

APPROFONDIMENTI

Il contenzioso nell'era dell'intelligenza artificiale

Rischi e responsabilità per imprese e
intermediari finanziari

Settembre 2026

Sara Biglieri, Partner, Chiomenti
Luca Ferrari, Partner, Chiomenti



Sara Biglieri, Partner, Chiomenti

Luca Ferrari, Partner, Chiomenti

> **Sara Biglieri**

Sara Biglieri è Partner e Co-Responsabile della Practice Area “Civil Litigation”. Assiste gruppi internazionali e primarie società operanti nei settori farmaceutico, grande distribuzione, beni di consumo, energy, immobiliare, finanziario e assicurativo nell’ambito di contenziosi giudiziali, incluse class action e arbitrati. Presta assistenza giudiziale e stragiudiziale nelle aree del diritto civile, commerciale, societario, immobiliare e bancario.

> **Luca Ferrari**

Luca Ferrari fornisce assistenza in controversie complesse sia dinanzi a giudici ordinari sia in procedimenti arbitrali, nei settori del diritto bancario e finanziario, dove rappresenta istituzioni finanziarie italiane e straniere, in particolare in controversie relative a prestazioni di servizi bancari e di investimento, operazioni di finanziamento, finanza strutturata, cartolarizzazioni, garanzie, nonché in materia esecutiva e d’insolvenza.

1. Dalla **compliance** alla prevenzione del contenzioso

L’intelligenza artificiale è ormai parte integrante dell’organizzazione delle imprese, dalla selezione del personale alla gestione della clientela, dalla valutazione del merito creditizio ai sistemi di supporto alle decisioni. Parallelamente, cresce il rischio di contenzioso. Le prime controversie internazionali dimostrano che la responsabilità non dipende esclusivamente dal malfunzionamento dell’algoritmo, ma coinvolge l’intero sistema di *governance* adottato dall’impresa.

In tale scenario si inserisce il nuovo quadro normativo europeo, costruito su strumenti tra loro complementari.

Il primo pilastro è rappresentato dal **Regolamento (UE) 2024/1689 (AI Act)**, entrato in vigore il 1° agosto 2024 e destinato a trovare applicazione progressiva. La disciplina regola l’immissione sul mercato, la messa in servizio e l’utilizzo dei sistemi di IA secondo un approccio basato sul rischio, prevedendo obblighi differenziati in relazione al possibile impatto di ciascun sistema sui diritti fondamentali, sulla sicurezza e sugli interessi degli utenti. Su questo impianto è successivamente intervenuto il **Regolamento (UE) 2026/1744** dell’8 luglio 2026, c.d. Omnibus digitale sull’IA, che ha apportato modifiche mirate all’AI Act per semplificarne l’attuazione e chiarire alcuni profili applicativi, mantenendo invariato il livello di tutela dei diritti fondamentali e della sicurezza.

A tale disciplina si affianca la **Direttiva (UE) 2024/2853 sulla responsabilità da prodotto difettoso**, che dovrà essere recepita dagli Stati membri entro il 9 dicembre 2026 e che include espressamente nella nozione di prodotto anche il *software*, compresi i sistemi di intelligenza artificiale, assoggettandoli al relativo regime europeo di responsabilità. In Italia, in attuazione della delega conferita dalla legge 17 marzo 2026, n. 36, il 4 agosto 2026 il Consiglio dei Ministri ha approvato, in esame preliminare, lo schema di decreto legislativo di recepimento della Direttiva, attualmente sottoposto all’esame parlamentare, destinato a modificare in modo organico la disciplina della responsabilità per danno da prodotti difettosi contenuta nel Codice del consumo, mediante la sostituzione degli artt. 114-127 e l’introduzione del nuovo art. 126-bis.

Tali strumenti operano su piani complementari: l’AI Act individua gli *standard* di progettazione e utilizzo dei sistemi di IA, mentre la Direttiva disciplina le conseguenze risarcitorie della loro eventuale difetto-

sità. Sebbene il Regolamento non disciplini direttamente la responsabilità civile, i suoi obblighi sono destinati a costituire un parametro di riferimento nella valutazione della diligenza dell'impresa e della conformità del sistema agli *standard* di sicurezza richiesti.

Il quadro europeo sarà inoltre completato dall'attuazione nazionale dell'AI Act. Il 4 agosto 2026 il Consiglio dei Ministri ha approvato, in esame definitivo, due decreti legislativi di adeguamento della normativa nazionale al Regolamento (UE) 2024/1689, in attuazione della delega di cui all'articolo 24 della legge 23 settembre 2025, n. 132. Il primo decreto disciplina l'utilizzo dei sistemi di intelligenza artificiale per l'attività di polizia e introduce disposizioni in materia di responsabilità civile e penale, inclusa l'acquisizione delle prove e la tutela processuale. Il secondo decreto definisce l'assetto di *governance* nazionale, attribuendo poteri di vigilanza e ispettivi alle autorità competenti (in particolare AgID, ACN, Banca d'Italia, CONSOB e IVASS) e regola l'impiego dell'intelligenza artificiale nei percorsi formativi, nelle professioni, nel lavoro, nella sanità e nella pubblica amministrazione.

Per le imprese, la conformità alla disciplina europea non costituisce dunque solo un adempimento regolatorio, ma rappresenta uno strumento di prevenzione del rischio legale. Documentare il funzionamento del sistema, la gestione dei rischi, la tracciabilità delle decisioni e il rispetto degli obblighi di trasparenza diventano elementi centrali della **litigation readiness**, destinati a un ruolo decisivo nell'eventuale contenzioso.

