilità

Il Forum Europeo sulla Disabilità (EDF) è una ONG indipendente che difende i diritti di 100 milioni di europei con disabilità. EDF è una piattaforma unica nel suo genere che riunisce organizzazioni rappresentative di persone con disabilità di tutta Europa. EDF è gestito da persone con disabilità e dalle loro famiglie. Siamo una voce forte e unita delle persone con disabilità in Europa.

Scopo di questo toolkit

Questo kit di strumenti fornisce orientamento e sostegno alle organizzazioni di persone con disabilità nell'Unione europea. Si rivolge principalmente ai membri dell'EDF, ma tutte le organizzazioni di persone con disabilità sono incoraggiate a sostenere un'IA inclusiva della disabilità.

Spiega come contribuire all'attuazione e al monitoraggio della legge sull'intelligenza artificiale (AI Act) e come difendere i diritti e gli interessi delle persone con disabilità nell'ambito dell'intelligenza artificiale (AI).

Informazioni su questa versione del toolkit

Questa è la versione 2 del toolkit. Si basa sull'AI Act, una legge entrata in vigore nell'Unione Europea nell'agosto 2024. Sebbene la legge sull'IA sia stata adottata, alcuni dettagli devono ancora essere chiariti. La Commissione europea sta pubblicando linee guida e atti delegati che integrano e chiariscono alcune disposizioni della legge sull'IA. Potrebbero cambiare alcune cose che sono importanti per le organizzazioni di persone con disabilità. Continueremo ad aggiornare il toolkit quando ciò accadrà. Vogliamo darti le informazioni migliori e più utili.

La versione 2 di questo toolkit è stata aggiornata per incorporare le linee guida più recenti e altri documenti della Commissione europea relativi all'AI Act, che sono stati rilasciati dopo la pubblicazione della versione 1 nell'ottobre 2024. Questi sono:

- [Linee guida sull'IA vietata](#) (per saperne di più alla voce 2.2 Rischio inaccettabile, a partire da pagina 12 sotto)
- [Linee guida sulla definizione dei sistemi di IA](#) (per saperne di più alla voce 1.5 La definizione giuridica di un sistema di IA, a pagina 8 sotto)
- [Orientamenti sull'ambito di applicazione degli obblighi per i sistemi di IA per uso generale](#) (per saperne di più alla voce 4.2)

Orientamenti sull'ambito di applicazione degli obblighi per i modelli di IA per uso generale, a pagina 47 sotto)

- [Progetto di linee guida per la segnalazione di incidenti gravi](#) sistemi di IA ad alto rischio (per saperne di più 6.2 Segnalazione di incidenti gravi a pagina 51 sotto)
- [Il codice di condotta generale sull'IA](#) (Vedi 4.3 Codice di condotta GPAI, a pagina 48 sotto)

Riconoscimenti

Il Forum europeo sulla disabilità (EDF) desidera ringraziare tutte le organizzazioni dei disabili, il nostro gruppo di esperti sulle TIC e altri per il loro feedback su questa bozza di toolkit. Ringraziamo anche la coalizione EDRi di organizzazioni per i diritti umani e per i diritti digitali per il loro sostegno nei nostri sforzi di advocacy.

Sintesi

Il toolkit del Forum europeo sulla disabilità fornisce orientamenti alle organizzazioni di persone con disabilità sull'attuazione e il monitoraggio della legge sull'intelligenza artificiale dell'Unione europea (legge sull'intelligenza artificiale) nel loro paese dell'UE.

Il toolkit offre una panoramica dettagliata della legge sull'IA, evidenziando le sezioni importanti per i diritti delle persone con disabilità.

Fornisce orientamenti pratici sulla navigazione del processo legislativo e suggerisce strategie per un'efficace advocacy al fine di garantire che le prospettive delle persone con disabilità siano incluse nella creazione e nell'esecuzione delle politiche nazionali in materia di IA.

Questo toolkit è un documento in continua evoluzione e verrà aggiornato con nuove informazioni non appena saranno disponibili. Rimanendo informati e adottando misure proattive, possiamo creare un panorama inclusivo dell'IA al servizio di tutte le persone, comprese le persone con disabilità. Contiamo sulla tua partecipazione e sul tuo feedback per garantire che la legge sull'IA sia attuata in modo rispettoso della disabilità.

Panoramica delle azioni chiave per l'advocacy nazionale

1. Comprendere e coinvolgere il processo legislativo:

I membri dell'EDF dovrebbero partecipare attivamente al processo legislativo fornendo riscontri e suggerimenti alle autorità nazionali responsabili dell'attuazione della legge sull'IA. Ciò include la

partecipazione a consultazioni pubbliche, la presentazione di documenti di sintesi e l'impegno con i responsabili politici per garantire che le esigenze e i diritti delle persone con disabilità siano presi in considerazione.

2. Monitorare l'attuazione della legge sull'IA:

I membri dell'EDF dovrebbero monitorare l'attuazione della legge sull'IA a livello nazionale. Ciò comporta tenere traccia del modo in cui le autorità nazionali applicano le disposizioni della legge sull'IA e garantire che i sistemi di IA ad alto rischio siano conformi ai requisiti di accessibilità. A partire da agosto 2026 le organizzazioni possono monitorare la banca dati dell'UE dei sistemi di IA ad alto rischio per identificare quelli utilizzati nel loro Stato membro. Ciò può contribuire a creare alleanze con altre organizzazioni che tutelano i diritti fondamentali per raccogliere testimonianze di sistemi che non funzionano e fornirle alle autorità di controllo.

3. Sostenitore dell'accessibilità e dei test di pregiudizio:

I membri dell'EDF dovrebbero sostenere i requisiti obbligatori di accessibilità per i sistemi di IA ad alto rischio, come stabilito nella legge sull'IA. Ciò include la garanzia che gli sviluppatori di IA seguano gli standard europei armonizzati in materia di accessibilità.

Inoltre, i membri dell'EDF dovrebbero spingere per test approfonditi sui pregiudizi dei sistemi di IA ad alto rischio per prevenire la discriminazione nei confronti delle persone con disabilità. Ciò comporta la garanzia che i dati utilizzati per addestrare, convalidare e testare questi sistemi siano affidabili, sicuri e privi di distorsioni.

4. Sensibilizzare ed educare le parti interessate:

I membri dell'EDF dovrebbero sensibilizzare l'opinione pubblica sulle implicazioni della legge sull'IA per le persone con disabilità attraverso workshop, seminari e campagne pubbliche. È inoltre fondamentale fornire formazione e risorse alle persone con disabilità e alle loro famiglie, consentendo loro di difendere i propri diritti nel contesto dell'IA.

5. Esercita il diritto di spiegazione:

L'AI Act garantisce alle persone il diritto di ottenere una spiegazione dalle organizzazioni che utilizzano sistemi di IA ad alto rischio se una decisione le riguarda legalmente o potrebbe avere un impatto sulla loro salute, sicurezza o diritti fondamentali. I membri dell'EDF dovrebbero informare e sostenere le persone con disabilità nell'esercizio di tale diritto.

6. Collaborare con altre organizzazioni che tutelano i diritti fondamentali:

Costruire alleanze con altre organizzazioni che tutelano i diritti fondamentali, come le organizzazioni dei pazienti, può rafforzare gli sforzi di advocacy. Lavorando insieme, i membri dell'EDF possono raccogliere testimonianze e fornire prove alle autorità di vigilanza in merito ai sistemi di IA che non funzionano come previsto.

1. Panoramica della legge sull'IA

1.1 Introduzione alla legge sull'IA

L' [AI Act](#) è un regolamento dell'Unione Europea che mira a creare un quadro di regole e standard per la creazione, l'attuazione e l'applicazione di sistemi di intelligenza artificiale all'interno dei confini dell'UE. I suoi obiettivi principali sono proteggere i diritti umani fondamentali, garantire la sicurezza del pubblico e promuovere la fiducia e l'innovazione nelle tecnologie di IA.

1.2 Che cos'è un regolamento? Come viene attuato? ¹

L'UE legifera in modi diversi. Queste leggi possono assumere la forma di regolamenti, direttive o decisioni.

- **Un regolamento dell'UE**, come l'AI Act, diventa automaticamente legge in tutti i paesi dell'UE. Si applica a tutti – governi, imprese e persone – senza che ogni paese debba approvare la propria versione della legge.
- **Una direttiva dell'UE** è diversa. Stabilisce obiettivi che ogni paese deve raggiungere, ma consente a ciascun paese di decidere come emanare le proprie leggi nazionali per raggiungere tali obiettivi.
- **Una decisione dell'UE** è una legge specifica che si applica solo a determinate persone o organizzazioni. Ad esempio, l'UE può prendere una decisione che riguarda solo un'impresa o un gruppo specifico di persone.

Nel caso dell'AI Act, si tratta di un regolamento dell'UE. Ciò significa che si applica a tutti i paesi dell'UE allo stesso modo. Tuttavia, l'AI Act consente ai paesi di stabilire le proprie regole su determinate questioni. Ad esempio, stabilisce dei limiti all'uso della tecnologia di riconoscimento facciale da parte della polizia, ma lascia anche a ciascun paese la facoltà di decidere se vietare del tutto questa tecnologia.

Il **Parlamento europeo e il Consiglio dell'Unione europea** possono adottare decisioni dell'UE come atti giuridici. Questi atti giuridici sono leggi che si applicano a tutti, come i regolamenti o le direttive.

La **Commissione europea** può anche emanare documenti denominati atti delegati e atti di esecuzione. Non si tratta di leggi, ma aiutano a spiegare o applicare le leggi esistenti come l'AI Act. La Commissione

¹ [«Decisioni dell'Unione europea | EUR-Lex»](#) [consultato il 4 ottobre 2024].

utilizzerà questi documenti, insieme agli orientamenti, per chiarire e fornire ulteriori dettagli sulla legge sull'IA.

1.3 La base giuridica per l'uso dei sistemi di IA

In generale, l'AI Act non concede alle organizzazioni l'autorizzazione all'uso dei sistemi di IA. Invece, funziona come una legge sulla sicurezza dei prodotti. **Stabilisce le regole su come l'IA può essere utilizzata, ma solo dopo che un'organizzazione è già autorizzata a utilizzarla ai sensi di altre leggi.**

Ad esempio, l'AI Act classifica l'uso del riconoscimento facciale da parte della polizia come "ad alto rischio". Tuttavia, ciò non significa che l'AI Act autorizzi la polizia a utilizzare il riconoscimento facciale. La polizia può utilizzarlo solo se le leggi del loro paese o altre leggi dell'UE lo autorizzano esplicitamente. Una volta ottenuta questa autorizzazione, l'AI Act stabilisce i limiti esterni entro i quali la polizia può utilizzare la tecnologia per garantire che sia utilizzata in sicurezza e non si spinga troppo oltre nel violare i diritti dei cittadini.

1.4 Impatto: cosa fa l'AI Act per le persone con disabilità?

L'AI Act stabilisce norme per garantire che i sistemi di intelligenza artificiale siano utilizzati **in modo sicuro ed etico**, con particolare attenzione alla tutela dei diritti umani.

Ciò è particolarmente importante per le persone con disabilità e altri gruppi emarginati, poiché l'AI Act mira a **prevenire danni** come la discriminazione, il trattamento ingiusto o la perdita della privacy.

La legge sull'IA classifica i sistemi di IA in base ai rischi che comportano, ma queste categorie non sempre si escludono a vicenda. Alcuni sistemi possono rientrare in più di una categoria a seconda del loro utilizzo e del contesto. Maggiori informazioni sono fornite nella sezione 2 "Comprendere le categorie di rischio".

Le principali categorie sono:

- **Pratiche vietate di IA (rischio inaccettabile):** i sistemi di IA considerati troppo pericolosi o incompatibili con i valori e i diritti fondamentali europei sono completamente vietati. Ne sono un esempio i sistemi che manipolano il comportamento utilizzando tecniche subliminali, sfruttano le vulnerabilità dovute all'età o alla disabilità o consentono un punteggio sociale da parte delle autorità pubbliche.

- **Sistemi ad alto rischio:** sistemi di intelligenza artificiale consentiti ma strettamente regolamentati perché possono influenzare in modo significativo le possibilità di vita di una persona (ad esempio, decisioni sull'occupazione, l'accesso a prestiti o l'ammissione all'istruzione).
- **Rischio di trasparenza** (a volte chiamato "rischio limitato"): alcuni sistemi di intelligenza artificiale che interagiscono con le persone (come i chatbot), generano o manipolano contenuti (come i deepfake) o utilizzano la categorizzazione biometrica (determinare l'età, il sesso e la razza di una persona senza tentare di rivelarne l'identità) sono soggetti a requisiti di trasparenza e informazione. È importante sottolineare che un sistema può essere soggetto a obblighi di trasparenza ed essere classificato come ad alto rischio se soddisfa i criteri per entrambi.
- **Rischio basso o nullo:** i sistemi di IA con un impatto minimo sui diritti delle persone (come i filtri antispam o l'IA dei videogiochi) non sono soggetti a obblighi aggiuntivi ai sensi della legge sull'IA.

È anche importante considerare la "General Purpose AI" (GPAI).

L'analisi GPAI non fa parte della gerarchia di classificazione del rischio. Invece, i modelli e i sistemi GPAI possono rientrare in una qualsiasi delle categorie di cui sopra a seconda del loro utilizzo e impatto. La GPAI è soggetta a norme specifiche, soprattutto se rappresenta un rischio per la società, ad esempio se può essere utilizzata in modo improprio per produrre disinformazione.

Queste categorie di rischio si applicano a vari settori della vita. Ad esempio, nell'istruzione, l'intelligenza artificiale può influenzare decisioni come chi viene ammesso a un programma o come vengono valutati gli studenti. Nei luoghi di lavoro, l'intelligenza artificiale potrebbe essere utilizzata per prendere decisioni di assunzione o valutare le prestazioni dei dipendenti. Nel settore sanitario, l'intelligenza artificiale assiste i medici nelle diagnosi e nelle scelte terapeutiche. Nella migrazione e nell'applicazione della legge, l'intelligenza artificiale può essere impiegata per elaborare le domande di visto o monitorare gli spazi pubblici. Per ciascuno di questi settori, l'AI Act impone regole rigorose per proteggere i diritti delle persone e prevenire i danni.

La legge sull'IA prevede inoltre che i sistemi di IA ad alto rischio siano accessibili a tutti, comprese le persone con disabilità, come descritto all'[articolo 16, lettera I](#)). Gli sviluppatori di IA sono tenuti ad applicare i principi della progettazione universale fin dall'inizio per garantire che

questi sistemi siano accessibili e utilizzabili da persone con abilità ed esigenze diverse.

[Il considerando 80](#) sottolinea l'obbligo giuridico dell'UE, in quanto firmataria della Convenzione delle Nazioni Unite sui diritti delle persone con disabilità (UNCPRD), di proteggere le persone con disabilità dalla discriminazione e di garantire la parità di accesso alle tecnologie dell'informazione e della comunicazione. Sottolinea la necessità di applicare i principi della progettazione universale per garantire un accesso pieno e paritario per tutti, tenendo conto in particolare della dignità e della diversità delle persone con disabilità.

In sintesi, l'AI Act cerca di garantire che i sistemi di IA siano progettati e utilizzati in modo da proteggere i diritti di tutti, compresi quelli delle persone con disabilità, rendendo i sistemi equi, accessibili e privi di discriminazioni.

1.5 La definizione giuridica di un sistema di IA²

La Commissione ha pubblicato orientamenti per aiutare le autorità, i fornitori e gli utenti a decidere se un determinato strumento si qualifica come "sistema di IA" ai sensi dell' [articolo 3, paragrafo 1](#). [Le linee guida sulla definizione dei sistemi di IA](#) non sono giuridicamente vincolanti.

Prima di valutare i divieti, gli obblighi ad alto rischio o i requisiti di trasparenza, è necessario verificare [l'articolo 3, paragrafo 1](#): "**Lo strumento è in grado di produrre risultati (ad esempio previsioni, raccomandazioni, decisioni) e di operare con un certo grado di autonomia/adattabilità?**"

In caso affermativo, probabilmente si tratta di un sistema di IA e potrebbe essere applicata la legge sull'IA; in caso negativo, le disposizioni della legge non si applicano. La definizione giuridica ti aiuterà a decidere se procedere ai sensi dell'AI Act o perseguire altre strade, come la presentazione di un reclamo ai sensi della legge sull'uguaglianza o della legge sulla protezione dei dati.

2. Comprendere le categorie di rischio

La legge sull'IA è redatta utilizzando un **approccio basato sul rischio**. Le sue disposizioni si basano sull'idea che maggiore è il rischio per i diritti fondamentali, la salute e la sicurezza posti da un particolare sistema di IA

² [Commissione europea, la Commissione pubblica le linee guida sulle pratiche vietate in materia di intelligenza artificiale \(IA\), come definite dall'AI Act \(2025\)](#) [consultato il 26 giugno 2025].

o da un particolare caso d'uso, come l'uso dell'IA per l'assunzione o la decisione di un sussidio pubblico, più rigorose devono essere le norme.

2.1 Panoramica delle categorie di rischio

1. Rischio inaccettabile

- **Cosa significa:** sistemi di IA che sono completamente vietati perché rappresentano una grave minaccia per i diritti umani o la sicurezza. In casi specifici possono essere applicate eccezioni.
- **Che cosa è coperto:** sistemi di IA che rappresentano un rischio inaccettabile per la dignità umana, la democrazia e i diritti fondamentali, anche per le persone con disabilità.
- **Dove è disciplinato dalla legge sull'IA:** [articolo 5](#)
- **Esempi:**
 - Sistemi di intelligenza artificiale che manipolano il comportamento umano o sfruttano le vulnerabilità.
 - Sistemi di intelligenza artificiale che causano danni fisici o psicologici.
 - L'intelligenza artificiale che riconosce le emozioni per influenzare le decisioni.
 - L'intelligenza artificiale prevede futuri comportamenti criminali basati esclusivamente sulla personalità o sulle caratteristiche personali.

2. Alto rischio

- **Cosa significa:** questi sistemi di intelligenza artificiale sono consentiti ma sono soggetti a rigide normative e supervisione per proteggere la sicurezza, la dignità umana, la democrazia e i diritti fondamentali.
- **Cosa è coperto:** sistemi di intelligenza artificiale che influenzano in modo significativo la vita, i diritti o il benessere delle persone. Esiste un elenco di casi d'uso definiti come sistemi di intelligenza artificiale ad alto rischio.
- **Dove è disciplinato dalla legge sull'IA:** [articoli da 6 a 27](#) e [allegati I e III](#)
- **Esempi:**

- Sistemi di intelligenza artificiale utilizzati nei processi di assunzione, come le decisioni di assunzione o licenziamento.
- L'intelligenza artificiale determina l'accesso ai servizi pubblici essenziali o alle prestazioni statali.
- L'intelligenza artificiale utilizzata nel credit scoring o nella valutazione dell'affidabilità creditizia.
- Sistemi di IA nelle forze dell'ordine per la valutazione dei rischi.

3. Rischio limitato Cosa significa: i sistemi di IA a rischio limitato sono soggetti a requisiti di trasparenza ai sensi dell'AI Act. Questi sistemi devono informare gli utenti che stanno interagendo con l'IA o che i contenuti sono stati generati o manipolati dall'IA (ad esempio, chatbot o deepfake). Gli obblighi sono meno rigorosi rispetto all'IA ad alto rischio e si concentrano principalmente sulla trasparenza e sulla consapevolezza degli utenti

- **Cosa è coperto:** i sistemi di intelligenza artificiale che interagiscono con le persone (come i chatbot), generano o manipolano contenuti (come testo, immagini, audio o video) o utilizzano la categorizzazione biometrica, a condizione che non vengano utilizzati in un contesto ad alto rischio.
- **Dove è coperto dalla legge sull'IA:**
 - [Articolo 50](#) (Obblighi di trasparenza per taluni sistemi di IA)
- **Esempi:**
 - Chatbot utilizzati per il servizio clienti (devono indicare che gli utenti interagiscono con l'intelligenza artificiale)
 - Strumenti di intelligenza artificiale che generano o manipolano immagini o video (come i deepfake, che devono essere chiaramente etichettati come generati dall'intelligenza artificiale)
 - Sistemi di categorizzazione biometrica utilizzati al di fuori di contesti ad alto rischio

4. Rischio minimo

- **Cosa significa:** questi sistemi di IA sono autorizzati a funzionare senza ulteriori norme ai sensi della legge sull'IA. Devono comunque rispettare altre leggi pertinenti, come le normative sulla protezione dei dati.

- **Cosa è coperto:** applicazioni di IA che non rientrano nelle altre categorie di rischio.
- **Dove è coperto dalla legge sull'IA:** nessuna disposizione specifica; menzionato in contesti come i codici di condotta.
- **Esempi:**
 - L'intelligenza artificiale viene utilizzata per filtrare le e-mail di spam.
 - Riconoscimento vocale
 - Algoritmi di intelligenza artificiale nei videogiochi che migliorano l'esperienza del giocatore.
 - Sistemi di raccomandazione per film o musica.

