

L'Intelligenza Artificiale: cos'è, tipi, rischi, opportunità e applicazioni bancarie

Un'immagine concettuale dell'Intelligenza Artificiale, spesso raffigurata come un cervello elettronico: in realtà l'AI odierna si basa su algoritmi matematico-statistici che elaborano enormi quantità di dati.

Cos'è l'AI?

L'**Intelligenza Artificiale** è la disciplina informatica che si occupa di creare sistemi e programmi capaci di svolgere compiti che richiederebbero l'intelligenza umana. In altre parole, l'AI consente alle macchine di **simulare capacità umane** come l'apprendimento, la pianificazione e una sorta di simulazione di processi di ragionamento e creatività. Ad esempio, un algoritmo di AI può analizzare dati, riconoscere modelli e prendere decisioni o fornire risposte basandosi su tali analisi, imitando in parte il modo in cui un essere umano risolverebbe un problema.

Per funzionare, un sistema di AI riceve **dati in input** (che possono essere forniti dall'uomo o raccolti tramite sensori, come fotocamere, microfoni, ecc.), li elabora mediante algoritmi, e produce un **output** che può essere una decisione, una previsione o un'azione. Un aspetto chiave dell'AI moderna è la **capacità di apprendere dall'esperienza**: molti sistemi di intelligenza artificiale adattano il proprio comportamento analizzando gli effetti delle azioni precedenti e migliorando le prestazioni nel tempo. Questo processo di apprendimento automatico ("**machine learning**") permette all'AI di migliorare continuamente, man mano che elabora nuovi dati.

Va sottolineato che l'intelligenza artificiale **non è una "magia"** o qualcosa di misterioso, anche se a volte viene presentata così. In realtà, come osservano molti esperti, ciò che chiamiamo AI non è altro che **statistica e calcolo su larga scala** applicati a grandi moli di dati. Ad esempio, Kate Crawford – studiosa del tema – ricorda che i sistemi odierni di AI sono essenzialmente **modelli matematici probabilistici**, capaci di produrre risultati elaborando probabilità, ma privi di coscienza o comprensione profonda. Demistificare l'AI è importante per comprenderla: quando un computer traduce una frase o riconosce il volto di una persona in una foto, non c'è sotto nessuna "magia" e, soprattutto, nessun pensiero, ma "solamente" un insieme di formule matematiche e dati di addestramento che guidano quella risposta.

Il termine "**Intelligenza Artificiale**" fu coniato già nel 1956 dallo scienziato John McCarthy. Da allora l'AI ha attraversato varie fasi di sviluppo: dagli albori con sistemi a regole fisse, fino agli **algoritmi avanzati di oggi** basati su reti neurali e *machine learning*. Oggi l'AI è diventata una **tecnologia pervasiva** nelle nostre vite quotidiane: è presente nei motori di ricerca che usiamo su Internet, nei suggerimenti di acquisto degli e-commerce, nei filtri anti-spam delle e-mail, nei sistemi di navigazione satellitare, nei social media che frequentiamo e molto altro. Probabilmente, senza rendercene conto, interagiamo con forme di AI ogni giorno – dal correttore automatico sullo smartphone, all'assistente vocale che risponde a una domanda, fino al feed di notizie personalizzato che leggiamo online.

Vari tipi di AI

Parlando di intelligenza artificiale, è importante distinguere **diversi tipi o categorie** di AI in base alle loro capacità e al loro campo d'azione. La classificazione principale è tra **AI “debole” (o ristretta)** e **AI “forte” (o generale)**.

- **Intelligenza Artificiale Debole (Narrow AI):** È l'AI di cui disponiamo oggi. Si tratta di sistemi progettati per svolgere **compiti specifici e “limitati”**. Queste AI eccellono in un ambito ristretto, ma non hanno coscienza né capacità di pensare in modo autonomo al di fuori delle loro funzioni. Esempi comuni di AI debole sono: gli algoritmi di raccomandazione di YouTube o Netflix (che suggeriscono contenuti in base ai gusti dell'utente), gli assistenti virtuali come Siri, Alexa o Google Assistant (che riconoscono comandi vocali e forniscono risposte), i programmi di navigazione GPS, i software di riconoscimento facciale o delle impronte digitali, e così via. Queste intelligenze artificiali **esistono già oggi in abbondanza** e supportano tantissime applicazioni: pensiamo ai sistemi che analizzano le immagini mediche per aiutare i dottori a identificare tumori, o agli algoritmi nelle automobili a guida autonoma che elaborano i dati dei sensori per “guidare” il veicolo. Ogni AI debole è **specializzata** in uno scopo, anche se le più recenti forme di intelligenza artificiale generativa, con la loro flessibilità e multimodalità, possono sembrare sempre più flessibili e orientate alla generalizzazione.
- **Intelligenza Artificiale Forte o Generale (General AI):** È una forma di AI **ipotetica e in fase di ricerca**, che punta a raggiungere un'intelligenza simile a quella umana, capace di apprendere e svolgere *qualsiasi* compito intellettuale. Viene anche chiamata **AGI** (dall'inglese *Artificial General Intelligence*) o **AI forte**. **Al momento non esiste** una vera AI generale: nessun computer oggi può vantare la versatilità e la coscienza di una mente umana. Tuttavia, scienziati e ricercatori stanno esplorando questa frontiera, e molti ritengono che in futuro potrebbe emergere un'AI in grado di pensare in modo autonomo e adattarsi a situazioni completamente nuove. Sarebbe un'AI in grado di passare da un compito all'altro (dalla matematica alla letteratura, dal cucinare al progettare un software) apprendendo e **migliorando da sé**, senza essere vincolata a un singolo dominio. Siamo ancora lontani da questo traguardo: come ha spiegato anche il divulgatore Ray Kurzweil, creare un'AI pari all'intelligenza umana richiede progressi ancora imprevedibili, e comporta anche tante **questioni etiche** – ma resta un tema affascinante su cui riflettere per il futuro.