2. I destinatari degli obblighi

Uno degli aspetti più innovativi dell'AI Act consiste nel superamento di una concezione di responsabilità concentrata solo sul produttore del *software*. Il Regolamento distribuisce infatti gli obblighi lungo l'intera filiera dell'intelligenza artificiale, attribuendo specifiche responsabilità ai diversi operatori economici coinvolti.

Il principale destinatario è il **provider**, ossia il soggetto che sviluppa il sistema o lo immette sul mercato con il proprio nome o marchio. Su di esso gravano gli obblighi più rilevanti, tra cui predisporre un sistema di gestione dei rischi, garantire la qualità dei dati, redigere la documentazione tecnica, assicurare la registrazione automatica degli eventi (*logging*), fornire istruzioni d'uso adeguate e monitorare il sistema anche dopo la commercializzazione.

Accanto al *provider* vi è il **deployer**, vale a dire il soggetto che utilizza professionalmente il sistema. Poiché nella maggior parte dei casi le imprese acquistano o danno in licenza sistemi di terzi integrandoli nei propri processi, l'AI Act chiarisce che ciò non limita il rischio sul fornitore in quanto anche il *deployer* deve utilizzare il sistema secondo le istruzioni ricevute, garantire un'adeguata supervisione umana, conservare i registri quando previsto e assicurare che il personale disponga delle competenze necessarie per un utilizzo corretto dello strumento. L'impresa che utilizza un sistema di terzi può quindi assumere obblighi rilevanti, specie se lo integra nel proprio prodotto, ne modifica la destinazione d'uso o lo presenta con il proprio marchio.

Il Regolamento (UE) 2026/1744 ha rafforzato tale impianto intervenendo sull'art. 25 dell'AI Act. Il fornitore iniziale che, a seguito della modifica sostanziale del sistema da parte di un terzo, non sia più qualificato come tale, deve cooperare con il nuovo fornitore, mettendo a disposizione la documentazione, l'accesso tecnico e l'assistenza necessari all'adempimento degli obblighi regolatori, inclusa la valutazione di conformità. Il nuovo par. 4 richiede inoltre che i rapporti tra il fornitore di un sistema ad alto rischio e i terzi che forniscono componenti, modelli o servizi integrati siano disciplinati da un accordo scritto.

Tale distribuzione delle responsabilità incide direttamente sui possibili scenari di contenzioso, rendendo probabili azioni promosse nei confronti di più operatori della filiera. Diventa quindi essenziale disciplinare contrattualmente la ripartizione delle responsabilità tramite clausole di manleva, obblighi di cooperazione nella gestione delle richieste risarcitorie e procedure condivise per la conservazione della documentazione tecnica.

3. La classificazione del rischio come primo parametro di responsabilità

L'approccio *risk-based* costituisce il tratto distintivo dell'AI Act, poiché il legislatore europeo calibra gli obblighi imposti ai diversi operatori in funzione del livello di rischio che ciascun sistema di IA può comportare per i diritti fondamentali, la sicurezza e gli interessi economici delle persone.

Per le imprese, la classificazione del sistema è il primo passaggio della *compliance* e, al tempo stesso, il primo profilo scrutinato in un eventuale giudizio: una classificazione errata o incompleta può costituire non solo una violazione regolatoria, ma anche l'indice della mancata adozione delle misure organizza-

tive richieste dalla normativa.

L'AI Act distingue quattro categorie di rischio.

La prima comprende i **sistemi a rischio inaccettabile**, vietati perché incompatibili con i valori fondamentali dell'Unione europea, come talune pratiche di manipolazione cognitiva, di sfruttamento delle vulnerabilità delle persone e di *social scoring*.

La seconda, di maggior rilievo pratico, riguarda i **sistemi ad alto rischio**: pur lecite, tali applicazioni possono incidere significativamente sui diritti e sulle opportunità delle persone e sono quindi assoggettate a un articolato sistema di obblighi preventivi.

L'Allegato III individua gli ambiti in cui tali sistemi sono più diffusi: selezione e gestione del personale, istruzione e formazione, infrastrutture critiche, servizi pubblici essenziali, sanità, giustizia, contrasto dei reati, nonché alcuni servizi bancari e assicurativi. Vi ricadono, ad esempio, i *software* per lo *screening* automatizzato dei *curricula* o per la valutazione dell'affidabilità economica dei clienti.

Il Regolamento (UE) 2026/1744 ha rinviato l'applicazione degli obblighi per i sistemi di IA ad alto rischio, prevedendo due scadenze: il 2 dicembre 2027 per la maggior parte dei sistemi ad alto rischio e il 2 agosto 2028 per quelli integrati in prodotti già soggetti alla normativa europea di sicurezza.

Al di fuori di queste due categorie di rischio, il Regolamento prevede obblighi più limitati per i sistemi soggetti a specifiche prescrizioni di trasparenza e non impone particolari adempimenti per quelli a rischio minimo, rispetto ai quali si limita a incoraggiare l'adozione volontaria di codici di condotta.

In tale prospettiva, quanto maggiore risulta il rischio attribuito dal legislatore, tanto più elevata la diligenza richiesta nella progettazione, implementazione e gestione del sistema. È prevedibile che, nelle future controversie, la natura "ad alto rischio" del sistema sia tra i primi elementi utilizzati dal giudice per valutare la condotta dell'operatore economico.

4. Gli obblighi dell'AI Act come parametro della diligenza dell'impresa

Per i sistemi classificati ad alto rischio, l'AI Act introduce un articolato insieme di obblighi che, pur non disciplinando direttamente la responsabilità civile, sono destinati a costituire il parametro per valutare la diligenza dell'impresa nell'eventuale giudizio risarcitorio.