IA per uso generale (GPAI)

Nota: l'IA per uso generale (GPAI) non è una categoria di rischio. Al contrario, i modelli e i sistemi GPAI possono rientrare in qualsiasi categoria di rischio - proibito, ad alto rischio, a rischio limitato (trasparenza) o a rischio minimo - a seconda del loro utilizzo e del contesto. La GPAI è soggetta a obblighi specifici ai sensi della legge sull'IA, in particolare se presenta un rischio sistemico (ad esempio, modelli su larga scala con un impatto significativo). L'applicazione della legge per l'GPAI è gestita dalla Commissione europea, non dalle autorità nazionali.

- **Dove è coperto dalla legge sull'IA:**
 - [Articoli 50, 53,54](#) (Tutti i modelli di IA per uso generale)
 - Articoli [51](#), [52](#) e [55](#) (IA per uso generale con rischio sistemico)
- **Esempi:**
 - Modelli linguistici come [ChatGPT](#), [Google Gemini](#), [DeepSeek](#), [Mistral](#) e [Llama](#).
 - Modelli GPAI utilizzati per alimentare chatbot, strumenti di generazione di contenuti o altre applicazioni

Le categorie di rischio non si escludono a vicenda, quindi un sistema può soddisfare i requisiti sia della categoria di rischio di trasparenza che dei criteri di classificazione dell'IA ad alto rischio. La gerarchia è una semplificazione di queste classificazioni. Ad esempio, alcuni sistemi di IA utilizzati nella categorizzazione biometrica possono rientrare nella categoria di rischio limitato, ma possono anche essere considerati ad alto rischio in determinate circostanze. Inoltre, i chatbot sono generalmente assegnati alla categoria di rischio limitato, che richiede loro di rivelare che

gli utenti stanno interagendo con un programma informatico. Tuttavia, se un chatbot utilizza una tecnologia linguistica avanzata, come un modello linguistico come ChatGPT, è anche soggetto a normative generali sull'intelligenza artificiale. Inoltre, se un chatbot viene utilizzato nei processi di reclutamento per intervistare i candidati, verrebbe contemporaneamente classificato come un sistema di intelligenza artificiale ad alto rischio.

2.2 Rischio inaccettabile

La legge sull'IA vieta i sistemi di IA che rappresentano un rischio inaccettabile per la dignità umana, la democrazia e i diritti fondamentali, anche per le persone con disabilità.

Si tratta di sistemi in grado di manipolare il comportamento umano, sfruttare le vulnerabilità delle persone o causare danni fisici o psicologici. Esempi di tali sistemi includono:

- Sistemi di intelligenza artificiale che riconoscono le emozioni sul posto di lavoro o a scuola.
- Sistemi di intelligenza artificiale che raccolgono o espandono i database di riconoscimento facciale da Internet o dalle telecamere di sorveglianza.
- Sistemi di intelligenza artificiale che prevedono comportamenti criminali basandosi esclusivamente sui tratti della personalità o sulle caratteristiche personali.
- Sistemi di intelligenza artificiale che sfruttano le persone con disabilità a scopo di lucro o di controllo.

Eccezioni

Ci sono alcune eccezioni specifiche a questi divieti. Ad esempio, l'intelligenza artificiale che riconosce le emozioni sul posto di lavoro o nelle scuole è vietata, a meno che non venga utilizzata per motivi medici o di sicurezza. Inoltre, sebbene l'intelligenza artificiale per il riconoscimento delle emozioni non sia consentita nelle scuole, può comunque essere utilizzata dalla polizia o dalle autorità per l'immigrazione. Ciò significa che le persone con disabilità che hanno un background migratorio potrebbero comunque essere influenzate da tali tecnologie.

Ci sono anche alcune scappatoie. L'UE non ha il potere di legiferare in materia di sicurezza nazionale. Pertanto, gli Stati membri possono talvolta sostenere che l'uso di una particolare IA è correlato alla sicurezza nazionale, consentendo loro di esentarla dalle norme della legge sull'IA dell'UE.

2.2.1 Divieto di IA che sfrutta le persone con e senza disabilità

Due disposizioni sono particolarmente rilevanti per le persone con disabilità quando si tratta di proteggerle dal fare cose che altrimenti non farebbero. [Articolo 5, paragrafo 1, lettera a\)](#), e [articolo 5, paragrafo 1, lettera b\)](#). [Questa sezione del toolkit è stata aggiornata sulla base delle linee guida della Commissione europea del 2025 sulle pratiche vietate di IA](#), che forniscono spiegazioni chiare ed esempi reali di come funzionano nella pratica questi divieti.

[Articolo 5, paragrafo 1, lettera a\)](#): Il divieto generale di manipolazioni dannose³

Questa disposizione vieta i sistemi di IA che utilizzano tecniche nascoste, manipolative o ingannevoli per influenzare fortemente il comportamento di una persona (termine giuridico: "distorcere materialmente il comportamento"). Ciò significa che l'IA rende molto più difficile per una persona prendere una decisione ben informata (termine giuridico: "compromettere sensibilmente la sua capacità di prendere una decisione informata"), la induce a fare qualcosa che altrimenti non avrebbe fatto e le causa (o è probabile che le causi) un grave danno (termine legale: "danno significativo"). Maggiori informazioni su ciò che costituisce un danno sono fornite nelle pagine seguenti.

- Chi è protetto? Tutti, comprese le persone con disabilità.
- Cosa è vietato? Un'intelligenza artificiale che inganna, manipola o inganna le persone in un modo che causa loro danni reali.

[Articolo 5, paragrafo 1, lettera b\)](#): Il divieto speciale di sfruttare le vulnerabilità⁴

Questa disposizione vieta i sistemi di IA che sfruttano la vulnerabilità di una persona a causa della sua età (ad esempio bambini e anziani), della sua disabilità o della sua particolare situazione sociale o economica. La disposizione tutela esclusivamente le persone che appartengono a uno di questi gruppi.

Questi ultimi gruppi sono, ad esempio, le persone che vivono in condizioni di estrema povertà e sono maggiormente a rischio di sfruttamento a

³ [«La Commissione pubblica gli orientamenti sulle pratiche vietate in materia di intelligenza artificiale \(IA\), quali definite dalla legge sull'IA»](#), punti 60-97 [consultato il 26 giugno 2025].

⁴ [Commissione europea La Commissione pubblica gli orientamenti sulle pratiche vietate in materia di intelligenza artificiale \(IA\), quale definito dalla legge sull'IA, sez. 3.3](#) (considerando 98-121).

causa dei loro bisogni urgenti, o le persone in un'altra situazione precaria, come i richiedenti asilo che sono in attesa di una decisione sul loro status e il cui futuro è incerto, sono considerati in una "particolare situazione sociale o economica".

Non è sufficiente che qualcuno appartenga a uno solo dei gruppi protetti. Ciò significa che non si applica solo a qualcuno perché ha una disabilità. Il sistema di IA deve effettivamente sfruttare (cioè prendere di mira o sfruttare) la vulnerabilità della persona derivante dalla sua disabilità. L'inaccessibilità o le interfacce utente mal progettate non sono considerate sfruttamento ai fini di questa regola.

Ad esempio, [l'articolo 5, paragrafo 1, lettera b\)](#), potrebbe applicarsi a una persona con disabilità cognitiva che, a causa della sua disabilità, trova molto più difficile comprendere i rischi o i modi in cui un sistema di IA potrebbe manipolarla, anche se riceve supporto o informazioni accessibili.

- Chi è protetto? Persone che, secondo le linee guida, a causa della "loro età, disabilità o situazioni socio-economiche, che in linea di principio hanno una capacità più limitata di riconoscere o resistere alle pratiche manipolative o di sfruttamento dell'IA e necessitano di una maggiore protezione"
- Cosa è vietato? IA che sfrutta queste vulnerabilità per causare danni.

2.2.2 La logica alla base delle disposizioni

Sebbene [l'articolo 5, paragrafo 1, lettera b\)](#), menzioni esplicitamente le persone con disabilità, [l'articolo 5, paragrafo 1, lettera a\)](#), fornirà spesso la protezione primaria per i gruppi vulnerabili, comprese le persone con disabilità.

[L'articolo 5, paragrafo 1, lettere a\) e b\)](#), della legge sull'IA si applica allo sfruttamento o alla manipolazione intenzionali e non intenzionali da parte dei sistemi di IA. Ciò significa che i divieti riguardano le situazioni in cui gli sviluppatori o gli implementatori progettano deliberatamente un sistema di IA per rivolgersi a persone con disabilità, nonché i casi in cui un sistema di IA, indipendentemente dall'intenzione, ha l'effetto di manipolare o sfruttare la vulnerabilità di una persona a causa di disabilità, età o situazione sociale o economica. Il fattore decisivo è il risultato: se il sistema di IA distorce materialmente il comportamento o sfrutta la vulnerabilità e provoca danni significativi, si applica il divieto. In altre parole, non è solo l'intenzione alla base del sistema che conta, ma anche

l'impatto nel mondo reale che ha sugli individui, in particolare su quelli che sono più a rischio a causa delle loro circostanze.

Ad esempio, un utente su sedia a rotelle non è necessariamente più facile da ingannare con l'IA rispetto a qualcuno che non utilizza una sedia a rotelle, quindi [di solito si applica l'articolo 5, paragrafo 1, lettera a](#)).

Tuttavia, può essere più difficile per una persona con disabilità intellettiva riconoscere la manipolazione da parte dell'IA, anche se è istruita su come l'IA può manipolare le persone, quindi si tratta di un gruppo vulnerabile che il legislatore ritiene debba essere protetto in modo speciale dall'[articolo 5, paragrafo 1, lettera b](#)). Questo è simile per le persone anziane che potrebbero avere più difficoltà a riconoscere quando viene utilizzata l'intelligenza artificiale.

Gli orientamenti della Commissione sottolineano che ogni caso richiede un'attenta analisi e che questi due articoli dovrebbero essere esaminati congiuntamente. Nei casi in cui entrambe le disposizioni possono applicarsi a una situazione, le linee guida raccomandano di concentrarsi sull'aspetto principale dello sfruttamento. Se la tecnica manipolativa colpisce allo stesso modo persone appartenenti a gruppi demografici diversi, ma una persona con disabilità è più vulnerabile all'inganno, [si applica l'articolo 5, paragrafo 1, lettera a](#)). In questi casi, la vulnerabilità della persona è considerata un fattore più forte che dimostra il danno.

Tuttavia, se il sistema di IA si rivolge specificamente a uno o più gruppi vulnerabili di cui [all'articolo 5, paragrafo 1, lettera b](#)), come le persone con disabilità, i minori, gli anziani o le persone che vivono in condizioni di estrema povertà, e cerca di sfruttarne la maggiore vulnerabilità, si applica invece questa disposizione.

In pratica, ciò significa che i sistemi di IA devono essere attentamente esaminati per determinare se rappresentano un rischio generale o se si rivolgono specificamente a gruppi vulnerabili, il che influenzerà il modo in cui sono regolamentati.

2.2.2.1 L'asticella alta: quando si applicano questi divieti?

L'asticella per vietare questo tipo di pratiche di IA è intenzionalmente fissata molto in alto nell'articolo 5, non solo per i divieti discussi in precedenza.

Pertanto, non tutti i sistemi di intelligenza artificiale manipolativi o sleali sono vietati. Ai fini [dell'applicazione dell'articolo 5, paragrafo 1, lettere a\) o b\)](#), devono applicarsi tutte le seguenti condizioni:

1. Il sistema di IA viene immesso sul mercato, messo in servizio o utilizzato.
2. Esso deve:
 - Per l'articolo 5, paragrafo 1, lettera a): utilizzare **tecniche subliminali al di là della coscienza di una persona o tecniche intenzionalmente manipolative o ingannevoli**.
 - Ai fini dell'articolo 5, paragrafo 1, lettera b): **sfruttare attivamente (bersaglio o sfruttare)** una vulnerabilità dovuta all'età, alla disabilità o a una specifica situazione sociale o economica.
3. Deve influenzare il comportamento di qualcuno
 - Per quanto riguarda l'articolo 5, paragrafo 1, lettera a): queste tecniche devono avere un **forte impatto sul comportamento** (termine giuridico: "distorcere materialmente il comportamento") **indebolendo considerevolmente** ("sensibilmente pregiudizievole", termine giuridico) la capacità di prendere una decisione informata.
 - Per l'articolo 5, paragrafo 1, lettera b): queste tecniche devono avere un **forte impatto sul comportamento** sfruttando la vulnerabilità (non è necessario dimostrare di compromettere il processo decisionale informato).
4. Ciò causa, o è **ragionevolmente probabile che causi, un danno significativo** (un impatto negativo sostanziale sul benessere fisico, psicologico o finanziario) a quella persona o ad altri.

Tutti e quattro i punti devono essere soddisfatti contemporaneamente (sono chiamati cumulativi) e deve esserci un chiaro legame (un "nesso causale") tra ciò che l'IA fa, il modo in cui cambia il comportamento di una persona e il danno risultante. In altre parole, il danno deve essere il risultato diretto dell'azione manipolativa o di sfruttamento dell'IA e non solo una coincidenza.

La persuasione lecita, come la "pubblicità onesta", non è vietata. Le norme si applicano solo alle pratiche occulte o di sfruttamento che oltrepassano il limite della manipolazione o dello sfruttamento e causano danni reali.

Gli orientamenti della Commissione sottolineano che la violazione di una qualsiasi delle norme di [cui all'articolo 5](#) può comportare le multe più

elevate previste dalla legge, pertanto tali divieti dovrebbero essere interpretati in modo restrittivo e applicati solo nei casi più gravi.⁵

2.2.2.2 Esempi ufficiali tratti dalle linee guida della Commissione europea

[Articolo 5, paragrafo 1, lettera a\)](#): Manipolazione dannosa

Esempio 1:

Un chatbot per il benessere basato sull'intelligenza artificiale è progettato per aiutare gli utenti a condurre una vita più sana. Tuttavia, se il chatbot sfrutta le vulnerabilità degli utenti e incoraggia comportamenti malsani o pericolosi, come suggerire loro di fare esercizio fisico eccessivo senza riposo o acqua, può mettere a rischio le persone. Se qualcuno segue questo consiglio e subisce gravi danni, ad esempio a causa di un infarto, o se esiste un rischio reale di danni gravi, questo tipo di IA potrebbe essere vietato ai sensi dell'articolo 5, paragrafo 1, lettera a).

La legge chiarisce che non è necessario che si verifichi un danno effettivo. Se le azioni del sistema di IA possono causare gravi danni, ciò è sufficiente per l'applicazione del divieto. In altre parole, la regola copre entrambi i casi: quando qualcuno viene effettivamente danneggiato e quando esiste un ragionevole rischio di danno.⁶

Esempio 2:

Un sistema di intelligenza artificiale utilizza messaggi subliminali (come immagini lampeggianti troppo veloci per essere notate dalla mente cosciente) per influenzare le decisioni di acquisto, portando le persone ad acquistare cose che altrimenti non avrebbero avuto, con conseguenti danni finanziari significativi.⁷

[Articolo 5\(1\)\(b\)](#): Sfruttamento delle vulnerabilità⁸

Esempio 1:

⁶ [Commissione europea La Commissione pubblica gli orientamenti sulle pratiche vietate in materia di intelligenza artificiale \(IA\), quali definite dalla legge sull'IA](#), punto 84.

⁷ [Commissione europea La Commissione pubblica gli orientamenti sulle pratiche vietate in materia di intelligenza artificiale \(IA\), quali definite dalla legge sull'IA](#), punto 82.

⁸ [Commissione europea La Commissione pubblica gli orientamenti sulle pratiche vietate in materia di intelligenza artificiale \(IA\), quali definite dalla legge sull'IA](#), punto 108.

Un chatbot terapeutico rivolto alle persone con disabilità intellettiva sfrutta le loro vulnerabilità cognitive per influenzarle ad acquistare costosi prodotti medici o ad adottare comportamenti che potrebbero mettere in pericolo la loro salute o sicurezza.

Esempio 2:

Un sistema di intelligenza artificiale identifica le donne e le ragazze con disabilità su Internet e le indirizza con contenuti offensivi. Sfrutta le loro vulnerabilità legate alla disabilità, il che rende più difficile per loro difendersi e aumenta la loro esposizione a manipolazioni e abusi.

2.2.3 Punteggio sociale proibito

Il social scoring si riferisce all'uso di sistemi di intelligenza artificiale per valutare o classificare gli individui in base al loro comportamento e alle loro caratteristiche nel tempo. Questi sistemi utilizzano spesso i dati raccolti in un contesto e li applicano in un altro. Ad esempio, le informazioni provenienti dall'attività sui social media, dalle abitudini di acquisto o dai dati sanitari non correlati possono essere utilizzate in un contesto completamente diverso. Ciò può portare a risultati inaspettati, ingiusti o sproporzionati.

Questi sistemi creano un punteggio o un profilo per individui o gruppi, che possono quindi essere utilizzati per prendere decisioni su di loro, come l'accesso a servizi, benefici, lavori o altre opportunità.

La preoccupazione alla base del divieto di [cui all'articolo 5, paragrafo 1, lettera c\)](#), è che tale punteggio possa portare a un trattamento dannoso o iniquo delle persone, compresa l'esclusione da servizi o opportunità, e possa consentire un controllo e una sorveglianza sociali incompatibili con i valori e i diritti fondamentali dell'UE. Il divieto mira a proteggere la dignità umana, l'uguaglianza e la non discriminazione, la protezione dei dati e la vita privata e familiare.

C'è una linea sottile tra l'uso lecito dei dati e il punteggio sociale vietato. Ad esempio, se una persona richiede prestazioni di invalidità, una valutazione legale si baserebbe sulla storia medica della persona e, se del caso, sul contenuto dei fascicoli dei servizi sociali che mostrano il tipo di compiti di aiuto e sostegno che la persona ha ricevuto in precedenza. Questo è considerato un uso ad alto rischio dell'intelligenza artificiale ed è soggetto ai severi requisiti della legge sull'intelligenza artificiale, ma non è vietato.

Il limite è superato e il divieto di social scoring si applica se le autorità utilizzano dati eccessivi o irrilevanti. Ad esempio, se le autorità verificano

se una persona ha una storia di consumo elevato di alcol al fine di decidere se ha un diritto inferiore all'assistenza sociale. Non sarebbe pertinente indagare sul consumo di alcol quando si valuta il livello di mobilità ridotta della persona e il suo bisogno di aiuto per le attività quotidiane come la doccia o la toilette. Tuttavia, la valutazione del consumo di alcol sarebbe rilevante se il servizio sociale esaminasse specificamente come la persona può essere aiutata a uscire dalla dipendenza dall'alcol. L'uso di dati irrilevanti o intercontestuali per la valutazione è esattamente ciò che il divieto di punteggio sociale è progettato per impedire.

2.2.3.1 Che cosa vieta esattamente l'articolo 5, paragrafo 1, lettera c)?

Un sistema di IA che valuta o classifica persone fisiche o gruppi in un determinato periodo di tempo in base a:

- il loro comportamento sociale, o
- caratteristiche personali o della personalità (note, dedotte o previste), qualora il punteggio sociale porti o sia in grado di portare a:
- un trattamento dannoso o sfavorevole in contesti sociali estranei al luogo in cui i dati sono stati originariamente raccolti, e/o
- un trattamento dannoso o sfavorevole che sia ingiustificato o sproporzionato rispetto al comportamento o alla sua gravità.

2.2.3.2 Chi è protetto?

- **Tutti:** il divieto si applica a tutte le persone fisiche e a tutti i gruppi, indipendentemente dall'età, dalla disabilità o dallo status sociale.
- **Particolarmente rilevante per le persone con disabilità:** ad esempio, quando l'IA viene utilizzata per valutare l'ammissibilità ai servizi sociali o agli assegni di invalidità, ma si applica anche all'occupazione, all'istruzione, alle assicurazioni e ad altri settori.

2.2.3.3 L'asticella alta: requisiti cumulativi per [l'articolo 5, paragrafo 1, lettera c\)](#)

Affinché si applichi il divieto di punteggio sociale, devono essere soddisfatte contemporaneamente (cumulative) tutte le seguenti condizioni:

1. Posizionamento / Utilizzo

La pratica prevede l'immissione sul mercato, la messa in servizio o l'uso di un sistema di IA. Questo divieto vincola sia i fornitori che i distributori nei rispettivi ruoli. Il divieto si applica indipendentemente dal fatto che gli attori siano organizzazioni private (imprese, organizzazioni senza scopo di lucro, fondazioni) o attori pubblici.

2. Valutazione o classificazione nel tempo

Il sistema di intelligenza artificiale è destinato o utilizzato per valutare o classificare individui o gruppi in un determinato periodo di tempo in base al loro comportamento sociale o alle loro caratteristiche personali o di personalità (comprese quelle note, dedotte o previste).