Il primo nucleo riguarda la **governance del rischio**: il *provider* deve adottare un sistema di gestione dei rischi lungo l'intero ciclo di vita del sistema, individuando e mitigando i rischi prevedibili per salute, sicurezza e diritti fondamentali. Vi si affiancano gli obblighi sulla qualità dei dati di addestramento, validazione e *test*, che devono essere adeguati, rappresentativi e possibilmente esenti da errori, poiché *dataset* incompleti o distorti possono generare *output* inesatti o discriminatori.

Un secondo gruppo riguarda la **tracciabilità** del sistema: l'AI Act impone documentazione tecnica, conservazione dei *log* e registrazione degli eventi rilevanti, per consentire la ricostruzione *ex post* del funzionamento del sistema e delle decisioni assunte.

Centrale è la **supervisione umana**: coerentemente con l'approccio antropocentrico del Regolamento, il controllo umano non può limitarsi a una verifica formale, ma deve consentire di comprendere, monitorare e correggere l'*output* del sistema, evitando l'eccessivo affidamento sulla decisione algoritmica. La centralità di tale presidio è stata ulteriormente rafforzata, sul piano nazionale, dal decreto legislativo approvato in via definitiva dal Consiglio dei Ministri il 4 agosto 2026, che prevede una specifica disciplina per l'utilizzatore professionale di sistemi di IA ad alto rischio che ometta intenzionalmente di adottare le prescritte misure di sorveglianza umana.

Rilevano anche gli obblighi di **trasparenza e spiegabilità**: l'art. 50 impone di informare gli utenti quando interagiscono con un sistema di IA, mentre l'art. 86 riconosce agli interessati il diritto a una spiegazione chiara del ruolo svolto dal sistema nella decisione, senza necessità di divulgare il codice sorgente o i segreti industriali, ma con un livello di comprensibilità sufficiente a esercitare efficacemente i propri diritti.

Nella stessa prospettiva si colloca l'obbligo di **AI literacy** previsto dall'art. 4 dell'AI Act – come riformulato dal Regolamento (UE) 2026/1744 – che impone a *provider* e *deployer* di adottare misure volte

a sostenere lo sviluppo dell'alfabetizzazione in materia di IA del proprio personale e di qualsiasi altra persona che si occupi del funzionamento e dell'utilizzo dei sistemi per loro conto, tenuto conto del contesto operativo e dei rischi connessi. La nuova formulazione precisa che tale obbligo non richiede il raggiungimento di un livello specifico di alfabetizzazione, ma impone l'adozione di misure concrete di supporto, con l'intervento di sostegno della Commissione e degli Stati membri.

Nel complesso, tali obblighi delineano un'infrastruttura di *governance* dell'IA. La predisposizione e conservazione della documentazione tecnica, la valutazione dei rischi, la supervisione, l'informazione agli utenti e la formazione del personale diventeranno gli strumenti con cui l'impresa potrà dimostrare di avere adottato le misure ragionevolmente esigibili per prevenire il danno.

5. Profili processuali: *disclosure*, presunzioni e valore della documentazione

L'evoluzione della disciplina europea non ridefinisce soltanto gli obblighi di *compliance*, ma introduce innovazioni anche sul piano processuale. Sia la Direttiva (UE) 2024/2853, sia lo schema di decreto legislativo italiano rafforzano gli strumenti probatori a disposizione del danneggiato, attribuendo un ruolo centrale alla documentazione tecnica dell'impresa. In tale contesto si inserisce anche lo schema preliminare di decreto legislativo di recepimento della Direttiva, che interviene specificamente sul piano probatorio, disciplinando, attraverso i nuovi artt. 119 e 120 del Codice del consumo, l'esibizione degli elementi di prova e il sistema delle presunzioni.

Il primo profilo riguarda l'**accesso alle prove**. L'art. 9 della Direttiva attribuisce al giudice il potere di ordinare la divulgazione delle prove nella disponibilità del convenuto quando l'attore abbia fornito elementi sufficienti a rendere plausibile la domanda. Tale meccanismo è recepito, nello schema di decreto legislativo di attuazione, dal nuovo art. 119 del Codice del consumo, che disciplina l'esibizione degli elementi di prova pertinenti. Analogamente, l'art. 17 dello schema di decreto legislativo di adeguamento della normativa nazionale all'AI Act consente al giudice di ordinare l'esibizione della documentazione sul funzionamento del sistema di IA, ove siano allegati fatti idonei a rendere verosimile la domanda e il collegamento tra *output* e danno lamentato.

Diverso è invece il meccanismo previsto in caso di mancata esibizione. Da un lato, la Direttiva introduce una **presunzione di difettosità del prodotto**, dall'altro lo schema italiano richiama l'art. 116 c.p.c., con-

sentendo al giudice di desumere **argomenti di prova** dal comportamento della parte e, in caso di mancata produzione di documenti essenziali (*log*, documentazione tecnica, sistemi di gestione dei rischi, informazioni sulla supervisione umana), di ritenere provati i fatti allegati dall'attore.

Ulteriore elemento di convergenza è rappresentato dal sistema delle **presunzioni**. La Direttiva prevede presunzioni relative alla difettosità del prodotto e al nesso causale, in particolare nei casi di violazione degli obblighi di sicurezza, mancata *disclosure* o difficoltà probatoria derivante dalla complessità tecnica del sistema. Tali previsioni trovano attuazione nel nuovo art. 120 del Codice del consumo, che disciplina le presunzioni relative alla difettosità del prodotto e al nesso causale. Parallelamente, l'art. 18 dello schema di decreto legislativo di adeguamento della normativa nazionale all'AI Act introduce una presunzione del nesso causale tra la violazione degli obblighi dell'AI Act e il danno lamentato, salvo prova contraria.