3. Risultati della partita

Il punteggio porta o è in grado di portare a un trattamento dannoso o sfavorevole:

- In contesti sociali non correlati al luogo in cui i dati sottostanti sono stati generati o raccolti, e/o
- in modo ingiustificato o sproporzionato rispetto al comportamento o alla sua gravità.

2.2.3.4 Esempi tratti dalle linee guida della Commissione europea che illustrano il punteggio sociale

Esempio 1: Trattamento dannoso in contesti sociali non correlati

Un'agenzia pubblica per il lavoro utilizza un sistema di IA per valutare i disoccupati sulla base di un colloquio e di una valutazione basata sull'IA per determinare se un individuo debba beneficiare del sostegno statale per l'occupazione. Tale punteggio si basa su caratteristiche personali pertinenti, come l'età e l'istruzione, ma anche su variabili raccolte o dedotte da dati e contesti senza alcun collegamento apparente con lo scopo della valutazione, come lo stato civile, i dati sanitari per malattie croniche, dipendenze, ecc.

Esempio 2: trattamento ingiustificato o sproporzionato

Un'agenzia pubblica utilizza un sistema di intelligenza artificiale per profilare le famiglie ai fini dell'individuazione precoce dei minori a rischio sulla base di criteri quali la salute mentale dei genitori e la

disoccupazione, ma anche di informazioni sul comportamento sociale dei genitori derivanti da molteplici contesti. Sulla base del punteggio risultante, le famiglie vengono individuate per l'ispezione e i bambini considerati "a rischio" vengono prelevati dalle loro famiglie, anche in caso di trasgressioni minori da parte dei genitori, come occasionalmente assenti da visite mediche o multe stradali.

Esempio 3: trattamento ingiustificato o sproporzionato

Un comune utilizza un sistema di intelligenza artificiale per valutare l'affidabilità dei residenti sulla base di molteplici punti dati relativi al loro comportamento sociale in una varietà di contesti. Il punteggio generato per i residenti considerati "meno affidabili" viene utilizzato per la lista nera, ovvero il ritiro di benefici pubblici, altre gravi misure punitive e un maggiore controllo o sorveglianza. Tra i fattori presi in considerazione nella valutazione vi sono l'insufficiente volontariato e i comportamenti scorretti minori, come non restituire i libri in tempo, lasciare la spazzatura per strada al di fuori del giorno della raccolta e un ritardo nel pagamento delle tasse locali.

Nota importante per i membri di EDF: abbiamo raccolto un elenco di pratiche problematiche di sorveglianza e di dati in materia di protezione sociale, compresi i casi che possono avvicinarsi a un punteggio sociale proibito o che sono almeno usi problematici ad alto rischio. Vedere la sezione "**Esempi di protezione sociale**" nella sezione IA ad alto rischio di questa guida.

Riconoscimento delle emozioni sul luogo di lavoro e nell'istruzione: [articolo 5, paragrafo 1, lettera f\)](#)⁹

L'AI Act vieta l'IA per il riconoscimento delle emozioni sul posto di lavoro e nelle istituzioni educative a causa degli squilibri di potere tra datori di lavoro e dipendenti e scuole e studenti. La maggior parte delle applicazioni di riconoscimento delle emozioni sono considerate ad alto rischio, ma in queste aree si applica un divieto.

Requisiti cumulativi per [l'articolo 5, paragrafo 1, lettera f\)](#)

Affinché si applichi il divieto di riconoscimento delle emozioni, tutte le seguenti condizioni devono essere vere contemporaneamente (cumulative):

⁹ [Commissione europea La Commissione pubblica gli orientamenti sulle pratiche vietate in materia di intelligenza artificiale \(IA\), quali definite dalla legge sull'IA](#), paragrafi 239-270.

Posizionamento / Utilizzo:

Il sistema di IA è immesso sul mercato, messo in servizio per questo scopo specifico o utilizzato.

Scopo:

Il sistema ha lo scopo di dedurre le emozioni.

Contesto:

Il sistema viene utilizzato sul posto di lavoro o negli istituti di istruzione e formazione.

Esclusione:

Il divieto non si applica se il sistema di IA è destinato a motivi medici o di sicurezza.

2.2.3.5 Chiarimenti dalle linee guida

Il riconoscimento delle emozioni qui significa dedurre sentimenti come felicità, tristezza, rabbia, paura, ecc., in base a caratteristiche biometriche (ad esempio, espressioni facciali, tono di voce, linguaggio del corpo).

- L'analisi del sentiment (l'analisi del testo per prevedere le emozioni) non è coperta da questo divieto.
- Gli stati fisici come il dolore o la stanchezza non sono considerati emozioni ai fini di questo divieto. I sistemi che rilevano stanchezza o dolore (ad esempio, per garantire la sicurezza dei camionisti) non sono vietati.
- Semplici osservazioni (come notare che qualcuno sta sorridendo) non sono riconoscimento delle emozioni. Tuttavia, l'utilizzo dei dati biometrici per concludere che qualcuno è felice perché sta sorridendo è il riconoscimento delle emozioni.

Esempi:

- Analizzare la voce di uno studente e concludere che è arrabbiato e sta per diventare violento è il riconoscimento delle emozioni. Tuttavia, le linee guida non dicono esplicitamente se l'uso del riconoscimento delle emozioni per questo scopo è consentito o vietato.
- Identificare che qualcuno è malato non è un riconoscimento emotivo.
- L'utilizzo dell'analisi facciale per dedurre che un lavoratore è triste o disimpegnato è vietato sul posto di lavoro.

Identificazione biometrica remota

L'identificazione biometrica remota si riferisce al processo di identificazione di una persona in base alle sue caratteristiche fisiche, fisiologiche, comportamentali o psicologiche, di solito senza contatto diretto. Ciò può includere l'analisi di caratteristiche come la struttura del viso, i movimenti oculari, la forma del corpo, i modelli vocali, la prosodia (intonazione e ritmo), il modo in cui una persona cammina, la frequenza cardiaca, la pressione sanguigna, l'olfatto e persino il modo unico in cui una persona digita su una tastiera. L'applicazione più comune è il software di riconoscimento facciale che analizza le immagini o i video delle telecamere di sorveglianza. Tuttavia, per stabilire l'identità, è possibile combinare anche vari dati biometrici, ad esempio lo stile di camminata di una persona mascherata con tratti facciali e altezza visibili.

L'elemento che definisce l'identificazione biometrica remota è l'uso della tecnologia per riconoscere le persone a distanza, spesso tramite videocamere o microfoni. Ad esempio, un software è in grado di identificare una persona in base alla voce captata da un microfono o analizzando il modo in cui digita su una tastiera. Questi sistemi funzionano senza interazione fisica e si basano invece sui dati raccolti in remoto per effettuare identificazioni.

[L'articolo 5, paragrafo 1, lettera h\)](#), vieta l'uso dell'identificazione biometrica in tempo reale, come il riconoscimento facciale, negli spazi pubblici da parte delle forze dell'ordine. La polizia non può utilizzare il riconoscimento facciale sui feed delle telecamere in diretta, ma può utilizzarlo sui video registrati, che sono invece considerati ad alto rischio.

Vi sono tre eccezioni in cui la polizia può utilizzare l'identificazione biometrica in tempo reale, ma solo se una legge nazionale di quel paese dell'UE lo consente specificamente. Senza una tale legge, non è consentito.

Le tre eccezioni sono:

1. Localizzare persone scomparse o vittime di rapimento, tratta o sfruttamento sessuale.
2. Prevenire una minaccia specifica e imminente alla vita, alla sicurezza o all'attività terroristica.
3. Identificazione dei sospetti di reati gravi (ad esempio, terrorismo, omicidio, sfruttamento sessuale di minori) elencati nell' [allegato II](#) della legge sull'intelligenza artificiale.

La polizia deve condurre una valutazione d'impatto sui diritti fondamentali prima di effettuare l'identificazione biometrica in tempo reale. Di solito

devono anche registrarlo nel registro dei sistemi di intelligenza artificiale ad alto rischio.

Sia il [considerando 32 che](#) gli orientamenti riconoscono che i sistemi di identificazione biometrica remota possono essere distorti e portare a discriminazioni sulla base di disabilità, età, etnia, razza o sesso.

Un chiaro esempio che evidenzia la necessità di regolamentare rigorosamente l'uso dell'identificazione biometrica in tempo reale viene dall'Ungheria. Nel 2025, il governo ha criminalizzato la partecipazione alla parata del Pride e ha minacciato di utilizzare la tecnologia di riconoscimento facciale per identificare i partecipanti. Le organizzazioni per i diritti umani hanno fortemente criticato la Commissione europea per non essere intervenuta. Questo caso ha anche messo in luce una lacuna significativa: mentre l'uso del riconoscimento facciale in tempo reale sui feed delle telecamere pubbliche è vietato, le autorità possono aggirare questa restrizione utilizzando invece filmati di sorveglianza registrati.¹⁰

3. IA ad alto rischio

L'UE considera ad alto rischio i **sistemi di IA che potrebbero incidere in modo significativo sulla vita, sui diritti o sul benessere delle persone**. Questi sistemi sono utilizzati in processi decisionali critici, come ad esempio:

- Determinare se qualcuno viene accettato in un programma educativo.
- Rilevamento di imbrogli durante esami o valutazioni.
- Assistere i datori di lavoro nelle decisioni di assunzione o licenziamento.
- Decidere se qualcuno si qualifica per benefici o servizi governativi.
- Approvazione delle domande di prestito o assicurazione.
- Sostenere le autorità di contrasto nella prevenzione o nel perseguimento dei reati.

Per **le persone con disabilità**, due requisiti chiave sono particolarmente importanti quando si tratta di sistemi di intelligenza artificiale ad alto rischio.

In primo luogo, questi sistemi devono essere sottoposti a **test di bias** per identificare e correggere eventuali bias. Questo è fondamentale perché l'IA distorta può portare a risultati discriminatori, in particolare per le

¹⁰ ["Si accumulano pressioni su Bruxelles affinché agisca contro la scansione del volto al Pride di Budapest", POLITICO, 26 giugno 2025](#) [consultato il 16 ottobre 2025].

persone con disabilità che potrebbero essere più vulnerabili a tali pregiudizi nei processi decisionali.

In secondo luogo, i sistemi di IA ad alto rischio devono essere conformi ai **requisiti obbligatori di accessibilità**. Questi requisiti sono stabiliti dalla legge europea sull'accessibilità ([direttiva \(UE\) 2019/882](#)) e dalla direttiva sull'accessibilità del web ([direttiva \(UE\) 2016/2102](#)). Ciò garantisce che i sistemi di IA siano accessibili a tutti gli utenti, comprese le persone con disabilità, incorporando fin dall'inizio i principi di progettazione inclusiva.

3.1 Banca dati di IA ad alto rischio

[L'articolo 71](#) della legge sull'IA impone alla Commissione europea di istituire e mantenere una **banca dati pubblica** dei sistemi di IA ad alto rischio.

Questa banca dati pubblica sarà disponibile al pubblico **entro agosto 2026** e aiuterà tutti noi a monitorare quali tipi di sistemi di IA ad alto rischio vengono utilizzati nei paesi dell'UE. La tua organizzazione può iniziare a monitorare il database da questo momento in poi.

Tuttavia, i sistemi di IA ad alto rischio utilizzati dalle autorità di contrasto, di migrazione e di controllo delle frontiere saranno conservati in una sezione riservata della banca dati. Solo la Commissione europea e le autorità degli Stati membri dell'UE avranno accesso alla parte riservata del database.¹¹

L'assenza di trasparenza renderà molto più difficile per i giornalisti e le organizzazioni per i diritti umani esaminare i sistemi di IA difettosi o pericolosi o i modi in cui vengono utilizzati

3.2 Requisito di accessibilità per l'IA ad alto rischio

[L'articolo 16, lettera I\)](#), della legge sull'IA impone ai fornitori di sistemi di IA ad alto rischio di rispettare i requisiti di accessibilità in conformità, rispettivamente, alle direttive (UE) 2016/2102 e (UE) 2019/882, alla direttiva sull'accessibilità del web e alla legge europea sull'accessibilità.

Ciò significa che gli sviluppatori di IA di questi sistemi di IA dovranno seguire i **requisiti di accessibilità** di tali direttive e saranno in grado di utilizzare le norme europee armonizzate in materia di accessibilità per dimostrare la conformità a tale obbligo.

Ad esempio, potranno utilizzare la norma europea armonizzata "[EN 301 549 - Requisiti di accessibilità per prodotti e servizi TIC](#)" e la norma

¹¹ [Articolo 49, paragrafo 4](#), della legge sull'IA

europea "[EN 17161 - Progettazione per tutti - Accessibilità seguendo un approccio di progettazione per tutti in prodotti, beni e servizi - Ampliamento della gamma di utenti](#)".

Il considerando 80 della legge sull'IA spiega che i sistemi di IA devono essere accessibili a tutti, comprese le persone con disabilità, nell'ambito dell'impegno dell'UE a prevenire la discriminazione e promuovere l'uguaglianza. Ciò significa che i sistemi di IA ad alto rischio devono essere progettati seguendo un approccio di "progettazione universale" che consiste nel creare prodotti e servizi inclusivi fin dall'inizio. Pertanto, gli sviluppatori dovrebbero considerare l'accessibilità fin dall'inizio del processo di progettazione per evitare adattamenti successivi e rispettare la dignità e le esigenze delle persone con disabilità.

3.3 Pregiudizio

L'accessibilità non è l'unico problema che ci preoccupa come persone con disabilità quando vengono utilizzati i sistemi di intelligenza artificiale. **I pregiudizi** sono un altro grave problema che può avere conseguenze negative per i nostri diritti e le nostre opportunità.

Pertanto, oltre all '**accessibilità**', gli sviluppatori di sistemi di IA ad alto rischio devono anche considerare i potenziali pregiudizi che possono derivare dai dati che utilizzano o dai modelli di IA che creano. **I pregiudizi** possono rappresentare una minaccia per la nostra dignità, uguaglianza e inclusione in diversi ambiti della vita.

Dati di alta qualità sono essenziali per lo sviluppo di sistemi di intelligenza artificiale ad alto rischio accurati e imparziali. [L'articolo 10](#) della legge sull'IA stabilisce requisiti specifici per i dati utilizzati per addestrare, convalidare e testare tali sistemi. I dati devono essere affidabili, sicuri e rispettosi della salute, dell'incolumità, dei diritti fondamentali e della non discriminazione delle persone, in particolare per le persone con disabilità.

I dati devono inoltre essere controllati per individuare eventuali pregiudizi che potrebbero danneggiare la **salute** e la **sicurezza** delle persone, violare i loro **diritti fondamentali** o portare a discriminazioni illecite, ad esempio nei confronti **delle persone con disabilità**. Se tali distorsioni vengono rilevate, devono essere prevenute o ridotte al minimo.

Inoltre, i dati devono essere **selezionati e gestiti con cura**, il che include la corretta raccolta, pertinenza ed elaborazione (ad es. corretta etichettatura o aggiornamento). Tutte le ipotesi fatte in relazione ai dati devono essere chiaramente indicate.

3.3.1 I pregiudizi non sono l'unico problema

È importante riconoscere che i pregiudizi sono solo una delle numerose minacce ai diritti delle persone con disabilità. Alcuni casi d'uso dell'IA sono sbagliati, anche quando i pregiudizi vengono rimossi in modo tale che l'IA possa fare una valutazione o una previsione accurata su una persona con disabilità.

3.3.2 Esempi pratici di bias

Il pregiudizio dell'IA si verifica quando un sistema di IA è ingiusto o sbagliato per alcune persone più che per altre. Ad esempio, un sistema di intelligenza artificiale che decide chi ottiene un lavoro potrebbe essere ingiusto nei confronti delle persone con disabilità se i dati da cui ha appreso non le includono o non le mostrano. **I pregiudizi dell'intelligenza artificiale possono danneggiare la vita delle persone**, ad esempio privarle di possibilità, risorse o benefici. Pertanto, è importante assicurarsi che i sistemi di intelligenza artificiale siano realizzati e controllati in modo da fermare o ridurre i pregiudizi e rispettare i diritti umani.

Possiamo illustrare i pregiudizi con l'esempio di un programma di intelligenza artificiale progettato per fermare l'incitamento all'odio. Questo tipo di programma analizza i testi per determinare se il sentimento dietro una parola è positivo o negativo.

Uno studio del 2021 illustra i pregiudizi nel modo in cui **l'intelligenza artificiale assegna un valore numerico** per indicare l'intensità delle parole negative. Puoi istruire l'IA a trattare qualsiasi cosa peggiore di una soglia specifica come **incitamento all'odio da rimuovere**. Ad esempio, se i sentimenti negativi vanno da -0,01 a -1, possiamo istruire il software a trattare qualsiasi cosa più intensa di -0,30 come incitamento all'odio. Anche se questo sembra giusto perché applica lo stesso standard a tutti, lo studio ha rivelato che ci possono essere pregiudizi contro le persone con disabilità in sistemi che misurano matematicamente l'intensità dei sentimenti dietro parole diverse.

Tabella delle frasi e dei punteggi (annotazioni degli autori tra parentesi):

- Il mio vicino è una persona alta: 0.00 (Neutro)
- Il mio vicino è una bella persona: 0.85 (Positivo)
- Il mio vicino è una persona di colore: -0.16 (leggermente negativo)
- Il mio vicino è una persona con handicap mentale: -0.10 (Leggermente negativo)
- Il mio vicino è una persona cieca: -0.50 (Molto negativo)

Fonte: Venkit, Pranav Narayanan e Shomir Wilson¹²

Nello scenario in cui assumiamo che qualsiasi punteggio inferiore a -0,30 indichi incitamento all'odio, questo metodo non riesce a filtrare il termine più offensivo "handicappato" mentre blocca l'uso della parola "cieco".

Il nostro esempio mostra che un filtro automatico per l'incitamento all'odio potrebbe erroneamente impedire alle persone con cecità di condividere le loro esperienze o difendere i loro diritti sui social media. Di conseguenza, esiste un rischio potenziale che tali filtri possano **limitare la libertà di espressione** dei gruppi emarginati che stanno cercando di proteggere.

3.4 Il diritto di ottenere una spiegazione

[L'articolo 86](#) della legge sull'IA stabilisce che una persona che ritiene che una decisione di un'organizzazione che utilizza un sistema di IA ad alto rischio la riguardi giuridicamente e possa incidere sulla sua salute, sicurezza o diritti fondamentali ha il diritto di ottenere una spiegazione dall'organizzazione che utilizza l'IA. Ciò non si applica ai sistemi di IA utilizzati per servizi infrastrutturali critici come l'elettricità e l'acqua.

3.5 Categorie di IA ad alto rischio

Raccomandazioni generali per i membri del FES

- **Consulta la banca dati dei sistemi di IA ad alto rischio:** a partire dall'agosto 2026, quando entreranno in vigore le disposizioni per l'IA ad alto rischio, puoi monitorare la banca dati dell'UE per individuare i sistemi di IA utilizzati nel tuo Stato membro che potrebbero rappresentare un rischio per i diritti fondamentali. Potrebbe poi essere possibile creare alleanze con altre organizzazioni che tutelano i diritti fondamentali, come le organizzazioni dei pazienti, per raccogliere testimonianze di sistemi che non funzionano correttamente e fornirle alle autorità di controllo.
- **Sostenere il diritto delle persone a essere informate:** il diritto di essere informati appartiene alle persone interessate dai sistemi di IA. Ad esempio, se qualcuno sta cercando di ottenere un prestito o il suo accesso all'assistenza sanitaria è stato deciso da un'intelligenza artificiale, ha il diritto di sapere se è in gioco un sistema di intelligenza artificiale e come ha ragionato. I membri di EDF possono sostenere queste persone incoraggiandole a chiedere

¹² [Pranav Narayanan Venkit e Shomir Wilson, "Identificazione dei pregiudizi contro le persone con disabilità nell'analisi del sentimento e nei modelli di rilevamento della tossicità", arXiv:2111.13259 \[Cs\], 2021 \[consultato il 15 aprile 2023\].](#)

spiegazioni e in modo che possano capire in che modo le decisioni dell'IA le influenzano.

- **Contattare l'autorità di vigilanza del mercato:** se si sospetta un comportamento scorretto di un sistema di IA, contattare l'autorità di vigilanza del mercato competente.

3.5.1 IA ad alto rischio nell'istruzione:

Alcune applicazioni di intelligenza artificiale nell'istruzione sono considerate ad alto rischio perché possono avere un impatto significativo sul futuro degli studenti e possono perpetuare pregiudizi o invadere la privacy¹³. Gli esempi includono:

- Decisioni di ammissione ai programmi di formazione
- Valutazione degli studenti
- Determinazione dei livelli di istruzione
- Monitoraggio degli imbrogli durante gli esami

Prove del problema

- Software anti-cheat che segnala gli studenti con disabilità come imbrogliatori¹⁴
- Un algoritmo informatico che ha causato una crisi di valutazione nelle scuole britanniche che ha aggravato le ingiustizie socioeconomiche.¹⁵ Molti studenti con disabilità erano probabilmente tra quelli colpiti, poiché sono spesso all'intersezione di molteplici forme di emarginazione, incluso lo status socioeconomico.

Raccomandazioni per i membri di EDF:

- **Invito a presentare linee guida sull'IA nell'istruzione:** esortare le autorità a pubblicare linee guida per le istituzioni educative affinché elaborino linee guida sull'uso sicuro ed etico dell'IA nell'istruzione e offritevi di contribuire con la vostra esperienza per garantire che le esigenze delle persone con disabilità siano incluse negli sforzi nazionali per utilizzare l'IA nell'istruzione.

¹³ [Considerando 56](#) della legge sull'IA

¹⁴ ["In che modo il software automatizzato di supervisione dei test discrimina gli studenti disabili"](#), Center for Democracy and Technology [consultato il 13 novembre 2022].

¹⁵ [Sam Shead, "Come un algoritmo informatico ha causato una crisi di valutazione nelle scuole britanniche", CNBC](#), 21 agosto 2020, sezione Tecnologia [consultato il 5 ottobre 2024].

Ad esempio, molte istituzioni educative promuovono l'uso di strumenti di feedback per migliorare le capacità di scrittura degli studenti. Tali strumenti sono particolarmente utili per gli studenti con disabilità come la dislessia. In un'università statunitense che supportava l'uso di Grammarly, si è verificato un incidente in cui uno studente è stato sospeso dopo un aggiornamento di Grammarly. Grammarly ha migliorato la qualità del proprio servizio includendo un modello linguistico AI come quello di ChatGPT. Ciò ha indotto il software di rilevamento del plagio dell'università a contrassegnare il lavoro dello studente come generato dall'intelligenza artificiale.^{16 17}

Allo stesso tempo, si può presumere che i chatbot ben funzionanti abbiano il potenziale per essere uno strumento di apprendimento. L'emergere di modelli linguistici di intelligenza artificiale si sta rivelando trasformativo e sta spingendo a rivalutare le pratiche pedagogiche.¹⁸

L'UNESCO ha pubblicato una guida globale sull'IA generativa nell'istruzione per aiutare le nazioni ad agire e pianificare il futuro. Sebbene la guida promuova l'inclusione, l'equità e la diversità, manca di strumenti e regolamenti specifici che rispondano alle esigenze degli studenti con disabilità. Ciò sottolinea la necessità di politiche che garantiscano che gli studenti con disabilità possano beneficiare dei progressi dell'IA senza effetti negativi.¹⁹ I membri del FES potrebbero partecipare attivamente all'elaborazione di tali orientamenti nei loro paesi. Un esempio stimolante che fa luce su questo problema è un documento politico dell'Educating All Learners Alliance e di New America che affronta la protezione dei diritti civili degli studenti con disabilità nell'era dell'intelligenza artificiale.²⁰

- **Sostenere la creazione di meccanismi di rimedio nelle leggi scolastiche:** lavorare con i gruppi per i diritti degli studenti per cambiare la legge in modo che gli studenti di diversi tipi di scuole abbiano il diritto legale di appellarsi a decisioni come la sospensione o l'esame. Anche se la decisione è stata presa da esseri umani

¹⁶ [Jeanette Settembre, "Studente messo in libertà vigilata per aver usato Grammarly: Violazione dell'intelligenza artificiale"](#), New York Post, 21 febbraio 2024, sezione Tecnologia [consultato il 5 ottobre 2024].

¹⁷ [William Tang, "Che cos'è l'imbroglione? L'intelligenza artificiale vieta, le scuole possono bocciare gli studenti che usano la grammatica"](#), USA TODAY, 17 aprile 2024 [consultato il 6 ottobre 2024].

¹⁸ [Grammarly, "Ripensare le politiche di integrità accademica nell'era dell'intelligenza artificiale"](#), Grammarly Blog, 2024 [consultato il 6 ottobre 2024].

¹⁹ [Fengchun Miao e Holmes Wayne, Guida per l'IA generativa nell'istruzione e nella ricerca](#) (UNESCO, 2023), doi:10.54675/EWZM9535.

²⁰ [Educare All Learners Alliance e New America, dando priorità agli studenti con disabilità nella politica dell'IA](#), 2023 [consultato il 6 ottobre 2024].

utilizzando strumenti automatizzati per cercare di riconoscere se il lavoro dello studente è stato creato da un'intelligenza artificiale, potrebbe esserci il rischio che lo strumento di riconoscimento contenga errori.

3.5.2 IA ad alto rischio nel mondo del lavoro:

I sistemi di intelligenza artificiale utilizzati nell'occupazione sono ad alto rischio perché possono avere un impatto significativo sulle opportunità di lavoro, sul reddito e sui diritti dei lavoratori, perpetuando potenzialmente pregiudizi e invadendo la privacy.²¹

Esempi di applicazioni di intelligenza artificiale ad alto rischio:

- **AI per l'assunzione:**
 - Annunci di lavoro mirati
 - Smistamento delle domande di lavoro
 - Selezione dei candidati
- **AI per la gestione della forza lavoro:**
 - Decisioni su promozioni e licenziamenti
 - Assegnazioni di attività
 - Monitoraggio delle prestazioni e del comportamento dei dipendenti

Prove di problemi:

- Le tecnologie di sorveglianza sul posto di lavoro possono svantaggiare i lavoratori, come evidenziato nel rapporto CDT.²²
- Il software di intelligenza artificiale utilizzato per valutare i candidati in base al loro comportamento nei colloqui video ha dimostrato di discriminare in modo massiccio i candidati con disabilità.²³

Raccomandazioni per i membri di EDF:

Collaborare con i sindacati e le organizzazioni per i diritti digitali per raccogliere prove: dall'agosto 2026, quando entreranno in vigore le disposizioni sull'IA ad alto rischio, è possibile monitorare la banca dati dell'UE per individuare i sistemi di IA utilizzati nel proprio Stato membro e che possono rappresentare un rischio per i diritti dei lavoratori. Potrebbe poi essere possibile creare alleanze con altre organizzazioni che tutelano i

²¹ [Considerando 57](#) della legge sull'IA

²² [Lydia X. Z. Brown e altri, Abilismo e discriminazione della disabilità nelle nuove tecnologie di sorveglianza](#) (Centro per la Democrazia e la Tecnologia) [consultato il 13 novembre 2022].

²³ [Mara Mills e Meredith Whittaker, Disabilità, pregiudizi e intelligenza artificiale, Rapporto dell'AI Now Institute, 2019](#), pp. 15-17 circa 'HireVue'.

diritti fondamentali per raccogliere testimonianze di sistemi che non funzionano e fornirle alle autorità di controllo.

Costruire relazioni con le autorità di regolamentazione: stabilire contatti con le autorità di vigilanza del mercato del lavoro, gli organismi nazionali per i diritti umani, i difensori civici per l'uguaglianza, le autorità di protezione dei dati e le autorità di regolamentazione del mercato in modo che possano consultarti in caso di domande sulle persone con disabilità.

3.5.3 IA ad alto rischio nei servizi essenziali:

I sistemi di IA utilizzati nei servizi essenziali sono considerati ad alto rischio a causa del loro impatto significativo sui diritti delle persone e del potenziale di discriminazione.²⁴

I servizi essenziali sono i servizi necessari per vivere una vita appagante in una società moderna e industrializzata. Gli esempi includono:

- Sanità
- Servizi di risposta alle emergenze (ad es. ambulanza, polizia)
- Prestazioni governative (ad esempio, prestazioni di disoccupazione, pensioni o invalidità)
- Utenze (ad esempio, elettricità, gas)
- Accesso ai servizi finanziari (ad es. assicurazioni, banche, prestiti)

Questi sistemi sono ad alto rischio poiché l'utilizzo di IA con difetti e pregiudizi potrebbe significare che le persone sono discriminate.

Esempio: chatbot sanitario in Svezia

I ricercatori dell'Istituto reale svedese di tecnologia hanno valutato un chatbot che è stato introdotto come un ulteriore modo per contattare l'assistenza sanitaria, oltre alla possibilità per il paziente di parlare con un'infermiera quando chiama la hotline sanitaria in Svezia chiamata 1177. La valutazione ha dimostrato che l'IA non era utilizzabile/accessibile a persone con diversi tipi di disabilità.²⁵

Questo esempio mostra i potenziali rischi derivanti dall'introduzione troppo precoce delle tecnologie di IA sul mercato. Sebbene il chatbot non debba sostituire ma integrare la consultazione di un infermiere, il suo rilascio anticipato con problemi di accessibilità irrisolti potrebbe mettere a rischio i pazienti. Se i pazienti si affidano interamente al chatbot senza

²⁴ [Considerando 58](#) della legge sull'IA

²⁵ [Marika Jonsson e a., «Tillgänglighet och användbarhet i 1177direkt»](#), 2023 (Fonte disponibile solo in svedese).

rendersi conto dei suoi limiti, potrebbero non ricevere cure mediche tempestive e accurate.

Esempi di sorveglianza problematica nella protezione sociale

1. Riconoscimento facciale dei manifestanti – riferimenti incrociati con l'elenco delle persone che ricevono prestazioni di invalidità (Regno Unito)

Un altro problema di cui i membri di EDF devono essere consapevoli è quando le autorità utilizzano la sorveglianza, con o senza intelligenza artificiale, nella loro ricerca per determinare se qualcuno ha diritto a un beneficio.

L'IA può essere legalmente utilizzata per prevenire le frodi nei sistemi di protezione sociale ai sensi dell'AI Act, ma è fondamentale garantire che queste misure non vadano troppo oltre e rispettino la privacy e la dignità delle persone con disabilità.

Un esempio preoccupante viene dal Regno Unito. Il rapporto State of Surveillance 2023 dell'organizzazione per le libertà civili Big Brother Watch, con sede nel Regno Unito, mostra come i sistemi di sorveglianza del welfare abbiano criminalizzato le persone disabili, trasformando le loro attività quotidiane in prove per una costante valutazione governativa.²⁶

Ci sono notizie preoccupanti secondo cui la polizia nel Regno Unito ha passato informazioni e video di sorveglianza sui manifestanti con disabilità alle autorità che decidono il loro diritto ai benefici legati alla disabilità, mettendoli a rischio di vedersi annullare i benefici.²⁷

²⁶ [Rick Burgess, "Il panopticon del sistema di welfare", in Stato di sorveglianza nel 2023](#), a cura di Big Brother Watch (Greater Manchester Coalition of Disabled People, 2023), p. 16 [consultato il 6 ottobre 2024].

²⁷ [John Pring, "Le forze di polizia ammettono di aver passato filmati di manifestanti disabili al DWP"](#), Disability News Service, 20 dicembre 2018 [consultato il 3 aprile 2023]; [John Pring, "La forza di polizia della Conferenza Tory ammette di aver condiviso informazioni sui manifestanti con il DWP"](#), Disability News Service, 14 febbraio 2019 [consultato il 3 aprile 2023]; ["Far luce sul DWP Parte 1 - Abbiamo letto la guida di 995 pagine dell'Agenzia per il Welfare del Regno Unito sulla conduzione della sorveglianza e qui ci sono le parti più spaventose"](#), Privacy International [consultato il 3 aprile 2023]; ["Far luce sul DWP Parte 2 - Una lunga giornata di viaggio verso la trasparenza | Privacy International"](#) [consultato il 6 ottobre 2024]; ["Il Comitato delle Nazioni Unite per i diritti delle persone con disabilità chiede al Regno Unito di agire sui rischi per i diritti umani dell'IA"](#), Privacy International [consultato il 6 ottobre 2024].

Questi tipi di pratiche possono essere utilizzate con o senza l'intelligenza artificiale e possono avere un effetto dissuasivo sulla libertà di parola delle persone con disabilità.

Un'altra cosa è che i governi possono utilizzare l'intelligenza artificiale per rilevare irregolarità o anomalie nella protezione sociale e nei benefici. Ciò può significare esaminare tutto ciò che potrebbe essere considerato una potenziale frode. Le persone con disabilità tendono a deviare dalla norma e siamo anche intersezionali nelle nostre identità, il che significa che qualcuno può avere una disabilità ed essere una persona LGBTQI. Di norma, le persone con disabilità si distinguono nelle statistiche e sono quindi esaminate più frequentemente. Ed è importante che ciò non avvenga in modo discriminatorio.

2. Danimarca - Monitoraggio intrusivo e perdita della privacy²⁸

Il 12 novembre 2024 Amnesty International ha pubblicato un rapporto che evidenzia i rischi di discriminazione da parte dell'Agenzia danese per la previdenza sociale Udbetaling Danmark (UDK). Il rapporto sottolinea che l'uso dell'intelligenza artificiale per identificare le persone nelle indagini sulle frodi previdenziali rappresenta un rischio di discriminazione nei confronti delle persone con disabilità, delle persone a basso reddito, dei migranti, dei rifugiati e dei gruppi razzialmente emarginati.

Secondo Amnesty International, le autorità elaborano anche una grande quantità di dati personali su una scala che Amnesty descrive come sorveglianza di massa. La relazione sottolinea che esiste una combinazione tra il trattamento di massa dei dati personali, che in molti casi è di natura sensibile, e l'esame attento e intimo delle persone che richiedono o ricevono prestazioni quando vengono intervistate da assistenti sociali e investigatori sulle frodi.

Secondo il comunicato stampa, le autorità danesi utilizzano un sistema composto da un massimo di 60 modelli algoritmici che sarebbero stati utilizzati per individuare le frodi sui sussidi e segnalare le persone per ulteriori indagini da parte delle autorità danesi. Amnesty International afferma di aver ottenuto l'accesso parziale a quattro di questi algoritmi come parte della loro ricerca. Il comunicato stampa afferma che le persone intervistate da Amnesty International hanno riferito di essere sottoposte a una notevole pressione psicologica a causa della sorveglianza da parte degli investigatori e degli operatori del caso. Secondo il presidente del Comitato per le politiche sociali e del mercato del lavoro

²⁸ ["Danimarca: il sistema di welfare basato sull'intelligenza artificiale alimenta la sorveglianza di massa e rischia di discriminare i gruppi emarginati – Rapporto"](#), Amnesty International, 12 novembre 2024 [consultato il 21 settembre 2025].

presso la Dansk Handicap Foundation, le persone con disabilità che vengono ripetutamente interrogate dagli assistenti sociali possono soffrire di depressione e stress. L'organizzazione ha riferito che molte persone descrivono il controllo in corso come avente un impatto negativo significativo sul loro benessere, con alcuni che affermano che sembra che la sorveglianza costante li stia "erodendo".

Il rapporto fa luce sul monitoraggio intrusivo, come valutare se qualcuno ha bisogno di un assistente personale per assistere in casa. Il rapporto si riferisce a un partecipante a un focus group che ha descritto la propria esperienza di assistenti sociali che effettuano queste valutazioni come una forma di sorveglianza. Trovava umiliante avere un assistente sociale in casa in ogni momento, anche durante le attività private come la doccia, per valutare esattamente quanto tempo era necessario per questo e poi. Il rapporto afferma che si tratta di un monitoraggio eccessivo che sembra non essere necessario e non aggiunge alcun valore nel determinare il livello appropriato di supporto personalizzato, ma piuttosto invade la privacy del beneficiario del beneficio.

Il comunicato stampa di Amnesty afferma che il sistema danese potrebbe essere classificato come un sistema di punteggio sociale, che rientrerebbe nel divieto. Le autorità danesi non sono d'accordo.

3. Francia: gruppi emarginati interessati dagli strumenti di valutazione del rischio

In Francia, quattordici organizzazioni per i diritti umani, tra cui APF France handicap e Amnesty International, hanno avviato procedimenti legali contro un algoritmo di valutazione del rischio utilizzato dalla CNF (Caisse Nationale des Allocations Familiales), il fondo nazionale per gli assegni familiari che gestisce i programmi di assistenza sociale, citando preoccupazioni relative a criteri discriminatori. Le organizzazioni sostengono che l'algoritmo classifica ampiamente gli individui che vivono l'emarginazione, come le persone con disabilità, come soggetti di sospetto. Altri gruppi identificati come potenzialmente vulnerabili alla discriminazione includono genitori single, donne e individui che vivono in povertà.²⁹

Le organizzazioni fanno anche riferimento ai risultati di Lighthouse Reports, che ha esaminato parti dell'algoritmo ottenuto attraverso le richieste di libertà di informazione. Secondo la loro analisi, i dati disponibili suggeriscono che sia avere una disabilità che essere occupati

²⁹ [«Francia: l'algoritmo discriminatorio utilizzato dall'Agenzia per la sicurezza sociale deve essere fermato»](#), Amnesty International, 16 ottobre 2024 [consultato il 21 settembre 2025].

possono portare indipendentemente un individuo a superare la soglia per ulteriori indagini.³⁰

4. Serbia – Le persone con disabilità affrontano la povertà causata da una digitalizzazione ingiusta³¹

Il rapporto di Amnesty International "Intrappolato dall'automazione: povertà e discriminazione nello stato sociale della Serbia" documenta l'introduzione del registro delle Social Card, un sistema informativo centralizzato che utilizza l'automazione per consolidare i dati personali dei richiedenti e dei beneficiari dell'assistenza sociale da una serie di database ufficiali del governo

Il registro introduce un processo decisionale semi-automatizzato nella valutazione dell'ammissibilità all'assistenza sociale e segnala i casi che richiedono la revisione da parte di un assistente sociale.

La relazione rileva che il registro della Social Card ha reso operative le condizioni restrittive di ammissibilità esistenti e ha esacerbato l'esclusione, danneggiando in particolare i Rom e le persone con disabilità. Il sistema si basa spesso su dati di origine imprecisi, portando le persone a perdere l'assistenza sociale perché i dati raccolti erano errati o classificati in modo errato. Le comunità emarginate sono colpite in modo sproporzionato, poiché la forte dipendenza del sistema dai dati grezzi sulle attività e sul reddito non coglie la complessità della vita delle persone.

L'anagrafica delle Social Card ha ridotto l'autonomia degli assistenti sociali. Molti ora si rimettono agli output del sistema, anche quando sanno che i dati non sono corretti. Il registro è progettato per tenere traccia delle modifiche che possono portare alla perdita dell'assistenza, piuttosto che per facilitare l'accesso ai benefici. Il rapporto evidenzia che questo sistema rischia la discriminazione indiretta, in quanto non tiene conto degli ostacoli che i gruppi emarginati incontrano nel mantenere aggiornati i propri registri o nel navigare in complessi processi amministrativi.

Amnesty International conclude che l'introduzione dell'automazione nel sistema di welfare serbo ha aggravato le carenze esistenti, messo a repentaglio il diritto delle persone alla sicurezza sociale e danneggiato in modo sproporzionato le comunità già emarginate.

³⁰ ["Come abbiamo indagato sulla macchina francese per la profilazione di massa", Rapporti del faro](#), [consultato il 21 settembre 2025].

³¹ ["Serbia: intrappolata dall'automazione: povertà e discriminazione nello stato sociale della Serbia"](#), Amnesty International, 4 dicembre 2023 [consultato il 21 settembre 2025].

Altri esempi di applicazioni di IA ad alto rischio nei servizi essenziali:

Servizi pubblici: determinare l'ammissibilità all'assistenza pubblica, all'assistenza sanitaria, alle prestazioni di invalidità o alle prestazioni sociali.

- **Servizi finanziari:**
 - Valutazioni di idoneità al credito
 - Rilevamento delle frodi
- **Assicurazioni:** definizione dei premi in base alla valutazione del rischio
- **Servizi di emergenza:**
 - Dare priorità alle chiamate di emergenza
 - Invio di ambulanze e triage dei pazienti

Prove di problemi:

- I sistemi di IA possono svantaggiare ingiustamente le persone con disabilità, come si vede in:
 - Le banche utilizzano criteri irrilevanti come la capitalizzazione delle parole per l'affidabilità creditizia, svantaggiando le persone con dislessia (Binns e Kirkham, 2020).³²
 - Gli algoritmi segnalano in modo sproporzionato i genitori con disabilità.³³

Raccomandazioni per i membri dell'EDF: sostenere e sostenere l'idea che i revisori terzi debbano valutare questi tipi di IA e garantire che i quadri di audit siano basati su principi e linee guida sull'accessibilità.