Ne deriva che registri, *log*, sistemi di gestione del rischio, documentazione tecnica, istruzioni d'uso e informazioni sulla supervisione umana non sono più meri adempimenti regolatori, ma il principale presidio difensivo dell'impresa nell'eventuale giudizio risarcitorio.

6. I nuovi rischi dell'intelligenza artificiale per banche e intermediari finanziari

La diffusione dell'intelligenza artificiale nel settore bancario e finanziario – dal *credit scoring* all'AML/CFT, fino al *trading algoritmico* – non si limita a trasformare i processi operativi, ma genera nuove aree di responsabilità che impongono un ripensamento della *governance* interna degli intermediari.

Sul piano regolatorio, l'AI Act classifica come "alto rischio" i sistemi di *credit scoring* e di valutazione del merito creditizio delle persone fisiche, con esclusione dei sistemi antifrode¹, imponendo agli intermediari l'adozione anticipata delle misure di gestione del rischio, tracciabilità e supervisione umana già descritte.

Il dato empirico conferma la portata del fenomeno: secondo l'indagine OCSE condotta con la Banca d'I-

¹ Art. 6 e Allegato III, punto 5, lett. b), AI Act; cfr. anche Considerando 58.

talia², il 39% degli intervistati utilizza l'IA nell'operatività quotidiana, con il settore assicurativo in testa (70%), seguito da quello bancario (59%). I casi d'uso più diffusi riguardano l'ottimizzazione dei processi interni, l'analisi dei dati, la generazione di contenuti testuali, l'AML/CFT, l'antifrode e i *chatbot* per la clientela, sebbene la maggior parte delle applicazioni di IA generativa resti sperimentale.

Un primo profilo di rischio emerge dalla dipendenza da fornitori terzi: il 75% degli intervistati utilizza servizi *cloud* di terze parti e il 39% modelli di *General-Purpose AI* di terzi, mentre una quota equivalente evita soluzioni *open source* per timori di sicurezza, concentrando il rischio operativo su pochi fornitori.

Un secondo profilo riguarda la *governance* dell'IA, ancora eterogenea e immatura. Solo il 16% degli intermediari ha adottato assetti di *governance* specifici, la metà si affida alla sola supervisione umana (*human-in-the-loop*), e quasi la metà non ha ancora adottato misure di protezione contro le minacce informatiche specifiche dell'IA.

Sul piano normativo, gli intermediari segnalano soprattutto incertezza regolamentare e dubbi sull'attuazione dell'AI Act e sulla sua interazione con la normativa esistente, oltre a preoccupazioni in materia di protezione dei dati, proprietà intellettuale, gestione dei rischi di terze parti e resilienza operativa.

Emergono inoltre barriere non regolamentari, tra cui le difficoltà nel reperire competenze specialistiche, la limitata comprensione dell'IA da parte del *top management*, la carenza di casi d'uso consolidati, i problemi di qualità e accesso ai dati, i costi elevati, i rischi operativi e la scarsa trasparenza dei modelli esterni, con conseguente rischio di responsabilità legali o danni alla clientela.

A fronte di tali rischi, le autorità di vigilanza europee e nazionali sono intervenute, come si vedrà di seguito, con indicazioni volte a favorire uno sviluppo sicuro e responsabile dell'IA nei servizi bancari e finanziari, richiedendo agli intermediari il rafforzamento tempestivo dei presidi di *governance*.

L'AI Act, pur non sostituendosi alla normativa finanziaria di settore, si intreccia con un quadro di regole preesistenti in materia di assetti organizzativi, controlli interni, esternalizzazione, tutela della clientela,

trasparenza, adeguatezza e resilienza operativa, con la conseguente esigenza di coordinamento tra i vari corpi normativi.

Si tratta, in altri termini, di un'innovazione non più sperimentale, ma strutturale, che incide su processi decisionali con rilevanza giuridica diretta sulla clientela e che deve essere coordinata, *inter alia*, con gli obblighi di sicurezza informatica già previsti da DORA, CRA e NIS 2³, nell'ambito di una *governance* integrata dei rischi tecnologici.

7. Il nuovo quadro di aspettative in materia di *governance* e cybersicurezza dell'IA

In tale contesto, un insieme coordinato di interventi delle autorità di vigilanza e di stabilità finanziaria – la BCE, nell'ambito del Meccanismo di Vigilanza Unico (SSM), e l'ESRB – ha delineato un quadro di aspettative più stringente per banche e intermediari finanziari sui rischi connessi all'intelligenza artificiale, con enfasi sulle minacce *cyber* derivanti dai *frontier AI models* (FAIM), i modelli avanzati di IA a uso generico capaci di influenzare le operazioni informatiche offensive o difensive.

Il Warning ESRB/2026/3 del 25 giugno 2026 e il successivo rapporto tecnico di luglio 2026 documentano un salto qualitativo nelle capacità offensive dei FAIM, ormai in grado di condurre *cyber* attacchi pienamente automatizzati, comprimendo il tempo tra scoperta e sfruttamento delle vulnerabilità. Ciò indebolisce la resilienza operativa del sistema finanziario lungo quattro direttrici: il tempo per il *patching*; la capacità dei sistemi di difesa di tenere il passo con le minacce; la concentrazione su pochi fornitori *AI/cloud*; la capacità delle autorità di coordinare la risposta a incidenti transfrontalieri.