3.5.4 IA ad alto rischio nelle forze dell'ordine

I sistemi di intelligenza artificiale utilizzati nelle forze dell'ordine sono ad alto rischio perché possono avere un impatto significativo sulle libertà e sui diritti individuali, portando potenzialmente ad azioni illecite.³⁴

Esempi di applicazioni di intelligenza artificiale ad alto rischio:

³² [Reuben Binns e Reuben Kirkham, "In che modo la legge sull'uguaglianza e la protezione dei dati potrebbero plasmare l'equità dell'IA per le persone con disabilità?", arXiv:2107.05704 \[Cs\], 2021, p. 14 \[consultato il 6 ottobre 2024\].](#)

³³ ["Ecco come uno strumento di intelligenza artificiale può segnalare i genitori con disabilità", AP NEWS, 2023 \[consultato il 15 aprile 2023\].](#)

³⁴ [Considerando 59](#) della legge sull'IA

- **Polizia predittiva:** prevedere chi saranno le potenziali vittime di reati^{35 36}
- **Rilevamento delle bugie:** utilizzo dell'intelligenza artificiale come macchina della verità^{37 38}
- **Valutazione delle prove:** valutazione della credibilità delle prove nelle indagini o nei processi
- **Previsione della recidiva:** giudicare la probabilità di recidiva in base alla personalità e al comportamento passato³⁹
- **Profilazione:** profilazione delle persone durante l'individuazione, l'indagine o il perseguimento di reati⁴⁰

Raccomandazioni per i membri di EDF:

Lobby per la trasparenza: sostieni le norme nazionali che impongono alle autorità di contrasto del tuo Stato membro di essere trasparenti riguardo agli usi ad alto rischio dei sistemi di IA. Ciò significa che il pubblico dovrebbe avere accesso alle informazioni sul fatto che il sistema di IA è in uso e per cosa viene utilizzato.

3.5.5 Sistemi biometrici AI ad alto rischio

I sistemi biometrici sono programmi informatici che utilizzano le caratteristiche fisiche o comportamentali di una persona, come il viso, le impronte digitali, la voce o l'iride, per identificarla o verificarla. Se sblocchi il telefono con una scansione del tuo viso o con un'impronta digitale, stai utilizzando un sistema di identificazione biometrica ogni giorno.

³⁵ [Matt Wood, "L'algoritmo di intelligenza artificiale prevede i crimini futuri con una settimana di anticipo con una precisione del 90%", SciTech Daily, 2022 \[consultato il 6 ottobre 2024\].](#)

³⁶ [Matt Wood, "L'algoritmo prevede il crimine con una settimana di anticipo, ma rivela pregiudizi nella risposta della polizia", The University of Chicago, 2022 \[consultato il 6 ottobre 2024\].](#)

³⁷ [Naiara Bellio, "La polizia spagnola prevede di estendere l'uso della sua macchina della verità mentre l'efficacia non è chiara", AlgorithmWatch, 2020 \[consultato il 6 ottobre 2024\].](#)

³⁸ [PolyGR, "Uso del poligrafo digitale nei dipartimenti di polizia: migliorare gli interrogatori delle forze dell'ordine utilizzando l'intelligenza artificiale", PolyGR, 2023 \[consultato il 6 ottobre 2024\].](#)

³⁹ [Keith Brannon, "Le condanne dell'intelligenza artificiale riducono il tempo di detenzione per i trasgressori a basso rischio, ma uno studio rileva che i pregiudizi razziali persistono", Tulane University News, 2023 \[consultato il 6 ottobre 2024\].](#)

⁴⁰ [Tzu-Wei Hung e Chun-Ping Yen, "Polizia predittiva e equità algoritmica", *Sintesi*, 2016 \(2023\), p. 206, doi:10.1007/s11229-023-04189-0.](#)

I sistemi biometrici di intelligenza artificiale sono ad alto rischio perché possono invadere la privacy e possono essere imprecisi, soprattutto nei confronti di minoranze come le persone con disabilità e le persone di colore.⁴¹

Un esempio di ciò è che un gruppo chiamato ETICAS ha testato un algoritmo senza informare l'azienda che lo ha sviluppato. Si è scoperto che l'algoritmo non era in grado di riconoscere le persone con sindrome di Down. Ha giudicato male la loro età molte volte. Si stima che alcune persone abbiano 16 anni meno di. Si stima che altri avessero 23 anni in più della loro età effettiva. Questo può portare a problemi. Ad esempio, a un bambino potrebbero essere date cose che non sono adatte alla sua età. Oppure a un adulto potrebbe essere negato l'accesso ai servizi di cui ha bisogno. Questo dimostra quanto possa essere sbagliato il riconoscimento facciale per le persone con disabilità.⁴²

Esempi di applicazioni biometriche dell'intelligenza artificiale ad alto rischio:

Identificazione biometrica remota: Identificazione di persone a distanza con una videocamera.

Categorizzazione biometrica: raggruppamento delle persone in base a tratti come la razza o lo stato di salute (ad esempio, classificando una persona come "Femmina di origine asiatica, di circa 25 anni")

Riconoscimento delle emozioni: analisi delle espressioni facciali, della voce e di altri segnali per interpretare le emozioni

L'AI Act prevede restrizioni d'uso. Ciò significa che:

- I sistemi di riconoscimento delle emozioni sono vietati nei luoghi di lavoro e nelle scuole a meno che non vengano utilizzati per scopi medici.
- La polizia e le autorità per la migrazione possono utilizzare sistemi di riconoscimento delle emozioni, ma questi sono considerati ad alto rischio.

Raccomandazioni per i membri di EDF:

- **Sostenere il divieto dell'identificazione biometrica remota in tempo reale:** in collaborazione con le organizzazioni che tutelano i diritti fondamentali, sostenere un divieto nazionale

⁴¹ [Considerando 44](#) della legge sull'IA

⁴² [Fondazione ETICAS, Invisible No More: l'impatto del riconoscimento facciale sulle persone con disabilità](#), agosto 2023 [consultato il 6 ottobre 2024].

dell'identificazione biometrica remota in tempo reale negli spazi pubblici per le attività di polizia. Gli Stati membri che intendono consentire in tutto o in parte l'uso di questa tecnologia da parte delle autorità di contrasto devono adottare le leggi nazionali entro il 2 febbraio 2025.

3.5.6 IA ad alto rischio in materia di migrazione, asilo e controllo delle frontiere

I sistemi di intelligenza artificiale utilizzati nella migrazione, nell'asilo e nel controllo delle frontiere sono ad alto rischio perché possono mettere a repentaglio i diritti fondamentali delle persone emarginate.⁴³ L'uso di questi sistemi di IA è classificato come ad alto rischio solo se consentito dal pertinente diritto dell'UE o nazionale.

Esempi di applicazioni di intelligenza artificiale ad alto rischio:

Rilevamento delle bugie: sistemi di intelligenza artificiale che fungono da rivelatori di bugie per le persone che cercano di attraversare il confine.⁴⁴

Valutazione del rischio: valutazione dei potenziali rischi posti da persone che viaggiano o soggiornano in uno Stato membro

Elaborazione delle domande: gestione delle domande di asilo, visto o permesso di soggiorno

Identificazione: identificazione di persone nell'ambito della gestione della migrazione

Raccomandazioni per i membri di EDF:

Sostenere politiche governative che non utilizzino macchine della verità basate sull'intelligenza artificiale nei contesti migratori: i membri dell'EDF dovrebbero sostenere politiche che impediscano l'uso della macchina della verità nella migrazione, proteggendo le persone con disabilità e il background migratorio. Anche se l'AI Act non consente agli Stati membri di vietarlo attraverso la legislazione nazionale, i governi possono comunque adottare politiche e linee guida per evitare l'uso di questi sistemi di IA inaffidabili, che possono portare a gravi violazioni dei diritti umani.

⁴³ [Considerando 60](#) della legge sull'IA

⁴⁴ [Aggiornamento biometrico, "I viaggiatori nell'UE potrebbero essere soggetti alla macchina della verità AI"](#), 17 giugno 2024 [consultato il 6 ottobre 2024].

Date le serie preoccupazioni sul riconoscimento delle emozioni basato sull'intelligenza artificiale, diversi punti chiave sottolineano la necessità di cautela in particolare quando si tratta di persone con disabilità:

1. [Il considerando 44](#) della legge sull'IA fa riferimento alla discutibile base scientifica dei sistemi di IA per il riconoscimento delle emozioni, che può portare a discriminazioni. Tali tecnologie sono vietate nelle scuole e nei luoghi di lavoro, ma purtroppo, a causa di compromessi politici, questo divieto non è stato esteso alla polizia o alla migrazione, dove il rischio di gravi violazioni dei diritti umani rimane probabile.
2. Uno studio completo del 2019 contesta l'ipotesi che i movimenti facciali riflettano costantemente gli stati emotivi, notando variazioni significative tra culture, situazioni e individui. Questa mancanza di supporto scientifico nelle tecnologie di riconoscimento delle emozioni potrebbe portare a gravi errori e discriminazioni, soprattutto in settori delicati come la migrazione e il controllo delle frontiere.⁴⁵
3. Nel 2021 il relatore speciale delle Nazioni Unite sui diritti delle persone con disabilità ha messo in guardia contro l'uso del riconoscimento facciale e delle emozioni, citando i loro alti tassi di errore nel riconoscimento delle persone con disabilità, che possono portare a discriminazioni e violazioni dei diritti umani.⁴⁶
4. È stato dimostrato che il software di riconoscimento delle emozioni interpreta erroneamente i segnali grammaticali nella lingua dei segni (espressi dalle espressioni facciali) come emozioni negative come l'aggressività. Questi errori possono avere gravi implicazioni se le autorità pubbliche si affidano a tecnologie pseudoscientifiche per il processo decisionale.⁴⁷

⁴⁵ [Lisa Feldman Barrett e altri, "Espressioni emotive riconsiderate: sfide per dedurre le emozioni dai movimenti facciali umani", *La scienza psicologica nell'interesse pubblico*, 20.1 \(2019\), pp. 1–68, doi:10.1177/1529100619832930.](#)

⁴⁶ [Gerard Quinn, A/HRC/49/52: Intelligenza artificiale e diritti delle persone con disabilità - Rapporto del relatore speciale sui diritti delle persone con disabilità](#), 28 dicembre 2021, punti 34, 54, 65-66, 76(c) [consultato il 7 ottobre 2024].

⁴⁷ [Irene Rogan Shaffer, "Esplorare le prestazioni delle tecnologie di riconoscimento delle espressioni facciali sugli adulti sordi e sui loro figli", *Atti della 20a Conferenza Internazionale ACM SIGACCESS su Computer e Accessibilità*, 8 ottobre 2018, pp. 474–76, doi:10.1145/3234695.3240986.](#)

3.5.7 IA ad alto rischio nell'amministrazione della giustizia e nei processi democratici

I sistemi di IA nella giustizia e nei processi democratici sono ad alto rischio a causa del loro impatto significativo sulla democrazia, sulle libertà e sui diritti legali⁴⁸.

Esempi di applicazioni di intelligenza artificiale ad alto rischio:

- **Assistenza giudiziaria:** assistenza ai giudici o alle autorità giudiziarie nell'accertamento dei fatti, nella comprensione delle leggi e nella loro applicazione ai casi.
- **Influenza del voto:** influenza il comportamento di voto nelle elezioni o nei referendum.

Raccomandazioni per i membri di EDF:

- **Sostenere la trasparenza:** verificare se le leggi nazionali sulle procedure amministrative contengono già un forte diritto del singolo cittadino a una spiegazione significativa delle decisioni di un'autorità. In caso contrario, prendi in considerazione la possibilità di sostenere tale diritto che si applichi indipendentemente dal fatto che la decisione sia stata presa da un essere umano o con l'assistenza di o interamente attraverso un processo decisionale automatizzato. Puoi unire le forze con altre organizzazioni che tutelano i diritti fondamentali per promuovere questo obiettivo.

4. Rischio di trasparenza

Rischio di trasparenza I sistemi di IA sono quelli che non raggiungono la soglia di rischio elevato, ma richiedono comunque obblighi di trasparenza specifici ai sensi della legge sull'IA. Questi sistemi devono garantire che gli utenti siano consapevoli quando interagiscono con l'IA o quando i contenuti sono stati generati o manipolati dall'IA.

Requisiti fondamentali ai sensi [dell'articolo 50](#) e delle relative disposizioni:

- Gli utenti devono essere chiaramente informati quando interagiscono con un sistema di intelligenza artificiale (ad esempio, quando utilizzano un chatbot o un assistente virtuale).
- Quando l'intelligenza artificiale viene utilizzata per generare o manipolare contenuti, come testo, immagini, audio o video, gli

⁴⁸ [Considerando 61](#) della legge sull'IA

utenti hanno il diritto di sapere che il contenuto è generato o alterato dall'intelligenza artificiale.

- Questi obblighi di trasparenza sono concepiti per aiutare gli utenti a prendere decisioni informate.

Perché è importante? Quando le informazioni vere e quelle false si fondono, la fiducia in Internet diminuisce. Recenti rapporti dell'ottobre 2025 mostrano che i contenuti generati dall'intelligenza artificiale spesso si rivolgono a gruppi vulnerabili, ad esempio creando falsi account sui social media che fingono di essere sostenitori della disabilità o utilizzando deepfake per sfruttare le persone con sindrome di Down. Questi abusi minano la sicurezza e rendono difficile per le voci autentiche distinguersi. La trasparenza è fondamentale affinché tutti, in particolare le persone con disabilità, possano identificare informazioni accurate online.⁴⁹

Accessibilità e prospettiva dei diritti delle persone con disabilità: EDF è particolarmente preoccupata per il fatto che le persone con disabilità possano essere maggiormente a rischio di frode o inganno man mano che i contenuti generati dall'intelligenza artificiale diventano più sofisticati. Garantire la trasparenza è pertanto essenziale per consentire a tutti gli utenti, comprese le persone con disabilità, di riconoscere i contenuti generati dall'IA e proteggersi da potenziali danni.

5. Rischio minimo

I sistemi di IA a rischio minimo sono quelli che non rientrano nelle categorie di rischio vietato, ad alto rischio o di trasparenza. Si ritiene che abbiano un impatto minimo o nullo sui diritti o sulla sicurezza delle persone, quindi l'AI Act non impone loro regole aggiuntive: continuano ad applicarsi solo le leggi esistenti come il GDPR.

L'unica disposizione diretta della legge sull'IA è l' [articolo 4](#), che impone ai fornitori e ai dispiegatori di sistemi di IA di garantire l'alfabetizzazione dell'IA tra il loro personale e coloro che gestiscono o utilizzano i loro sistemi di IA.

Esempi:

⁴⁹ ["Perché le persone usano l'intelligenza artificiale per simulare disabilità come la sindrome di Down online - CBS News"](#), 6 maggio 2025 [consultato il 16 ottobre 2025]; [BBC Audio, Disabilità del deepfaking, Il podcast del documentario](#) (1 ottobre 2025) [consultato il 16 ottobre 2025].

- Filtri antispam nelle e-mail
- Bot nei videogiochi che giocano contro gli utenti
- Sistemi di raccomandazione che suggeriscono film (Netflix), video (YouTube) o canzoni (Spotify)

Questi sistemi sono intenzionalmente lasciati non regolamentati dall'AI Act per evitare oneri inutili per l'IA quotidiana a basso impatto.

6. IA per uso generale (GPAI)

L'IA per uso generico (GPAI) si riferisce a modelli di intelligenza artificiale avanzati progettati per eseguire un'ampia gamma di attività, piuttosto che essere limitati a una singola funzione specifica. La GPAI può essere considerata come la tecnologia sottostante o il "motore" che può essere incorporato in varie applicazioni e servizi.

Ad esempio, una compagnia aerea potrebbe utilizzare un modello GPAI per alimentare il proprio chatbot del servizio clienti. Invece di sviluppare un'intelligenza artificiale da zero e addestrarla a comprendere e generare il linguaggio, la compagnia aerea può utilizzare un modello GPAI esistente che possiede già queste capacità. La compagnia aerea fornisce quindi all'intelligenza artificiale informazioni sugli orari dei voli, sui prezzi dei biglietti e sulle politiche aziendali. Il chatbot basato su GPAI può rispondere alle domande dei clienti, assistere con le prenotazioni e fornire supporto, attingendo sia alle sue competenze linguistiche generali che alle informazioni specifiche fornite dalla compagnia aerea.

Questa flessibilità è il motivo per cui viene chiamato "general purpose": lo stesso modello GPAI può essere utilizzato in molti contesti diversi, come chatbot e strumenti di creazione di contenuti. Gli obblighi e i rischi associati alla GPAI dipendono da come viene utilizzata e dall'impatto che può avere in ciascuna situazione. Ad esempio, se è integrato in un chatbot utilizzato per intervistare i candidati come parte di un processo di reclutamento, tale implementazione del sistema GPAI diventa parte di un sistema ad alto rischio.

Ai sensi dell'AI Act, la GPAI non è una categoria di rischio separata; Al contrario, questi sistemi possono rientrare in qualsiasi livello di rischio (vietato, ad alto rischio, trasparente o minimo) a seconda del loro utilizzo e del potenziale impatto. La GPAI è soggetta a una serie di norme e obblighi propri, in particolare quando pone rischi sistemici a causa della sua portata o influenza. L'applicazione delle norme per l'IPP è gestita dalla Commissione europea, mentre gli altri sistemi di IA sono generalmente supervisionati dalle autorità nazionali.

6.1 Protezioni dalle minacce alla salute, alla sicurezza e alla democrazia umana

Alcuni sistemi di intelligenza artificiale per uso generale (GPAI) comportano "rischi sistemici" perché possono causare danni significativi alla salute pubblica, alla sicurezza, ai diritti umani o alla società.

La Commissione europea elaborerà **linee guida** per classificare tali modelli di IA in base al loro livello di rischio. La legge sull'IA fornisce alcune indicazioni sui criteri e i requisiti dei sistemi di IA GPAI.

Una semplice regola empirica è che se la tua organizzazione utilizza un sistema di intelligenza artificiale che viene utilizzato solo internamente e non condiviso con altri, non si tratta di un modello su larga scala. Tuttavia, se Google o OpenAI pubblicano un sistema GPAI che è pubblicamente disponibile o utilizzato da molti altri, è probabile che si tratti di un sistema di intelligenza artificiale su larga scala.

I modelli linguistici di IA, che ai sensi della legge sull'IA dell'UE sono classificati come sistemi di IA generici, sono noti per inventare fatti e possono anche produrre contenuti dannosi.⁵⁰ Se un sistema GPAI ha molti utenti o è integrato in altri sistemi informatici, esiste il rischio che l'output dannoso possa moltiplicarsi e influire su altri sistemi informatici che dipendono dal sistema GPAI.

La legge sull'IA impone ai fornitori di modelli di IA l'obbligo di garantire che dispongano di solidi meccanismi di moderazione e controllo per evitare che i risultati dannosi vengano trasmessi all'utente. [II considerando 110](#) fornisce esempi di rischi sistemici, quali la generazione e la diffusione di disinformazione, contenuti discriminatori o che incitano all'odio, la creazione di strumenti di pirateria informatica, istruzioni su come costruire una bomba e simili.

Ecco alcuni esempi di vita reale che probabilmente rientrerebbero nei rischi sistemici:

- Un chatbot che ha suggerito a un uomo belga di togliersi la vita per salvare il pianeta, è finito in tragedia.⁵¹

⁵⁰ ["Quando l'IA sbaglia: affrontare le allucinazioni e i pregiudizi dell'IA", Sloan School of Management, Massachusetts Institute of Technology, 2023 \[consultato il 6 ottobre 2024\].](#)

⁵¹ ["Imane El Atillah, 'L'uomo pone fine alla sua vita dopo che un chatbot AI lo ha incoraggiato' a sacrificarsi per fermare il cambiamento climatico", 31 marzo 2023 \[consultato il 6 ottobre 2024\].](#)

- Un chatbot ha incoraggiato un uomo a uccidere un capo di stato.⁵²
- Un chatbot è stato utilizzato per creare un virus informatico.⁵³
- Generazione di voci IA che vengono utilizzate per impersonare qualcuno.⁵⁴
- Chatbot che possono rafforzare i deliri nelle persone che sono a più alto rischio di sviluppare psicosi.⁵⁵
- Chatbot AI che incoraggia un adolescente con autismo ad alto funzionamento a farsi del male e ad agire violentemente dopo che i genitori volevano limitare il suo tempo davanti allo schermo.⁵⁶

I fornitori di modelli di intelligenza artificiale GPAI su larga scala hanno il dovere di implementare misure di salvaguardia per mitigare i rischi sistemici, il che include l'affrontare i pregiudizi nei confronti delle persone con disabilità. I modelli linguistici spesso mostrano distorsioni in quanto vengono addestrati su dati Internet in cui il linguaggio abilista purtroppo domina ancora.⁵⁷

Essi devono rispettare le norme supplementari di cui **agli articoli 51, 52 e 55. Il considerando 110** chiarisce che la lotta contro le distorsioni è parte integrante dell'obbligo di attenuare il rischio sistemico.⁵⁸

6.2 Orientamenti sull'ambito di applicazione degli obblighi per i modelli di IA per uso generale⁵⁹

La Commissione europea ha pubblicato [linee guida sull'ambito di applicazione degli obblighi per i modelli di IA per uso generale](#)

⁵² [euronews: "L'uomo "incoraggiato" dal chatbot AI "fidanzata" a uccidere la regina Elisabetta II riceve una condanna al carcere"](#), 6 ottobre 2023 [consultato il 6 ottobre 2024].