Dinanzi a tali minacce, l'ESRB raccomanda agli intermediari la revisione e l'aggiornamento del *framework* interno di sicurezza informatica, affinché tenga conto delle vulnerabilità evidenziate e delle nuove capacità offensive rese possibili dall'IA.

Con la lettera SSM-2026-0301 del 7 luglio 2026, indirizzata ai CEO delle *significant institutions*, la BCE ha tradotto tali preoccupazioni in un obbligo concreto, evidenziando che i modelli di IA emergenti in-

² OECD (2026), L'intelligenza artificiale nei mercati finanziari italiani, OECD Publishing, Paris, <https://doi.org/10.1787/7ecf1246-it>.

³ Reg. (UE) 2022/2554 (*Digital Operational Resilience Act* – DORA), il Reg. (UE) 2024/2847 (*Cyber Resilience Act* – CRA) e la Dir. (UE) 2022/2555 (*Network and Information Security 2* – NIS 2).

dividuano vulnerabilità e generano *exploit*⁴ a velocità senza precedenti, e attribuendo la responsabilità primaria della reazione agli organi di gestione delle banche.

Le *significant institutions* devono pertanto sviluppare “senza indugio” un piano d’azione, da presentare al Joint Supervisory Team (JST) entro il 31 ottobre 2026, articolato su misure a breve e a lungo termine.

Quanto alle misure a breve termine, la BCE richiede di: dare priorità alla protezione delle aree vulnerabili agli attacchi; accelerare la gestione delle vulnerabilità e le attività di *patching*; potenziare le capacità di monitoraggio, rilevamento e difesa basate sull’IA; rafforzare *governance*, finanziamenti e formazione dedicati alla prevenzione, nonché la qualità della catena di approvvigionamento.

Quanto alle azioni a medio-lungo termine, l’Autorità suggerisce il rafforzamento della difesa a più livelli e delle buone pratiche di sicurezza informatica, la modernizzazione delle infrastrutture e il miglioramento della resilienza operativa, inclusa la gestione delle crisi e i meccanismi di condivisione delle informazioni tra intermediari e autorità.

8. Le buone pratiche di *governance* per gli enti finanziari

Sulla stessa linea, il Financial Stability Board (FSB) ha posto in consultazione, il 10 giugno 2026, una serie di *sound practices* per la gestione responsabile dei rischi connessi all’IA, completando il quadro di aspettative già delineato da BCE ed ESRB.

Le dodici buone pratiche, da applicarsi nel rispetto del principio di proporzionalità, riguardano la *governance* a livello organizzativo e la gestione delle diverse fasi del ciclo di vita dell’IA.

Il Consiglio di amministrazione e il *senior management* sono incoraggiati a fare riferimento a tali pratiche nella definizione della strategia aziendale, nell’adozione delle tecnologie e nella gestione dei rischi, in un contesto operativo sempre più permeato dall’IA.

⁴ Per tale deve intendersi uno strumento, una tecnica o un metodo utilizzato per sfruttare una o più vulnerabilità al fine di ottenere accesso non autorizzato ai sistemi, modificarli o comprometterne il funzionamento, oppure causare in altro modo rischi informatici o incidenti legati alle tecnologie dell’informazione e della comunicazione.

Il primo gruppo di pratiche definisce le fondamenta di *governance* da porre in essere prima dei singoli casi d’uso: l’adozione dell’IA non può discendere da iniziative isolate, ma da una chiara direzione strategica, con il *board* chiamato ad allinearne l’uso al modello di *business* e alla propensione al rischio, definendo aspettative, integrando i rischi nella *risk tolerance* e assicurando risorse adeguate in termini di competenze, tecnologia e infrastrutture.

Su tale base si innesta un impianto di *governance* con ruoli precisi, tipicamente secondo il modello delle tre linee di difesa, e meccanismi di coordinamento tra le funzioni coinvolte. I rischi IA vanno integrati nel *framework* di *risk management* esistente, con *policy* e controlli calibrati su materialità e rischio, e documentazione lungo l’intero ciclo di vita. È richiesta inoltre un’adattabilità organizzativa continua, tramite formazione, monitoraggio degli sviluppi esterni e revisione periodica dei controlli.

Il secondo blocco di *sound practices* entra nel merito operativo, seguendo idealmente le sette fasi del ciclo di vita di un sistema IA (*inception*, progettazione e sviluppo, verifica e validazione, *deployment*, operatività e monitoraggio, rivalutazione, ritiro).

Il punto di partenza è la valutazione di materialità e rischio, condotta all’avvio e in modo continuativo, che orienta la selezione del modello IA in base a obiettivi di business, esigenze operative e *trade-off* tra *performance* e spiegabilità.

A monte e a valle della selezione si colloca la *data governance*, per garantire dati accurati, completi e sicuri. Un tema trasversale è la spiegabilità e trasparenza: le istituzioni devono comprendere le differenze tra le diverse tipologie di IA, privilegiando soluzioni più spiegabili o controlli compensativi, e modulando il livello di trasparenza in base allo *stakeholder* di riferimento.

Segue la gestione della *performance*, proporzionata al rischio tramite *test*, monitoraggio continuo, *benchmarking*, *red-teaming* e supervisione umana calibrata su materialità, rischio, autonomia e complessità (*human-in-the-loop*, *AI-in-the-loop*, *human-on-the-loop*, *human-in-command*, *kill switch*, *contestability*), con accorgimenti ulteriori per *GenAI* e IA agentica, come identificatori individuali per gli agenti IA, confini operativi netti e monitoraggio dei passaggi intermedi decisionali.

Le ultime due pratiche riguardano i rischi esterni: sul fronte *cyber* e ICT si richiede di adattare le pra-

tiche di sicurezza ai rischi specifici dell'IA, sfruttare l'IA per la gestione del rischio *cyber* e includere scenari *cyber-IA* nei test, in linea con ESRB e BCE. Sul fronte dei rischi di terze parti, l'attenzione si concentra su *performance*, trasparenza, qualità dei dati, concentrazione della *supply chain* e continuità dei fornitori, da presidiare con *due diligence*, clausole contrattuali adeguate, strumenti di trasparenza (*model card*, certificazioni come ISO/IEC 42001), controlli compensativi e gestione attiva della concentrazione verso singoli provider.

Merita, infine, richiamare la lettera del Presidente del FSB, Andrew Bailey, del 28 agosto 2026, indirizzata ai Ministri delle Finanze e ai Governatori delle Banche centrali del G20, che segnala il permanere di vulnerabilità nel sistema finanziario globale, tra cui la sopravvalutazione degli asset legati all'intelligenza artificiale, l'aumento della leva finanziaria nei mercati azionari e la crescente concentrazione degli investimenti incrociati tra società di IA e *hyperscaler*. In particolare, Bailey osserva che i FAIM potrebbero avere la capacità di alterare in modo sostanziale la velocità, la portata e gli aspetti economici del rischio informatico, il che potrebbe minare la fiducia del mercato a livello sistemico, specialmente a causa dell'elevata concentrazione dei fornitori di servizi di terze parti. Ne discende la necessità che istituzioni finanziarie, infrastrutture di mercato e fornitori tecnologici rafforzino la gestione delle vulnerabilità, le capacità di risposta e ripristino e la preparazione a scenari di *disruption* simultanea in più operatori o dipendenze tecnologiche condivise. La lettera richiama, da ultimo, la necessità di protocolli su scala globale per il rilascio e l'impiego sicuri e responsabili dei FAIM, rilevando che molte giurisdizioni non dispongono ancora di un quadro adeguato.

9. La Comunicazione delle ESAs sui rischi ICT dei FAIM

Il 31 luglio 2026 le Autorità europee di vigilanza (EBA, EIOPA ed ESMA) hanno pubblicato la comunicazione JC 2026 25, intitolata "*Toward a consistent and risk-based approach for ICT risks from frontier AI models*", con l'obiettivo di promuovere un approccio di vigilanza coordinato e basato sui rischi derivanti dai FAIM. La Comunicazione rileva che le capacità avanzate dei più recenti modelli accelerano significativamente i rischi cibernetici, consentendo di individuare e sfruttare rapidamente le vulnerabilità, di colpire le debolezze delle infrastrutture condivise e di fare leva sui punti singoli di malfunzionamento o di compromissione (*single points of failure*). Il documento si pone in continuità con gli avvertimenti

dell'ESRB, le raccomandazioni dell'ENISA⁵, la lettera della BCE ai CEO e il Piano d'azione della Commissione europea sulla cybersicurezza e l'intelligenza artificiale del 7 luglio 2026⁶.

Le ESAs confermano che l'attuale quadro regolamentare dell'Unione – in particolare il Regolamento DORA e l'AI Act – offre una base solida per fronteggiare i rischi connessi al rilascio di modelli di IA altamente capaci. DORA mantiene piena rilevanza attraverso il sistema di gestione dei rischi ICT, i test di resilienza, la gestione degli incidenti e la gestione dei rischi derivanti dai fornitori terzi ICT; l'AI Act, a sua volta, prevede obblighi aggiuntivi per i fornitori di modelli di IA per finalità generali con rischio sistemico, tra cui trasparenza, documentazione tecnica e cybersicurezza. Tuttavia, la riduzione dei tempi di scoperta e sfruttamento delle vulnerabilità impone agli intermediari di agire rapidamente e in via proattiva, adottando misure proporzionate alla dimensione, al profilo di rischio e alla natura, scala e complessità dei servizi, delle attività e delle operazioni svolte.

La Comunicazione individua tre strategie di mitigazione del rischio. La prima è la prevenzione, da perseguire mediante inventari completi e costantemente aggiornati degli asset IT, principi di *secure-by-design*, riduzione della superficie di attacco, gestione rigorosa degli accessi, applicazione proattiva delle *patch* e standard di cybersicurezza estesi all'intera catena di fornitura. La seconda è la rilevazione, attraverso scansioni continue delle vulnerabilità e aumento della frequenza di test e verifiche di conformità. La terza è la gestione, che richiede l'aggiornamento dei presidi di risposta agli incidenti e di continuità operativa per le minacce assistite dall'IA, test periodici di resilienza operativa con scenari potenziati dall'IA, l'evoluzione della responsabilità degli organi di gestione verso un coinvolgimento continuo e una consapevolezza in materia di cybersicurezza più reattiva e proattiva.

Le ESAs sottolineano che gli organi di gestione degli intermediari devono essere pienamente impegnati nella mitigazione dei rischi cibernetici alimentati dai FAIM, mediante chiari assetti di *governance* e responsabilità, piani di risposta tempestivi e investimenti interni sufficienti. Il Risk Appetite Framework deve essere aggiornato per incorporare metriche, soglie di tolleranza e misure di controllo coerenti con il profilo di rischio in evoluzione, anche in relazione all'uso interno dei modelli e all'esposizione indiretta

⁵ <https://www.enisa.europa.eu/publications/enisas-view-on-cybersecurity-in-the-frontier-ai-era>

⁶ <https://digital-strategy.ec.europa.eu/en/library/eu-action-plan-cybersecurity-and-artificial-intelligence>

agli stessi. Tali strategie devono essere attuate in modo proporzionato, in conformità all'art. 4 del Regolamento DORA.