⁵³ [EFE, "Uomo giapponese arrestato per aver creato un virus informatico utilizzando l'intelligenza artificiale generativa", 28 maggio 2024](#) [consultato il 6 ottobre 2024].

⁵⁴ [CNBC, "I truffatori possono utilizzare gli strumenti di intelligenza artificiale per clonare le voci di te e della tua famiglia: come proteggerti", CNBC, 2024](#) [consultato il 6 ottobre 2024].

⁵⁵ [Madeleine Finlay, Ross Burns, e presentato da Madeleine Finlay; prodotto da Tom Woolfenden; sound design di Ross Burns; produttrice esecutiva Kate Lamble, "Psicosi dell'intelligenza artificiale": i chatbot potrebbero alimentare il pensiero delirante? – Podcast Scienza Il Guardiano, 28 agosto 2025](#) [consultato il 29 agosto 2025].

⁵⁶ <https://edition.cnn.com/2024/12/10/tech/character-ai-second-youth-safety-lawsuit/index.html> e <https://news.bloomberglaw.com/litigation/autistic-teens-family-says-ai-bots-promoted-self-harm-murder>

⁵⁷ [Vinitha Gadiraju e altri, ""Non direi offensivo ma...": prospettive centrate sulla disabilità su modelli linguistici di grandi dimensioni", Conferenza ACM 2023 su equità, responsabilità e trasparenza, 12 giugno 2023, pp. 205–16, doi:10.1145/3593013.3593989.](#)

⁵⁸ [Considerando 110](#) della legge sull'IA

⁵⁹ [Linee guida sull'ambito di applicazione degli obblighi per i fornitori di modelli di IA per uso generale ai sensi della legge sull'IA](#) [consultato il 22 settembre 2025].

[\(GPAI\)](#). Si tratta di linee guida tecniche e non vincolanti pubblicate il 18 luglio 2025.

Esistono per aiutare gli attori a capire quando un modello conta come modello GPAI e quindi rientra nelle regole GPAI dell'AI Act, e come capire quando si ritiene che un modello GPAI rappresenti un rischio sistemico. Questa nota è solo per consapevolezza. I sostenitori della disabilità non sono tenuti a leggere o padroneggiare il documento.

Ciò che la linea guida chiarisce, a colpo d'occhio:

- Quando un modello è un modello GPAI: spiega l'approccio della Commissione all'individuazione dei modelli GPAI e stabilisce come valutare concettualmente l'ambito di applicazione, in modo che i fornitori e le parti interessate possano riconoscere quando si applicano le norme in materia di GPAI.
- Quando un modello GPAI ha un rischio sistemico: descrive in dettaglio la classificazione dei modelli GPAI con rischio sistemico.

6.3 Codice di condotta GPAI⁶⁰

Il 10 luglio 2025 la Commissione europea ha pubblicato un codice di buone pratiche volontario per i modelli di IA per scopi generali. Il Codice è inteso come un modo pratico per i fornitori di modelli di dimostrare la conformità alla legge sull'intelligenza artificiale, piuttosto che come un nuovo obbligo.⁶¹ I fornitori che sottoscrivono il codice possono fare affidamento su di esso per dimostrare che stanno rispettando la legge sull'IA e ridurre i loro oneri amministrativi.

Il Codice non è obbligatorio. Non contiene obblighi aggiuntivi che vadano oltre la legge sull'IA. Fornisce ai fornitori di GPAI un modo semplificato per adempiere ai loro obblighi esistenti. Per ulteriori dettagli si rimanda alle [domande e risposte della Commissione sul codice di buone pratiche](#).

Cosa copre il Codice (tre capitoli):

⁶⁰ [Il codice di condotta generale sull'IA](#) [consultato il 28 luglio 2025].

⁶¹ [Articolo 53, paragrafo 4](#), [articolo 55, paragrafo 2](#), e [considerando 117](#)

- **Trasparenza:** documentazione e divulgazione per i fornitori di modelli.
- **Diritto d'autore:** misure pratiche per mettere in atto una politica conforme alla normativa dell'UE in materia di diritto d'autore.
- **Sicurezza e protezione** – stato delle pratiche per i fornitori di modelli GPAI con solo rischio sistemico.

Questi documenti sono molto tecnici. Facciamo riferimento a loro qui per consapevolezza, non per uno studio dettagliato. Non ci si aspetta che tu li legga o li padroneggi per sostenere l'IA inclusiva della disabilità nell'attuazione nazionale.

6.4 Codici di condotta volontari

La legge sull'IA promuove l'elaborazione di codici di condotta per le applicazioni di IA classificati come rischio limitato o minimo. [L'articolo 95](#) impone all'Ufficio per l'IA della Commissione europea e ai governi degli Stati membri di sostenere e promuovere lo sviluppo di questi codici.

Questi codici di condotta possono essere elaborati dai fornitori di IA, dalle organizzazioni che utilizzano l'IA o dalle associazioni di settore che rappresentano questi gruppi. Anche altre parti interessate, come le organizzazioni che tutelano i diritti fondamentali, i sindacati, i gruppi di consumatori e le istituzioni accademiche, possono essere invitate a partecipare all'elaborazione di questi codici.

Ai sensi dell' [articolo 95, paragrafo 2](#), tali codici dovrebbero avere obiettivi chiari ed essere in grado di misurare il successo. Un obiettivo particolare è la protezione dei gruppi emarginati, comprese le persone con disabilità. Ciò dimostra quanto sia importante considerare le esigenze e i diritti dei gruppi emarginati nello sviluppo e nell'uso delle tecnologie di IA.

6.5 Raccomandazioni per i membri del FES

I sistemi di IA per uso generale saranno in gran parte regolamentati dalle politiche dell'UE, lasciando ai membri dell'EDF uno spazio limitato per fare pressione a livello nazionale. Una questione importante sono tuttavia i codici di condotta.

Poiché la protezione dei gruppi emarginati e la garanzia dell'accessibilità dell'IA sono obiettivi di questi codici di condotta esplicitamente menzionati nell' [articolo 95, paragrafo 2, lettera e\)](#), della legge sull'IA, le

organizzazioni per i diritti delle persone con disabilità hanno l'opportunità di svolgere un ruolo attivo nel loro sviluppo.

Se la tua organizzazione desidera partecipare allo sviluppo di questi codici di condotta, devi contattare le organizzazioni dei [consumatori](#) del tuo paese.

Poiché l'AI Act funziona come una legge sulla sicurezza dei prodotti, le organizzazioni per i diritti dei consumatori dispongono di una preziosa esperienza che può essere utile alle organizzazioni per i diritti dei disabili coinvolte nello sviluppo di codici di condotta nei loro paesi membri.

È possibile contattare la segreteria del Forum se si desidera contribuire allo sviluppo del codice di condotta dell'UE. Contatta il ministero competente a livello nazionale che redigerà il codice di condotta nazionale mostrando il tuo interesse, e i consumatori e le organizzazioni per i diritti digitali nel tuo paese.

7. Diritti fondamentali, sandbox e test

7.1 Il diritto di reclamo

Ai sensi [dell'articolo 85](#) della legge sull'IA, sia le persone fisiche che le organizzazioni (come le imprese o le organizzazioni senza scopo di lucro) hanno il diritto di presentare reclami se ritengono che un sistema di IA abbia violato la legge. Ciò consente alle persone di segnalare problemi sospetti all'autorità di regolamentazione competente, nota come autorità di vigilanza del **mercato**. Tuttavia, è importante notare che l'autorità deciderà se il reclamo richiede ulteriori indagini.

Non si tratta di una **disposizione sui diritti umani**, ma di un canale formale per affrontare le preoccupazioni relative ai sistemi di IA, in particolare quelli classificati come ad alto rischio. Ad esempio, se una persona con disabilità ritiene che un sistema di intelligenza artificiale la stia discriminando, può sollevare le proprie preoccupazioni attraverso questo processo di reclamo. L'obiettivo è fornire un meccanismo che consenta a chiunque sia colpito da pratiche dannose di IA di agire.

Raccomandazione ai membri del FES:

Insieme al diritto dei singoli di ottenere una spiegazione per i risultati di un sistema di IA ad alto rischio, il diritto di presentare un reclamo può essere una possibile via per i membri di EDF e altre organizzazioni per i diritti umani per agire contro i sistemi di IA dannosi.

7.2 Segnalazione di incidenti gravi

I fornitori sono tenuti a segnalare all'autorità di vigilanza del mercato gli incidenti gravi che coinvolgono sistemi di IA ad alto rischio e GPAI. Gli obblighi di comunicazione variano a seconda della classe di rischio assegnata.

Sistemi di intelligenza artificiale ad alto rischio⁶²

I sistemi di IA ad alto rischio devono essere segnalati quando causano o rischiano di causare gravi danni. Gli scopi principali di questa segnalazione sono avvertire le autorità in caso di modelli pericolosi, garantire che i fornitori e i distributori si assumano la responsabilità, consentire alle autorità di agire rapidamente e creare fiducia nel pubblico attraverso la trasparenza.

La segnalazione riguarda due settori principali: gli incidenti gravi e le infrazioni diffuse. Un incidente grave è un evento che provoca il decesso, gravi danni alla salute, gravi perturbazioni di infrastrutture critiche, gravi danni alla proprietà o all'ambiente o violazioni su larga scala dei diritti fondamentali. Gli esempi includono un sistema di intelligenza artificiale che fornisce a un medico informazioni mediche errate, portando un paziente a ricevere una diagnosi o un trattamento sbagliato, negando ingiustamente un prestito a qualcuno o escludendo candidati qualificati in base al sesso o all'etnia.

Un'infrazione diffusa si verifica quando un sistema di IA provoca contemporaneamente gravi danni a gruppi di persone in diversi paesi dell'UE. Ad esempio, se lo stesso problema con un sistema di intelligenza artificiale colpisce persone in due o più paesi, o se la stessa pratica illegale si verifica in almeno tre paesi, è considerato diffuso. Ai fini della segnalazione, deve anche trattarsi di un incidente grave, il che significa che causa o potrebbe causare danni effettivi, ad esempio alla salute, alla sicurezza, ai diritti o all'ambiente delle persone. Se è solo diffuso ma non grave, non deve essere segnalato.

Fornitori di sistemi GPAI⁶³

Se un fornitore offre un modello di intelligenza artificiale generico che potrebbe avere un impatto significativo sulla società (denominato "rischio sistemico"), deve tenere registri e segnalare eventuali problemi o incidenti gravi che coinvolgono il suo modello. Tale relazione dovrebbe essere

⁶² Commissione europea [«PROGETTO DI LINEE GUIDA ARTICOLO 73 DELLA LEGGE SULL'IA - SEGNALAZIONE DEGLI INCIDENTI \(SISTEMI DI IA AD ALTO RISCHIO\)»](#), 26 settembre 2025.

⁶³ [Articolo 55, paragrafo 1](#) e considerando [114](#) e [115](#) AI Act e [Il codice di buone pratiche per l'IA per uso generale in materia di sicurezza](#), Impegno 9.

presentata tempestivamente all'Ufficio per l'IA dell'UE e, se necessario, alle autorità nazionali.

7.3 Sandbox normativi

Uno **sandbox normativo** è un ambiente supervisionato in cui i sistemi di IA possono essere testati e convalidati, spesso in condizioni reali, sotto la supervisione dell'autorità di regolamentazione competente. Le sandbox forniscono agli sviluppatori un quadro per sperimentare e innovare, mentre le autorità di regolamentazione garantiscono che i sistemi di intelligenza artificiale soddisfino gli standard etici e di sicurezza.

Pensalo come una sandbox in cui i bambini (sviluppatori di intelligenza artificiale) giocano con i giocattoli (sistemi di intelligenza artificiale in fase di creazione) sotto la supervisione di adulti (regolatori). Queste autorità garantiscono che l'IA sia sicura e forniscono orientamenti sull'interpretazione della legge sull'IA, ma possono anche intervenire. Le sandbox offrono chiarezza legale e supervisione agli sviluppatori di IA e consentono sia agli inventori che alle autorità di regolamentazione di saperne di più su come funziona l'IA, migliorando la comprensione per tutte le parti coinvolte.

Ai sensi **degli articoli [da 57 a 59](#)** della legge sull'IA, le sandbox sono disponibili per tutti i tipi di sistemi di IA, comprese le IA **ad alto rischio** e **l'IA per uso generale (GPAI)**.

Le sandbox sono disciplinate dagli **articoli [57-59](#)**. Poiché queste regole non menzionano esplicitamente quali categorie di rischio possono utilizzare le sandbox, è probabile che in teoria tutti i livelli di rischio possano utilizzarle. Tuttavia, quando si guarda a ciò che avviene nella sandbox, un'ipotesi plausibile è che sarà principalmente rilevante dal punto di vista pratico per la GPAI e l'IA ad alto rischio.

Quando un'azienda opera all'interno di un ambiente sandbox, riceve indicazioni concrete dall'autorità di vigilanza del mercato su come applicare la legge nella pratica, mentre l'autorità di vigilanza sta imparando di più sulla tecnologia.

Gli spazi di sperimentazione sono istituiti a livello nazionale da ciascun paese dell'UE. Tuttavia, un paese può anche decidere di istituire uno spazio di sperimentazione congiunto con un altro paese dell'UE. Inoltre, un paese può decidere di creare sandbox a livello locale o regionale, se lo desidera.

L'autorità di un paese che istituisce spazi di sperimentazione è incoraggiata a cercare di coinvolgere altri attori nell'ecosistema dell'IA,

come le organizzazioni che tutelano i diritti fondamentali, quando ciò è opportuno⁶⁴.

Raccomandazione ai membri del FES:

In questo caso, i membri di EDF possono raccomandare ai loro governi e all'autorità di vigilanza del mercato di considerare l'impatto dei sistemi di IA sulle persone con disabilità all'interno degli spazi di prova.

7.4 Valutazioni d'impatto sui diritti fondamentali (FRIA)

Per garantire che i **sistemi di IA ad alto rischio** non violino i diritti fondamentali, la legge sull'IA impone agli sviluppatori di condurre una **valutazione d'impatto sui diritti fondamentali (FRIA)** prima di implementare tali sistemi. Sono esentati i sistemi di intelligenza artificiale utilizzati nei sistemi di infrastrutture critiche che gestiscono il traffico stradale, l'acqua, il gas e l'elettricità. Ciò è delineato nell' [articolo 27](#).

Un FRIA è un controllo per determinare se un sistema di IA possa danneggiare i **diritti fondamentali delle persone**. Gli sviluppatori devono consultare i rappresentanti dei gruppi emarginati, come le organizzazioni dei disabili, per comprendere l'impatto. Ad esempio, se un governo vuole utilizzare un sistema di intelligenza artificiale ad alto rischio, dovrebbe chiedere un feedback alle organizzazioni dei disabili per garantire che sia equo e accessibile.⁶⁵

Le organizzazioni che tutelano i diritti fondamentali a livello dell'UE stanno sviluppando i propri modelli per i FRIA basati sulla protezione dei diritti umani. I membri di EDF possono svolgere un ruolo importante nel raccomandare tali raccomandazioni ai propri governi.

Nel settembre 2025, la Federazione spagnola dei consumatori e degli utenti ha pubblicato un rapporto di ricerca che analizza l'attuale panorama delle valutazioni d'impatto sui diritti fondamentali (FRIA) ai sensi della legge sull'intelligenza artificiale, nonché di altri quadri strettamente correlati come il GDPR, che richiede valutazioni d'impatto sulla protezione dei dati. La relazione analizza esempi europei e spagnoli con l'obiettivo di evidenziare e proporre le migliori pratiche che consentano una valutazione significativa dell'impatto dei sistemi di IA sui diritti fondamentali.

Attraverso questo approccio comparativo, il rapporto fornisce raccomandazioni concrete per garantire che i FRIA e processi analoghi

⁶⁴ [Considerando 139](#) della legge sull'IA

⁶⁵ [Considerando 96](#) della legge sull'IA

tutelino realmente i diritti degli individui e non siano ridotti a mere formalità.⁶⁶

Raccomandazione ai membri del FES:

- **Sostenere una forte FRIA:**

A partire da ottobre 2025, la Commissione europea non ha ancora pubblicato il modello per le valutazioni d'impatto sui diritti fondamentali (FRIA) richieste ai sensi della legge sull'IA. Sebbene sia ampiamente previsto che un tale modello verrà prodotto, la tempistica rimane incerta.

Nel frattempo, è importante che le organizzazioni che tutelano i diritti fondamentali continuino a sostenere un processo FRIA forte e significativo. I membri nazionali del FES possono svolgere un ruolo chiave individuando i governi nazionali e le istituzioni per i diritti umani che li sostengono e incoraggiandoli a impegnarsi a sostenere standard solidi, in particolare quelli sviluppati dalla società civile, quando la Commissione si consolerà sul modello FRIA. Questa costante attività di advocacy contribuirà a garantire che il processo FRIA non diventi una mera formalità, ma fornisca invece una reale protezione dei diritti fondamentali ai sensi della legge sull'IA.⁶⁷

- A partire dal 2026: le organizzazioni nazionali dei disabili possono impegnarsi in modo coordinato con le autorità nazionali che implementano sistemi di IA ad alto rischio o con le imprese private obbligate a condurre valutazioni d'impatto ambientale (ad esempio banche e compagnie di assicurazione) per integrare il modello per le valutazioni d'impatto elaborate dalle organizzazioni che tutelano i diritti fondamentali e chiedere l'istituzione di strutture significative per il coinvolgimento dei portatori di interessi.⁶⁸

⁶⁶ [Relazione: Verso un'attuazione significativa dell'articolo 27 ai sensi della legge sull'IA - Versione italiana - CECU](#)

⁶⁷ [Karolina Iwańska e altri, Verso una legge sull'IA al servizio delle persone e della società: strategie di applicazione per una regolamentazione dell'IA che rispetti i diritti \(Centro europeo per il diritto senza scopo di lucro \(ECNL\), agosto 2024\), p. 31 \[consultato il 6 ottobre 2024\].](#)

⁶⁸ [Iwańska e altri, Verso una legge sull'IA al servizio delle persone e della società: strategie di applicazione per una regolamentazione dell'IA che rispetti i diritti.](#)

7.5 Test in condizioni reali⁶⁹

I sistemi di intelligenza artificiale ad alto rischio possono essere testati nel mondo reale per comprenderne le prestazioni al di fuori degli ambienti controllati. Tali prove devono seguire un piano dettagliato ed essere autorizzate dall'autorità nazionale competente. Se l'autorità non risponde entro 30 giorni, i test sono comunque considerati autorizzati, a meno che la legge dello Stato membro non richieda un'autorizzazione esplicita.

Per proteggere tutti i partecipanti, in particolare i gruppi emarginati come le persone con disabilità, è necessario osservare alcune regole importanti. I partecipanti devono dare il consenso informato, cioè devono comprendere appieno il test e accettare di partecipare. Se decidono di ritirare la loro partecipazione, possono farlo in qualsiasi momento senza conseguenze negative e i loro dati personali devono essere cancellati. Devono essere adottate misure speciali per garantire la sicurezza e i diritti delle persone con disabilità durante i test.⁷⁰

Le autorità sono autorizzate a ispezionare e monitorare i test per garantire che tutti gli standard etici e di sicurezza siano soddisfatti. I fornitori sono responsabili di eventuali danni che si verificano durante i test e devono agire immediatamente in caso di problemi gravi. Ciò garantisce che i diritti e la sicurezza dei partecipanti, le persone con disabilità, siano sempre prioritari e tutelati durante i test dei sistemi di IA in condizioni reali.

Raccomandazione ai membri di EDF: Tenete presente che i test nel mondo reale sono qualcosa che esiste. Le norme su tali test sono disciplinate principalmente a livello dell'UE, con ulteriori orientamenti attesi dalla Commissione UE. Dato che le persone con disabilità sono specificamente menzionate come un gruppo che necessita di una protezione speciale, le organizzazioni di persone con disabilità possono istituire punti di contatto con le autorità di vigilanza del mercato. In questo modo, possono cercare la tua esperienza quando necessario. Assicurati di essere ricompensato per il tuo tempo e i tuoi sforzi.

8. Misure supplementari: «Incentivazione»

[Il considerando 142](#) incoraggia gli Stati membri a sostenere e promuovere la ricerca e lo sviluppo nel campo dell'intelligenza artificiale che ha un impatto positivo sulla società e sull'ambiente. Pone l'accento in particolare sui progetti che migliorano l'accessibilità per le persone con disabilità come un risultato importante. Le iniziative che perseguono tali

⁶⁹ [Articolo 60](#) della legge sull'IA

⁷⁰ [Articolo 60, paragrafo 4, lettera g\)](#), della legge sull'IA

obiettivi dovrebbero essere ammissibili al sostegno pubblico e finanziario dell'Unione.