Le ESAs, in qualità di *Lead Overseers*, hanno inoltre avviato interlocuzioni mirate con i fornitori terzi critici di servizi ICT (CTPPs) per comprenderne la gestione delle nuove sfide e hanno iniziato a integrare i rischi connessi all'IA nella propria *Oversight Examination Methodology* per il 2027. La Comunicazione completa così il quadro di vigilanza emergente sui rischi dell'IA nel settore finanziario, insieme agli interventi dell'ESRB, della BCE e del FSB già esaminati, coerentemente con le iniziative di vigilanza nazionali – come le Comunicazioni della Banca d'Italia e di IVASS illustrate nel capitolo successivo – che traducono tali aspettative europee in requisiti operativi specifici per gli intermediari domestici.

10. La Comunicazione della Banca d'Italia in materia di resilienza operativa digitale e modelli avanzati di IA

Sul piano interno, la Banca d'Italia è intervenuta con una Comunicazione al mercato dedicata alla resilienza operativa digitale e ai modelli avanzati di IA, indirizzata a banche e gruppi bancari meno significativi, istituti di pagamento e moneta elettronica, imprese di investimento, fornitori di servizi crypto, gestori di FIA, SGR, fornitori di *crowdfunding* e intermediari finanziari ex art. 106 TUB. La Comunicazione rileva che i modelli di IA di ultima generazione individuano vulnerabilità e generano modalità di sfruttamento in tempi estremamente ridotti e senza competenze tecniche specifiche, accrescendo l'asimmetria tra attaccanti e difensori.

In coerenza con la Vigilanza unica europea e con i requisiti del DORA, la Banca d'Italia invita gli intermediari a rafforzare tempestivamente i presidi in sei aree: *governance*, con revisione del RAF per incorporare i rischi delle tecnologie di frontiera e adeguare le competenze tecnologiche in seno agli organi di amministrazione; igiene informatica, tramite *defence-in-depth*, autenticazione rigorosa e segmentazione delle reti; gestione delle risorse informatiche, con inventari aggiornati e modernizzazione delle infrastrutture; gestione delle vulnerabilità e delle *patch*, anche con il supporto di strumenti di IA; monitoraggio e difesa integrati con strumenti di IA; test di resilienza operativa digitale, coerenti con i requisiti previsti dal Regolamento DORA.

La Comunicazione richiede che il Consiglio di Amministrazione, in seduta congiunta con il Collegio Sin-

dacale, esamini il contenuto, avvii la redazione di una Relazione che rappresenti, per ciascuna area, il livello di esposizione al rischio, l'adeguatezza dei presidi esistenti, le principali carenze e definisca un piano di lavoro proporzionato al profilo di rischio, con azioni, tempi e investimenti; ogni intermediario deve individuare e comunicare all'Organo di Vigilanza uno specifico punto di responsabilità aziendale. La Relazione, con il piano di lavoro allegato, va trasmessa alla Vigilanza entro il 31 dicembre 2026.

Nel complesso, gli interventi di ESRB, BCE, FSB e Banca d'Italia confermano che, per banche e intermediari, la sfida dell'intelligenza artificiale non è più solo tecnologica ma organizzativa: i nuovi rischi – dalla concentrazione su pochi fornitori alle minacce *cyber* dei *frontier AI models*, sino alle lacune di *governance* ancora diffuse – richiedono risposte tempestive da parte degli organi di gestione. Rafforzare gli assetti di *governance*, integrare i rischi IA nel *risk management* esistente e presidiare la catena di fornitura tecnologica rappresentano quindi non solo un adempimento regolatorio, ma una condizione essenziale di resilienza e tutela della clientela, oltre che uno strumento fondamentale per prevenire e mitigare il rischio di contenzioso.

In data 17 luglio 2026, l'IVASS ha diffuso una comunicazione analoga indirizzata alle imprese di assicurazione e di riassicurazione con sede legale in Italia nonché alle rappresentanze generali per l'Italia delle imprese con sede legale in uno Stato terzo rispetto allo S.E.E.

11. Le principali aree di rischio per le imprese

Le prime esperienze applicative statunitensi consentono già di individuare alcune direttrici lungo le quali potrebbe essere destinato a svilupparsi il futuro contenzioso in materia di intelligenza artificiale.

Una prima area riguarda la **responsabilità precontrattuale**. La mancata informazione sull'utilizzo di sistemi di IA, soprattutto quando incidano sulla decisione o sulle condizioni offerte alla controparte, può violare il principio di buona fede nelle trattative. Si pensi a una banca che valuti il merito creditizio con un sistema di IA senza informare adeguatamente il cliente dei suoi limiti, esponendosi, in caso di erroneo diniego del credito, a responsabilità ai sensi dell'art. 1337 c.c. Per converso, la concessione di un finanziamento fondata su una valutazione erronea o negligente del merito creditizio effettuata mediante il sistema algoritmico potrebbe esporre l'istituto bancario a responsabilità per concessione abusiva del credito. Inoltre, l'utilizzo di sistemi di IA per la valutazione del merito creditizio espone

l'intermediario al rischio di contestazioni fondate sulla discriminazione algoritmica, qualora il sistema determini un trattamento deteriore per specifiche categorie di clienti sulla base di correlazioni statistiche non trasparenti.

Di rilievo in tal senso è la vicenda statunitense della Earnest Operations LLC, società di prestiti agli studenti con sede nel Delaware, risolta con un accordo transattivo da 2,5 milioni di dollari con l'ufficio del Procuratore Generale del Massachusetts, alla quale erano contestate, tra l'altro, la mancata adozione di misure ragionevoli per mitigare i rischi di *fair lending*, la mancata validazione dei modelli rispetto al rischio di *disparate impact* e l'impiego di variabili, quale il "Cohort Default Rate", suscettibili di penalizzare in modo sproporzionato i richiedenti di colore e ispanici in termini di approvazioni e condizioni di prestito.

Una seconda area concerne la **responsabilità contrattuale**. Quando l'impresa utilizza un sistema di IA nell'esecuzione della prestazione, l'eventuale errore del sistema ricade normalmente nella sua sfera organizzativa. Si pensi a una banca o un intermediario finanziario che utilizzi *chatbot* per informare la clientela su prodotti finanziari, fornendo informazioni imprecise o raccomandazioni non coerenti con il profilo del cliente – l'intermediario potrebbe rispondere a titolo di responsabilità contrattuale, non potendo l'automazione escludere l'obbligo di garantire correttezza e adeguatezza delle informazioni rese. Criticità analoghe emergono con l'IA agentica applicata alla gestione degli investimenti: piattaforme di *robo-advisory* o sistemi di *trading* algoritmico dotati di crescente autonomia decisionale potrebbero eseguire operazioni non coerenti con il profilo di rischio del cliente o con le istruzioni ricevute. L'impresa dovrà quindi dimostrare di avere utilizzato il sistema con la diligenza richiesta, rispettando le istruzioni d'uso, verificando gli *output* e garantendo un controllo umano effettivo.

Rappresentativa in tal senso è la vicenda *Estate of Lokken et al. v. UnitedHealth Group*⁷, nella quale una compagnia assicurativa sanitaria è stata convenuta in giudizio per l'asserito utilizzo di un sistema di IA destinato a determinare il diniego o l'interruzione della copertura di cure post-acute in contrasto con le valutazioni dei medici curanti. La Corte ha ritenuto ammissibili le domande fondate sull'inadempimento

⁷ Cfr. *Estate of Gene B. Lokken et al. v. UnitedHealth Group, Inc. et al.*, 766 F. Supp. 3d 835, 848–849 (D. Minn. 2025), 2025 WL 491148, at *8, Civil No. 23-3514 (JRT/DJF), Memorandum Opinion and Order del 13 febbraio 2025.

mento contrattuale e sulla violazione del dovere di buona fede, evidenziando come il ricorso a sistemi algoritmici non esoneri l'impresa dall'obbligo di garantire un controllo umano effettivo sul processo decisionale.

Una terza area riguarda la **responsabilità extracontrattuale e da prodotto difettoso**. La qualificazione del *software* come prodotto operata dalla Direttiva (UE) 2024/2853 apre infatti la strada ad azioni risarcitorie fondate sulla difettosità dei sistemi di IA.

Il caso *Garcia v. Character Technologies*⁸ ne rappresenta un esempio significativo: la controversia, promossa a seguito del suicidio di un minore che aveva sviluppato una relazione di dipendenza emotiva con un *chatbot companion*, ha consentito la prosecuzione di domande fondate, tra l'altro, su *product liability*, *design defect*, *failure to warn*, *negligence* e *wrongful death*. Il caso dimostra come il rischio di responsabilità possa derivare non solo dall'*output* del sistema, ma anche dalle scelte di progettazione, dall'assenza di adeguati presidi di sicurezza, dalla mancata informazione sui rischi e dalla prevedibilità del danno nei confronti di utenti vulnerabili.

Un'ultima area concerne il **contenzioso collettivo** e la comunicazione commerciale delle capacità dell'IA. La diffusione di tali sistemi nei rapporti *business to consumer* rende sempre più plausibili azioni fondate su informative insufficienti, rappresentazioni commerciali fuorvianti o discrepanze tra le funzionalità promesse e quelle effettivamente disponibili.

In questa prospettiva si colloca il caso *Landsheft v. Apple*⁹, relativo alle funzionalità "Apple Intelligence", nel quale il contenzioso non nasceva da un malfunzionamento del sistema, bensì dalla contestata rappresentazione al mercato di funzionalità di IA non ancora disponibili al momento dell'acquisto. La vicenda evidenzia come la responsabilità possa riguardare anche la comunicazione commerciale e l'affidamento ingenerato nei consumatori.

⁸ Cfr. *Garcia v. Character Technologies, Inc.*, 785 F. Supp. 3d 1157, 1179–1182 (M.D. Fla. 2025), 2025 WL 1461721, Case No. 6:24-cv-1903-ACC-UAM, Order del 20 maggio 2025, depositato il 21 maggio 2025.

⁹ Cfr. *Landsheft v. Apple Inc.*, No. 5:25-cv-02668-NW (N.D. Cal.), *Second Consolidated Class Action Complaint*, ECF No. 76, 1° maggio 2026, parr. 458–464; già *Class Action Complaint*, ECF No. 1, 19 marzo 2025, parr. 2–8. Il procedimento è stato oggetto di un accordo transattivo approvato in via preliminare con *Order Granting Preliminary Approval of Class Action Settlement*, ECF No. 94, 17 luglio 2026.

Le prime esperienze applicative mostrano dunque come il contenzioso in materia di intelligenza artificiale si stia progressivamente spostando dall'errore algoritmico alla valutazione complessiva dell'organizzazione dell'impresa. Diventano centrali la qualità della *governance* adottata, la capacità di documentare il ruolo dell'algoritmo nel processo decisionale, l'effettività del controllo umano, la correttezza delle informazioni rese agli utenti e la coerenza tra le capacità promesse e quelle realmente offerte dal sistema.



DB non solo
diritto
bancario

A NEW DIGITAL EXPERIENCE

 **dirittobancario.it**