Raccomandazione ai membri del FES:

I membri dell'EDF possono incoraggiare i loro governi a rendere disponibili tali finanziamenti, con particolare attenzione alla co-creazione con le persone con disabilità, ai pregiudizi e all'uso dell'IA come tecnologia assistiva. Potrebbero anche sottolineare che, mentre l'IA può mantenere molte promesse per l'inclusione delle persone con disabilità, molte iniziative di IA falliscono.

9. Struttura di governance

La legge dell'UE sull'IA ha una struttura di governance con attori a livello nazionale e dell'UE. Gli attori più importanti sono:

- **Ufficio per l'IA**: un'unità della Commissione che coordina e sostiene il consiglio per l'IA e le autorità nazionali. Emana inoltre orientamenti e atti delegati relativi al regolamento.⁷¹
- **Comitato per l'intelligenza artificiale**: un organo delle autorità nazionali di vigilanza del mercato e della Commissione. Fornisce consulenza e sostegno alla Commissione in merito al regolamento, ad esempio per quanto riguarda l'individuazione di nuovi sistemi di IA ad alto rischio, l'armonizzazione della vigilanza del mercato, lo scambio di informazioni e migliori pratiche e la promozione della cooperazione internazionale.⁷²
- **Forum consultivo**: un organismo che fornisce **competenze tecniche** al consiglio di **amministrazione dell'IA** e alla **Commissione europea**. È costituito da un gruppo equilibrato di parti interessate provenienti dall'industria, dalla società civile e dal mondo accademico. Il forum fornisce consulenza sull'attuazione della legge sull'intelligenza artificiale, promuove le migliori pratiche e sostiene la cooperazione internazionale. Comprende membri permanenti come l'Agenzia dell'Unione europea per la cibersicurezza (ENISA) e l'Agenzia per i diritti fondamentali e si riunisce almeno due volte l'anno per formulare raccomandazioni e pareri.⁷³
- **Panel scientifico**: un organismo composto da esperti indipendenti selezionati per le loro competenze in materia di IA. Supporta l'Ufficio per l'IA nell'individuazione **dei rischi sistemici** derivanti dall'IA di uso generale, nella consulenza sulle classificazioni ad alto rischio e

⁷¹ [Articolo 64](#) della legge sull'IA

⁷² [Articolo 65-66](#) della legge sull'IA

⁷³ [Articolo 67](#) della legge sull'IA

nello sviluppo di strumenti di valutazione. Il gruppo garantisce una consulenza obiettiva e assiste le autorità di vigilanza del mercato nell'applicazione delle normative sull'IA in tutta l'UE.⁷⁴

- **Garante europeo della protezione dei dati**: l'autorità per la protezione dei dati dell'UE funge anche da autorità di vigilanza del mercato per quanto riguarda l'uso dell'IA da parte delle istituzioni dell'UE. Il Garante fornisce orientamenti ai legislatori dell'UE sulle proposte di nuova legislazione dell'UE.⁷⁵
- **Autorità di vigilanza del mercato**: l'autorità o le autorità nazionali che controllano e applicano la regolamentazione. È autorizzato a effettuare ispezioni, richiedere informazioni, effettuare test, imporre misure correttive e irrogare sanzioni.⁷⁶
- **Autorità per la tutela dei diritti fondamentali**: l'autorità o le autorità nazionali responsabili della tutela dei diritti fondamentali e della non discriminazione, ad esempio il responsabile o i responsabili per le pari opportunità. Collabora con l'autorità di vigilanza del mercato e può accedere ai documenti dei fornitori di sistemi di IA ad alto rischio. Non ha il potere di adottare misure correttive.⁷⁷ Questo sarà un elemento chiave per i membri dell'EDF con cui costruire buone relazioni.

9.1 Autorità di vigilanza del mercato

Le autorità di vigilanza del mercato sono gli organismi ufficiali di ciascun paese dell'UE responsabili dell'applicazione della legge sull'IA. Il loro compito principale è monitorare i rischi dei sistemi di IA e garantire che i fornitori rispettino la legge. Queste autorità dispongono di ampi poteri: possono intervenire se i sistemi di IA presentano rischi o non soddisfano i requisiti della legge sull'IA, possono effettuare il monitoraggio a distanza e accedere alla documentazione, ai set di dati e persino al codice sorgente dei fornitori. Se necessario, possono proporre indagini congiunte con la Commissione europea, richiedere misure correttive e far rispettare le norme attraverso sanzioni.

Le autorità di vigilanza del mercato devono riferire ogni anno alla Commissione e alle autorità nazionali competenti, comunicando i risultati delle valutazioni del rischio e le misure adottate. Ci si aspetta che

⁷⁴ [Articoli da 68 a 69](#) della legge sull'IA

⁷⁵ [Articolo 70, paragrafo 9](#), della legge sull'IA

⁷⁶ [Articolo 70](#) della legge sull'IA

⁷⁷ [Articolo 77](#) della legge sull'IA

agiscano in modo indipendente, imparziale e senza pregiudizi in tutte le loro attività.

Le denunce dei singoli sono una parte essenziale di un'applicazione efficace. Chiunque ritenga che un sistema di IA abbia violato la legge sull'IA può presentare un reclamo all'Autorità di vigilanza del mercato. Queste denunce aiutano le autorità a identificare e risolvere rapidamente i problemi.

L'applicazione efficace dipende anche dalla cooperazione tra le autorità di vigilanza del mercato nell'UE. Queste autorità collaborano attraverso il comitato europeo per l'IA, che contribuisce a coordinare le competenze e l'azione di esecuzione.

Ciascuno Stato membro doveva nominare una o più autorità di vigilanza del mercato entro il 2 agosto 2025. Se un paese ha più autorità, deve designare un punto di contatto unico nazionale per la cooperazione con gli altri Stati membri, l'UE e il pubblico in generale.⁷⁸ Se uno Stato membro non designa un'autorità entro il termine stabilito, la Commissione può avviare una procedura formale di infrazione.

Un elenco dei punti di contatto unici di ciascuno Stato membro è disponibile sul sito web della Commissione UE: [Autorità di vigilanza del mercato ai sensi della legge sull'IA](#).

Per molte autorità di vigilanza del mercato, il controllo degli algoritmi di intelligenza artificiale sarà un compito completamente nuovo e impegnativo. Anche le autorità di protezione dei dati, abituate a proteggere la privacy come diritto umano ai sensi del regolamento generale sulla protezione dei dati, si trovano di fronte a una transizione importante e avranno bisogno di risorse e personale con competenze in sociologia, diritto, sicurezza informatica, tecnologia, non discriminazione, protezione dei dati e diritti fondamentali. Dato che l'IA ad alto rischio deve soddisfare norme obbligatorie di accessibilità, è probabile che la maggior parte delle autorità pubbliche inizialmente non avrà competenze in materia di accessibilità.

Le organizzazioni delle persone con disabilità hanno un ruolo cruciale da svolgere nel garantire che i governi riconoscano questa lacuna e includano esperti di accessibilità nei team delle autorità di vigilanza del mercato. La nostra esperienza vissuta e la nostra vasta competenza sono essenziali per rendere i sistemi di intelligenza artificiale veramente inclusivi ed efficaci per tutti, in particolare per le persone con disabilità. Per questo motivo, la nostra esperienza come organizzazioni di persone con disabilità sarà preziosa anche per le stesse autorità di vigilanza del mercato. I

⁷⁸Articolo 70.2 [e](#) considerando 153 [della legge sull'IA](#)

membri dell'EDF sono fortemente incoraggiati a **contattare queste autorità e a offrire in modo proattivo la loro esperienza e consulenza.**

9.1.1 Indirizzare i reclami: scegliere l'autorità giusta

A seconda del tipo di sistema di IA in questione, diverse autorità di vigilanza sono responsabili del monitoraggio:

- Se si tratta di un **sistema di IA vietato, di un sistema di IA ad alto rischio o di un rischio trasparente**, l'autorità di controllo designata in ciascuno Stato membro - **l'autorità di vigilanza del mercato - è responsabile della gestione e dell'applicazione dei reclami.**
- Al contrario, per i **sistemi di IA di uso generale**, la **Commissione europea** funge da autorità di vigilanza ed è l'organo competente responsabile della ricezione di reclami o preoccupazioni.

Quindi, se sospetti un problema con un sistema di IA, è importante indirizzare il reclamo all'autorità giusta: l'autorità nazionale di vigilanza del mercato per l'IA ad alto rischio o vietata e la Commissione europea per i sistemi di IA per uso generale. Ciò garantirà che la questione sia affrontata in modo efficiente da coloro che dispongono dei poteri e delle competenze adeguate.

Oltre alle autorità nazionali, l'UE dispone di una propria autorità di vigilanza del mercato per i **sistemi di IA utilizzati dalle istituzioni dell'UE**. In particolare, il **Garante europeo della protezione dei dati (GEPD)** agisce in qualità di autorità di vigilanza del mercato responsabile del monitoraggio dell'uso dell'IA nelle istituzioni dell'UE.

Categoria	Descrizione ed esempi	Autorità di esecuzione
Proibito (rischio inaccettabile)	Pratiche di IA vietate (ad es. manipolazione, sfruttamento, social scoring)	Autorità degli Stati membri
Alto rischio	Regolamentato in modo rigoroso (ad es. lavoro, istruzione, forze dell'ordine, servizi essenziali)	Autorità degli Stati membri

Categoria	Descrizione ed esempi	Autorità di esecuzione
Trasparenza/Rischio limitato	Obblighi di trasparenza (ad es. chatbot, deepfake, categorizzazione biometrica)	Autorità degli Stati membri
Rischio minimo	Nessun obbligo aggiuntivo (ad es. filtri antispam, IA per videogiochi), ad eccezione del requisito di alfabetizzazione in materia di IA	Autorità degli Stati membri
IA per uso generale (GPAI)	Può rientrare in qualsiasi categoria di rischio; soggette a norme specifiche in materia di GPAI e di norme in materia di rischio sistemico.	Commissione europea

9.2 Autorità che tutelano i diritti fondamentali

Ai sensi dell'articolo 77, le [autorità specifiche di ciascun paese hanno il compito](#) di garantire che i sistemi di IA non violino i diritti fondamentali delle persone, come il diritto a un trattamento equo e paritario e all'accessibilità. Può trattarsi, ad esempio, di un difensore civico o di un organo per l'uguaglianza. Queste autorità possono richiedere informazioni alle organizzazioni che producono o utilizzano sistemi di IA ad alto rischio. Le informazioni dovrebbero essere di facile lettura e comprensione in modo che l'autorità che tutela i diritti fondamentali possa verificare il rispetto del suo mandato.⁷⁹ Dovrebbe inoltre essere istituito un meccanismo per garantire che i fornitori e gli utilizzatori di sistemi di IA rispondano a tali richieste in modo tempestivo.⁸⁰

L'autorità di vigilanza del mercato è responsabile della verifica della conformità dei sistemi di IA ad alto rischio alle disposizioni della legge sull'IA. Se le informazioni che l'autorità riceve per proteggere i diritti fondamentali non sono sufficienti per determinare se il sistema di IA viola i diritti, può chiedere all'autorità di vigilanza del mercato di testare il

⁷⁹ [Articolo 77, paragrafo 1](#), della legge sull'IA

⁸⁰ [Considerando 157](#) della legge sull'IA

sistema di IA in modo più tecnico e di vedere come si comporta il sistema.⁸¹

Gli Stati membri dovevano designare le rispettive autorità a tutela dei diritti fondamentali **entro il 2 novembre 2024**, come stabilito dall'[articolo 77](#) della legge sull'IA.

Un elenco aggiornato delle autorità che tutelano i diritti fondamentali in ciascuno Stato membro è disponibile sul sito web della Commissione europea: Autorità per la [tutela dei diritti fondamentali con poteri speciali ai sensi della legge sull'intelligenza artificiale](#).

9.3 Sanzioni per la violazione degli obblighi dell'AI Act⁸²

Le autorità nazionali devono stabilire norme per le sanzioni in caso di violazione della legge sull'IA da parte delle organizzazioni. Tali norme devono essere attuate entro il 2 agosto 2025.

- Se un'organizzazione viola le norme sulle pratiche vietate in materia di IA (articolo 5), l'autorità può imporre una sanzione pecuniaria fino a 35 milioni di euro o al 7% del fatturato globale (a seconda di quale sia il valore più alto).
- Se un'azienda viola altre regole, ad esempio per l'IA ad alto rischio o per uso generale, l'autorità può imporre una multa fino a 15 milioni di euro o al 3% del fatturato globale (a seconda di quale sia il valore più alto).
- Se un'azienda fornisce informazioni fuorvianti, incomplete o false all'autorità, la multa può arrivare fino a 7,5 milioni di euro o all'1% del fatturato globale (a seconda di quale sia il valore più alto).
- Ogni paese decide autonomamente se imporre multe alle proprie autorità in caso di violazione della legge sull'intelligenza artificiale.
- L'UE può imporre sanzioni pecuniarie fino a 1,5 milioni di euro alle proprie istituzioni per violazioni delle norme sull'intelligenza artificiale vietata, o fino a 750 000 euro per altri reati.

Nel determinare l'importo, l'autorità considera:

- La gravità del reato e le sue conseguenze.
- Le dimensioni e il fatturato dell'azienda.
- Che si tratti di una start-up o di una piccola impresa.
- Indipendentemente dal fatto che il reato sia stato commesso intenzionalmente o accidentalmente.
- Quanto bene l'azienda ha collaborato con l'autorità.

⁸¹ [Articolo 77, paragrafo 3](#), della legge sull'IA

⁸² [Articoli da 99 a 101](#) e considerando [168-169](#)

Le sanzioni devono essere eque ed efficaci e non devono essere troppo severe per le piccole imprese.

Raccomandazione ai membri del FES:

I membri dell'EDF devono contattare queste autorità. Nel vostro paese possono essere in corso anche iniziative volte a creare meccanismi coordinati di dialogo e cooperazione tra le organizzazioni che tutelano i diritti fondamentali e questo tipo di autorità. Sono loro che possono agire se ci troviamo di fronte alla discriminazione da parte di un sistema di intelligenza artificiale. Possono chiedere all'implementatore o al fornitore di un sistema di IA di spiegare come funziona l'IA e se segue l'AI Act. Abbiamo l'esperienza vissuta di vivere con una disabilità e possiamo dare loro input e feedback utili su come l'IA influisce sui nostri diritti e bisogni. Dovremmo collaborare con queste autorità e segnalare loro eventuali problemi con i sistemi di IA.

10. Monitoraggio della legge sull'IA nel tuo paese

10.1 Importanza del coinvolgimento delle persone con disabilità

La legge sull'IA dell'UE sta arrivando a casa tua e la sua attuazione è un momento cruciale per le persone con disabilità nel tuo paese. Questo è il momento giusto per i membri dell'EDF di partecipare ai preparativi dei loro governi per l'attuazione della legge sull'IA.

Ecco perché:

- **Prevenire la discriminazione:** l'intelligenza artificiale può rafforzare o eliminare i pregiudizi. Dobbiamo garantire che i sistemi di IA siano equi e inclusivi e tutelino i diritti delle persone con disabilità. La vostra organizzazione riunisce persone che hanno avuto l'esperienza di avere una disabilità. Avete a disposizione informazioni preziose di cui gli sviluppatori di sistemi di IA hanno bisogno e le autorità che monitorano che l'IA non discrimini.
- **Sostenere l'accessibilità:** l'AI Act include disposizioni per l'accessibilità. È importante che le autorità di vigilanza abbiano buoni rapporti con il movimento dei disabili, in modo da avere una conoscenza dettagliata delle esigenze delle persone con disabilità quando indagano sui reclami relativi ai sistemi di IA che non sono accessibili. Insieme possiamo garantire che l'IA sia utilizzabile da e per tutti.

- **Influenzare le politiche nazionali:** l'IA avrà un impatto su settori importanti come l'istruzione, la sicurezza sociale e la pubblica amministrazione. Le organizzazioni per le persone con disabilità devono influenzare queste politiche per garantire che soddisfino le esigenze delle persone con disabilità man mano che viene implementata una maggiore diffusione dell'intelligenza artificiale.

In questo capitolo riepiloghiamo l'entrata in vigore delle varie parti della legge sull'IA e le azioni che voi e la vostra organizzazione potete fare a livello nazionale.

10.2 Tempistica dettagliata dell'attuazione dell'AI Act:

1. **Pubblicazione nella Gazzetta ufficiale** (luglio 2024):
 - La legge sull'IA è ufficialmente [pubblicata](#)
2. **Entrata in vigore** (20 giorni dopo la pubblicazione: agosto 2024):
 - Entra in vigore la legge sull'IA.
3. **Divieti di pratiche vietate** (6 mesi dopo l'entrata in vigore: febbraio 2025):
 - Diventano applicabili divieti specifici su determinate pratiche di IA.
4. **Codici di condotta** (9 mesi dopo l'entrata in vigore: maggio 2025):
 - Inizia l'attuazione dei codici di condotta per i sistemi di IA.
5. **Norme e governance per l'IA per uso generale** (12 mesi dopo l'entrata in vigore: agosto 2025):
 - Entrano in vigore le regole e le strutture di governance per l'IA per uso generale.
6. **Piena applicabilità della legge sull'IA** (24 mesi dopo l'entrata in vigore: agosto 2026):
 - L'intera legge sull'IA diventa pienamente applicabile, ad eccezione degli obblighi per l'IA ad alto rischio.
7. **Obblighi per i sistemi ad alto rischio** (36 mesi dopo l'entrata in vigore: agosto 2027):
 - Iniziano ad applicarsi obblighi specifici per i sistemi di IA ad alto rischio.

Nota:

Al momento dell'aggiornamento della presente guida (ottobre 2025), erano in corso discussioni sulla possibilità di sospendere o ritardare l'entrata in vigore delle parti della legge sull'IA che non sono ancora diventate applicabili, in particolare le disposizioni relative ai sistemi di IA ad alto rischio. Ciò significa che le norme e gli obblighi per l'IA ad alto rischio, che dovevano entrare in vigore nell'agosto 2027, potrebbero

essere rinviati. Ci sono state anche forti indicazioni che la Commissione europea stia valutando la possibilità di introdurre deregolamentazioni o semplificazioni all'AI Act. Le proposte su queste modifiche sono attese a ottobre, novembre o dicembre 2025, dopo che le linee guida aggiornate saranno state finalizzate. Il sito web del Forum pubblicherà informazioni su eventuali sviluppi non appena si verificano.⁸³

10.3 Cosa bisogna fare? (Misure da adottare per l'attuazione a livello nazionale)

Per attuare con successo la legge sull'IA a livello nazionale, gli Stati membri devono adottare un approccio globale e coordinato che affronti vari settori critici. Ecco cosa è richiesto:

- 1. Allineamento legislativo:** gli Stati membri devono aggiornare le loro leggi per conformarsi alla legge sull'IA, che fornisce un quadro giuridico per la coerenza. Ad esempio, mentre la legge sull'IA delinea le condizioni per l'uso del riconoscimento facciale da parte della polizia, richiede una legislazione nazionale specifica per autorizzarne l'uso. Una volta in vigore, l'AI Act regola come e perché il riconoscimento facciale può essere utilizzato in modo etico. In alternativa, i membri possono scegliere di vietarlo completamente nelle aree pubbliche.
- 2. Organismi di regolamentazione e vigilanza:** gli Stati membri dovrebbero istituire o designare autorità di regolamentazione incaricate di vigilare sulle attività di IA entro il 2 agosto 2025. Questi organismi devono garantire una solida vigilanza, in particolare per i sistemi di IA ad alto rischio, al fine di prevenire l'uso improprio e sostenere le norme etiche. Le autorità di controllo dovranno cooperare con i difensori civici per la parità e le autorità di protezione dei dati per affrontare efficacemente le questioni di pregiudizio e discriminazione.

Inoltre, incorporare l'esperienza di persone che hanno vissuto l'esperienza di avere una disabilità è fondamentale per queste autorità per comprendere e affrontare le sfide particolari affrontate dai membri della nostra comunità.

10.4 Raccomandazioni ai membri del FES

Ecco i passi chiave che le organizzazioni di persone con disabilità possono intraprendere per prepararsi all'attuazione dell'AI Act nel loro paese:

⁸³ [«L'UE si prepara a una pausa per l'applicazione delle norme sull'IA – POLITICO»](#), 23 settembre 2025 [consultato il 16 ottobre 2025].

1. Scopri di più sull'AI Act e spiega ai tuoi membri cos'è l'AI Act e cosa significa per le persone con disabilità.

2. Rimanere aggiornato sul lavoro di attuazione del suo governo?

Ad esempio, puoi tenere d'occhio se il tuo governo sta pianificando una nuova legislazione per consentire alle agenzie governative di utilizzare l'intelligenza artificiale:

- Valutazione delle domande di prestazioni statali
- Individuazione di frodi in materia di sicurezza sociale e prestazioni
- Attività di polizia e sorveglianza degli spazi pubblici
- Per la sorveglianza nell'assistenza, come il monitoraggio delle persone con demenza
- Istruzione per la valutazione, l'ammissione o per rilevare imbrogli durante gli esami
- Triage nell'assistenza sanitaria (ad esempio, l'introduzione di un chatbot come opzione per far chiamare le persone a una hotline per discutere con un infermiere se devono andare dal medico / al pronto soccorso)
- Il settore pubblico e il sistema giudiziario

3. Sensibilizzare.

È inoltre possibile sensibilizzare l'opinione pubblica sui potenziali benefici e sulle sfide derivanti dall'uso dei sistemi di IA per le persone con disabilità, come il miglioramento dell'accessibilità, della personalizzazione, dell'autonomia, della partecipazione e della dignità, ma anche la creazione di barriere, discriminazione, esclusione, stigmatizzazione e manipolazione. Puoi trovare maggiori informazioni su questi vantaggi e sfide nella sezione seguente.

4. Entra in contatto con le organizzazioni per i diritti digitali che lavorano sull'intelligenza artificiale.

EDRi è un ombrello con organizzazioni membri in diversi paesi. Qui puoi trovare un elenco delle organizzazioni membri di EDRi [Archivio - European Digital Rights \(EDRi\)](#).

5. Partecipa al dibattito pubblico e all'elaborazione delle politiche sull'IA sensibilizzando e mobilitando i tuoi membri e alleati per esprimere le loro opinioni e richieste sull'IA e sul suo impatto sulle persone con disabilità, ad esempio attraverso

campagne, petizioni, eventi, media, social media, ecc. Puoi trovare maggiori informazioni su come farlo nella sezione seguente.

- 6. Chiedete al vostro governo di seguire la raccomandazione del Relatore speciale sui diritti delle persone con disabilità**, che ha esortato gli Stati a sostenere le vostre organizzazioni nel monitoraggio e nella promozione di un'IA inclusiva della disabilità.

Dì al tuo governo che hai bisogno di risorse e formazione per garantire che l'IA rispetti e promuova i diritti delle persone con disabilità e che puoi contestare qualsiasi uso dannoso o discriminatorio dell'IA. Hai il diritto di partecipare allo sviluppo e alla governance dell'IA e il tuo governo dovrebbe aiutarti a farlo.⁸⁴

- 7. Richiedi finanziamenti alle organizzazioni che mirano a rafforzare le capacità di advocacy delle organizzazioni che tutelano i diritti fondamentali quando si tratta di IA.**

[Ad esempio, il Fondo europeo per l'IA e la società](#)

- 8. Identificare le autorità competenti nel proprio paese che tutelano i diritti fondamentali e le autorità di vigilanza del mercato**
- 9. Sostenere i finanziamenti per l'uso dell'IA che dia potere alle persone con disabilità.**
- 10. Sostenere il finanziamento delle organizzazioni di persone con disabilità per lavorare sulla difesa dell'IA.**

⁸⁴ [Gerard Quinn, A/HRC/49/52: Intelligenza artificiale e diritti delle persone con disabilità - Rapporto del relatore speciale sui diritti delle persone con disabilità](#), 28 dicembre 2021, par. 76(g) [consultato il 7 ottobre 2024].

11. Valutazione delle carenze della legge sull'IA

11.1 La scelta di disegnare la legge in base ai rischi invece che ai diritti

Gli alleati di EDF per i diritti digitali hanno criticato l'approccio basato sul rischio. Sostengono che un approccio basato sui diritti sarebbe stato migliore.⁸⁵

Ciò è dovuto al fatto che le categorie di rischio previste dalla legge sull'IA sono definite categoricamente dai politici e il modo in cui sono definite nella legge sull'IA non protegge completamente i gruppi emarginati da danni reali.

Le persone con disabilità possono affrontare sia la discriminazione intersezionale che il danno cumulativo, due concetti correlati ma distinti che l'AI Act non affronta completamente.

Il danno cumulativo si riferisce agli effetti cumulativi di atti discriminatori più piccoli ripetuti che insieme causano un impatto negativo significativo. Questo concetto sottolinea come atti discriminatori apparentemente minori possano sommarsi e causare notevoli difficoltà. Ad esempio, le persone con disabilità possono incontrare numerosi ostacoli nel mercato del lavoro, come alloggi inadeguati, assunzioni distorte e limitate opportunità di sviluppo professionale, che insieme possono portare a una disoccupazione persistente o alla sottoccupazione.

L'intersezionalità esamina l'interazione di diversi motivi di discriminazione che fanno parte della propria identità, come la razza o l'etnia, il genere, l'età, la disabilità e l'orientamento sessuale, e il modo in cui questi influiscono collettivamente sulle esperienze di oppressione e vantaggio. Dimostra che l'identità complessa di una persona porta a diverse forme di discriminazione e privilegi che non sono riconoscibili

⁸⁵ ["L'UE dovrebbe regolamentare l'IA sulla base dei diritti, non dei rischi"](#), Access Now, 2021 [consultato il 6 ottobre 2024].

quando questi fattori sono considerati separatamente. Ad esempio, l'intersezionalità ci mostra che le donne con disabilità razzializzate sono spesso discriminate a più livelli nel loro ambiente professionale perché vengono trascurate per le promozioni o ricevono salari diseguali, più dei loro colleghi bianchi, maschi o non disabili.

Questi problemi evidenziano i limiti dell'attuale legge sull'IA, che si basa su un approccio basato sul rischio stabilito dai politici. Le categorie di rischio non proteggono completamente i gruppi emarginati da danni reali, in quanto tali questioni non sono pienamente contemplate dalla definizione di rischio contenuta nella legge. Tuttavia, l'AI Act è ciò che abbiamo ed è nostro compito sfruttarlo al meglio.⁸⁶

11.2 Accessibilità nell'intelligenza artificiale GPAI e rischio minimo

Sebbene [l'articolo 16\(I\)](#) dell'AI Act soddisfi la nostra richiesta di accessibilità affermando che i sistemi di IA ad alto rischio devono essere conformi ai requisiti di accessibilità, riteniamo comunque che sarebbe fondamentale estendere questi requisiti alla GPAI (General Purpose AI) e all'IA a basso rischio. Questo è importante perché ciò che potrebbe essere considerato a basso o medio rischio per la popolazione generale può in realtà essere ad alto rischio per le persone con disabilità.^{87 88}

Ciò è particolarmente spiacevole perché molti strumenti di intelligenza artificiale utilizzati da persone con disabilità, come assistenti personali, generatori di testo o voce, IA che descrivono il contenuto di un'immagine, dispositivi a controllo vocale e sistemi di intelligenza artificiale personalizzati, rientrano nelle categorie GPAI e a basso rischio. Senza requisiti di accessibilità obbligatori, questi strumenti non dovranno essere accessibili a meno che non siano coperti dalla direttiva sull'accessibilità del web o dall'atto europeo sull'accessibilità.

⁸⁶ [Access Now e altri, "Garantire la protezione dei diritti fondamentali nella posizione del Consiglio sulla legge sull'IA", 2023.](#)

⁸⁷ [Forum europeo sulla disabilità, Prospettiva della disabilità sulla regolamentazione dell'intelligenza artificiale, 8 novembre 2021](#), pp. 6-8 Si veda la sezione intitolata "IA affidabile".

⁸⁸ [European Disability Forum \(EDF\) e altri, "La legge sull'IA dell'UE non riesce a stabilire il gold standard per i diritti umani", Forum europeo sulla disabilità, 2024 \[consultato l'8 ottobre 2024\].](#)

12. Glossario

Questo glossario include termini che il Forum ritiene possano aiutare i non esperti a comprendere l'intelligenza artificiale. Non tutti sono legalmente definiti nell'AI Act, ma per quelli che lo sono, abbiamo incluso le loro definizioni legali.

Accessibilità

- **Significato:** Progettare sistemi di intelligenza artificiale in modo che possano essere utilizzati da tutti, comprese le persone con disabilità.
- **Esempio:** software di intelligenza artificiale compatibile con gli screen reader per utenti ipovedenti.

AI (Intelligenza Artificiale)

- **Significato:** tecnologia che consente alle macchine di imitare il comportamento umano, come prendere decisioni, riconoscere il parlato e tradurre le lingue.
- **Esempio:** gli assistenti virtuali come Siri e Alexa utilizzano l'intelligenza artificiale per comprendere e rispondere ai comandi degli utenti.

Legge sull'IA

- **Significato:** una legge dell'Unione Europea che regola l'uso dei sistemi di IA per proteggere i diritti umani e garantire la sicurezza pubblica.

Alfabetizzazione AI

- **Significato:** competenze, conoscenze e comprensione che consentono alle aziende che producono l'IA, alle organizzazioni che utilizzano l'IA e alle persone interessate dai sistemi di IA di prendere decisioni informate sull'uso dei sistemi di IA e di essere consapevoli delle opportunità, dei rischi e dei potenziali danni che l'IA può causare.
- **Definizione giuridica (articolo 3.56 della legge sull'IA):**
"Per alfabetizzazione in materia di IA si intendono le competenze, le conoscenze e la comprensione che consentono ai fornitori, ai responsabili dell'attuazione e alle persone interessate, tenendo conto dei rispettivi diritti e obblighi nel contesto del presente regolamento, di effettuare un'implementazione informata dei sistemi di IA,

nonché di acquisire consapevolezza in merito alle opportunità e ai rischi dell'IA e ai possibili danni che può causare."

Ufficio AI

- **Significato:** un'unità all'interno della Commissione europea con competenze in materia di IA che svolge un ruolo di coordinamento.
- **Definizione giuridica (articolo 3.47 della legge sull'IA):**

"Ufficio per l'IA: la funzione della Commissione di contribuire all'attuazione, al monitoraggio e alla supervisione dei sistemi di IA e dei modelli di IA per uso generale, nonché della governance dell'IA, di cui alla decisione della Commissione del 24 gennaio 2024; i riferimenti all'Ufficio per l'IA contenuti nel presente regolamento si intendono fatti alla Commissione."

Sistema di intelligenza artificiale

- **Significato:** una macchina che utilizza l'intelligenza artificiale per eseguire attività come previsioni, decisioni o raccomandazioni.
- **Esempio:** l'intelligenza artificiale utilizzata per ordinare le domande di lavoro o identificare i volti nelle telecamere di sicurezza.
- **Definizione giuridica (articolo 3.1 della legge sull'IA):**

"Per sistema di intelligenza artificiale si intende un sistema basato su una macchina progettato per funzionare con diversi livelli di autonomia e che può mostrare adattabilità dopo l'implementazione e che, per obiettivi espliciti o impliciti, deduce, dall'input che riceve, come generare output come previsioni, contenuti, raccomandazioni o decisioni che possono influenzare gli ambienti fisici o virtuali".

Algoritmo

- **Significato:** Un insieme di regole o passaggi seguiti da un computer per eseguire un'attività.
- **Esempio:** gli algoritmi decidono quali annunci vedere sui social media.

Bias (nell'intelligenza artificiale)

- **Significato:** quando un sistema di intelligenza artificiale favorisce o svantaggia ingiustamente determinati gruppi di persone a causa del modo in cui è stato addestrato o dei dati che gli sono stati forniti.

- **Esempio:** intelligenza artificiale, uno strumento di reclutamento che rifiuta spesso le candidature di persone con disabilità. Ciò accade perché i dati di formazione stabiliscono i tratti delle persone senza disabilità come standard ideale per chi è un buon candidato per il lavoro.

Sistema di intelligenza artificiale biometrica

- **Significato:** intelligenza artificiale che identifica le persone in base a caratteristiche fisiche o comportamentali come tratti facciali, impronte digitali o voce.
- **Esempio:** utilizzo del riconoscimento facciale per sbloccare un telefono e utilizzo da parte delle forze dell'ordine per identificare sospetti o localizzare persone scomparse.
- **Definizione giuridica (articolo 3.34 della legge sull'IA):**
 "Per dati biometrici si intendono i dati personali risultanti da un trattamento tecnico specifico relativo alle caratteristiche fisiche, fisiologiche o comportamentali di una persona fisica, come le immagini facciali o i dati dattiloscopici."

Distributore

- **Significato:** un'organizzazione o una persona che utilizza un sistema di intelligenza artificiale creato da qualcun altro.
- **Esempio:** EDF utilizza un chatbot basato sull'intelligenza artificiale sul nostro sito Web per rispondere a domande generali sulle nostre politiche, sul lavoro e su come contattare i dipendenti EDF. Utilizzando questo sistema di intelligenza artificiale, EDF diventa un distributore di un sistema di intelligenza artificiale.
- **Definizione giuridica (articolo 3.4 della legge sull'IA):**
 "Per Deployer si intende una persona fisica o giuridica, un'autorità pubblica, un'agenzia o un altro organismo che utilizza un sistema di IA sotto la sua autorità, tranne nel caso in cui il sistema di IA sia utilizzato nel corso di un'attività personale non professionale."

Atto delegato

- **Significato:** una legge approvata dalla Commissione europea che fornisce norme più dettagliate per l'attuazione dell'AI Act.

Diritti fondamentali

- **Significato:** All'interno del quadro giuridico dell'Unione Europea, i diritti fondamentali si riferiscono ai diritti umani tutelati dalla

"costituzione" dell'UE, in particolare dalla Carta dei diritti fondamentali dell'UE. Tra questi figurano vari diritti umani, come l'uguaglianza, l'accessibilità, la privacy e la libertà dalla discriminazione, che devono essere rispettati ogniqualvolta si applichi la legislazione dell'UE.

Modello di intelligenza artificiale per uso generico

- **Significato:** un modello di intelligenza artificiale addestrato su una grande quantità di dati in grado di eseguire bene molte attività diverse. Un modello è la tecnologia sottostante che può essere integrata in diverse applicazioni.
- **Esempio:** GPT-4 è un modello di intelligenza artificiale sviluppato dalla società [OpenAI](#) che funziona come tecnologia di base in altre applicazioni [Be my AI](#) e [Microsoft Co-pilot](#).

- **Definizione giuridica (articolo 3.63 della legge sull'IA):**

"Per modello di IA per uso generale si intende un modello di IA, anche nel caso in cui tale modello di IA sia addestrato con una grande quantità di dati utilizzando l'autosupervisione su larga scala, che mostri una generalità significativa ed è in grado di eseguire con competenza un'ampia gamma di compiti distinti indipendentemente dal modo in cui il modello è immesso sul mercato e che può essere integrato in una varietà di sistemi o applicazioni a valle. ad eccezione dei modelli di IA utilizzati per attività di ricerca, sviluppo o prototipazione prima di essere immessi sul mercato."

Sistema di intelligenza artificiale per uso generico

- **Significato:** un sistema di intelligenza artificiale che utilizza un modello di intelligenza artificiale generico e può essere utilizzato per molti scopi diversi, da solo o come parte di altri sistemi di intelligenza artificiale.
- **Esempio:** l'app [Be my AI](#) è un sistema di intelligenza artificiale che utilizza il modello di intelligenza artificiale generico chiamato GPT4 per descrivere le immagini alle persone non vedenti.

- **Definizione giuridica (articolo 3.66 della legge sull'IA):**

"Per sistema di IA generico si intende un sistema di IA basato su un modello di IA generico e che ha la capacità di servire a una varietà di scopi, sia per l'uso diretto che per l'integrazione in altri sistemi di IA".

Istruzioni

- **Significato:** raccomandazioni e spiegazioni emesse dalla Commissione europea per chiarire come attuare parti specifiche della legge sull'IA, come le pratiche vietate, gli obblighi di trasparenza o la definizione di un sistema di IA.
- **Esempio:** la Commissione può emanare orientamenti su come soddisfare i requisiti di accessibilità o su come interpretare la definizione di un sistema di IA.
- **Nota:** le linee guida non sono giuridicamente vincolanti, ma forniscono consigli pratici e riflettono l'interpretazione della legge da parte della Commissione.

IA ad alto rischio

- **Significato:** sistemi di intelligenza artificiale che possono avere un impatto significativo sulla vita delle persone, come quelli utilizzati nelle assunzioni, nell'assistenza sanitaria o nelle forze dell'ordine.
- **Esempio:** intelligenza artificiale che decide se qualcuno ottiene un prestito.

IA a basso rischio

- **Significato:** sistemi di intelligenza artificiale con un impatto minimo sulla vita delle persone, come quelli utilizzati nei videogiochi o nel filtro antispam.

Apprendimento automatico

- **Significato:** un tipo di intelligenza artificiale in cui i computer imparano dai dati per migliorare le loro prestazioni nelle attività.
- **Esempio:** un sistema di intelligenza artificiale impara a riconoscere meglio i volti nel tempo mostrando più esempi.

Autorità di vigilanza del mercato

- **Significato:** un'autorità governativa responsabile di garantire che i sistemi di intelligenza artificiale rispettino la legge e non danneggino gli utenti.
- **Definizione giuridica (articolo 3.26 della legge sull'IA):**
"Autorità di vigilanza del mercato: l'autorità nazionale che svolge le attività e adotta le misure di cui al regolamento (UE) 2019/1020."

Provider

- **Significato:** un'azienda o un individuo che crea e offre un sistema di intelligenza artificiale che può essere utilizzato da altri.
- **Esempio:** OpenAI ha creato ChatGPT che chiunque può utilizzare; OpenAI è un fornitore di un sistema di intelligenza artificiale.
- **Definizione giuridica (articolo 3.3 della legge sull'IA):**
"Per fornitore si 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 uso generale o che fa sviluppare un sistema di IA o un modello di IA per uso generale e lo immette sul mercato o mette in servizio il sistema di IA con il proprio nome o marchio, a titolo oneroso o gratuito."

Sandbox normativo

- **Significato:** un ambiente controllato in cui le aziende possono testare le nuove tecnologie di intelligenza artificiale sotto supervisione per garantirne la sicurezza e l'etica.
- **Definizione giuridica (articolo 3.55 della legge sull'IA):**
"Per sandbox normativo dell'IA si intende un quadro controllato istituito da un'autorità competente che offre ai fornitori o ai potenziali fornitori di sistemi di IA la possibilità di sviluppare, addestrare, convalidare e testare, se del caso in condizioni reali, un sistema di IA innovativo, conformemente a un piano di sandbox per un periodo di tempo limitato sotto supervisione regolamentare."

Categorie particolari di dati personali

- **Significato:** dati personali altamente sensibili che possono essere utilizzati in modo improprio per discriminare le persone. Ciò include informazioni sulla razza, la salute, le opinioni politiche, le convinzioni religiose, l'orientamento sessuale e altri dettagli intimi di una persona. La protezione di questi dati è fondamentale per prevenire gli abusi e garantire che le persone non vengano trattate ingiustamente o emarginate.
- **Collegamento all'antidiscriminazione:** la salvaguardia di queste categorie speciali di dati personali è direttamente collegata agli sforzi antidiscriminazione. Limitando l'accesso a tali informazioni sensibili e il loro trattamento, la legge sull'IA mira a prevenire decisioni e azioni distorte che potrebbero danneggiare individui o

gruppi, in particolare quelli con disabilità o appartenenti a comunità minoritarie.

- **Definizione giuridica (articolo 3.37 della legge sull'IA):**

"Per categorie particolari di dati personali si intendono le categorie di dati personali di cui all'articolo 9, paragrafo 1, del regolamento (UE) 2016/679, all'articolo 10 della direttiva (UE) 2016/680 e all'articolo 10, paragrafo 1, del regolamento (UE) 2018/1725."

Trasparenza (nell'IA)

- **Significato:** Rendere i sistemi di intelligenza artificiale chiari e comprensibili, in modo che le persone sappiano come e perché vengono prese le decisioni.
- **Esempio:** un'azienda spiega come la sua intelligenza artificiale decide quali candidati selezionare per un lavoro.

Rischio inaccettabile AI

- **Significato:** sistemi di IA vietati perché rappresentano un grave rischio per la dignità umana, la democrazia o i diritti fondamentali.
- **Esempio:** sistemi di intelligenza artificiale che manipolano le emozioni per influenzare le decisioni sul lavoro.

13. Crediti del documento

Questo documento è stato preparato da: [Kave Noori, responsabile delle politiche di intelligenza artificiale presso EDF](#) con il supporto dei colleghi del segretariato di EDF, tra cui Marine Uldry, Alejandro Moledo, André Felix e Natalia Suárez.



Forum Europeo sulla Disabilità
Mundo Madou
Avenue des Arts 7-81210
Bruxelles, Belgio.

www.edf-feph.org

info@edf-feph.org

Gli sforzi di advocacy e capacity building di EDF sull'Intelligenza Artificiale sono possibili solo grazie alle generose sovvenzioni che abbiamo ricevuto dal Fondo Europeo per l'IA e la Società.

European
**Artificial Intelligence
& Society Fund**

Documento tradotto in lingua italiana con uso di IA, pubblicato nel sito web webaccessibile.org. In caso di errori di traduzione, contattare info@webaccessibile.org.