Un'altra distinzione utile riguarda i **metodi e modelli** usati nell'AI odierna. La maggior parte dei sistemi di intelligenza artificiale attuali si basa sul *machine learning* (apprendimento automatico) e in particolare sul **deep learning**, che utilizza **reti neurali artificiali** (algoritmi vagamente ispirati alla struttura del cervello biologico). Esistono vari modi in cui un'AI può **imparare** a svolgere un compito:

- **Apprendimento supervisionato:** all'AI vengono forniti molti esempi (dati di input) associati alle risposte corrette (output desiderato). La macchina impara a generalizzare dagli esempi e a produrre risultati corretti su nuovi dati simili. Esempio: un algoritmo a cui vengono

mostrati migliaia di foto di gatti etichettate come “gatto” imparerà a riconoscere un gatto in una nuova foto (questa volta non etichettata). Molte applicazioni pratiche usano questo metodo, dai riconoscitori di immagini ai filtri antispam (addestrati su e-mail etichettate come spam o non spam).

- **Apprendimento non supervisionato:** all'AI si forniscono grandi quantità di dati non etichettati, e il sistema deve trovare autonomamente schemi o raggruppamenti nascosti in quei dati. È come dire alla macchina: “ecco tanti dati, scopri da solo come sono strutturati”. Questo approccio è utile, ad esempio, per segmentare clienti con abitudini simili nel marketing, o per scoprire anomalie nei dati. Pur essendo inizialmente più lento e difficile, questo metodo può portare l'AI a scoperte inattese, trovando correlazioni che sfuggono agli esseri umani.
- **Apprendimento tramite rinforzo:** il sistema di AI impara per tentativi ed errori, ricevendo ricompense o penalità in base alle sue azioni. Questo metodo è ispirato al modo in cui si può addestrare un animale: se compie l'azione giusta ottiene un premio, se sbaglia, invece, non lo consegue. Nel contesto dell'AI, un esempio classico sono i programmi che imparano a giocare ai videogiochi o a scacchi: all'inizio agiscono a caso, ma gradualmente capiscono quali mosse portano a vincere (ricompensa) e quali a perdere (penalità), ottimizzando quindi la strategia. L'apprendimento per rinforzo è stato fondamentale per traguardi come il software AlphaGo (che ha appreso da zero a giocare a Go a livello superiore al campione mondiale) e viene usato anche per ottimizzare processi complessi, dalla gestione del traffico nelle città, all'allocazione di risorse nei data center, fino alla definizione di strategie finanziarie ottimali (trading algoritmico, allocazione di portafoglio, ecc.).

Ognuna di queste tecniche può anche essere combinata e usata insieme (si parla di apprendimento **ibrido**) per costruire sistemi AI più robusti e flessibili. Ad esempio, un assistente virtuale potrebbe usare modelli supervisionati per riconoscere la voce, modelli non supervisionati per apprendere nuove modalità di linguaggio dai dati di internet, e apprendimento per rinforzo per migliorare le interazioni col tempo in base ai feedback degli utenti.

I modelli generativi e “generalisti” di AI

Negli ultimi anni avrete probabilmente sentito parlare spesso di **AI generativa**: sono quei modelli di intelligenza artificiale capaci di **creare nuovi contenuti** – testi, immagini, musica, video – sulla base di prompt o richieste dell'utente. Questi modelli vengono detti a volte anche “**modelli generalisti**” perché si presentano come strumenti versatili, utilizzabili in tanti contesti. Esempi famosi includono **ChatGPT** (modello linguistico che può sostenere conversazioni e produrre testi su svariati argomenti), **DALL-E** e **Midjourney** (che generano immagini da descrizioni testuali), nonché modelli per creare audio e video realistici (deepfake). Molti di questi sistemi sfruttano architetture di *reti neurali profonde* note come **Transformer** e sono addestrati su quantità enormi di dati.

Questi modelli hanno capacità sorprendenti: ChatGPT, ad esempio, può rispondere a domande su quasi ogni argomento, aiutare a scrivere saggi o codice di programmazione, riassumere

documenti e così via. Tuttavia, è fondamentale capire **come funzionano** e quali **limiti** hanno, soprattutto per utilizzarli con consapevolezza. I modelli linguistici come ChatGPT in realtà non “pensano” come un umano: lavorano prevedendo **quale parola (o sequenza di parole) è più probabile che segua** in un testo, sulla base di quello che hanno appreso da milioni di documenti. Hanno quindi costruito un gigantesco “modello statistico” della lingua, che permette loro di produrre frasi di senso compiuto e argomentazioni che *sembrano* logiche.

Il comportamento di questi modelli dipende fortemente dai **dati su cui sono stati addestrati**. Se la base di conoscenza è ampia e ben bilanciata, l’AI riuscirà a imitare risposte molto ricche e variegata, persino includendo riferimenti culturali o detti popolari. Ad esempio, un modello addestrato su un vasto corpus potrebbe sapere che alla domanda “Cosa succede se si rompe uno specchio?” molti umani risponderebbero “porta sette anni di sventura”, rifacendosi a una superstizione. Un modello addestrato su dati più limitati o tecnici forse direbbe solo “bisogna comprarne un altro”, ignorando l’aspetto folkloristico. Questo esempio mostra un punto cruciale: **più è completa e diversificata la base dati**, più l’AI può cogliere aspetti “di umanità” nelle risposte. Ma attenzione: saper riportare un modo di dire non significa che l’AI *capisca* cosa sia la sfortuna; sta semplicemente riproducendo un’associazione appresa.

Un **grosso problema** dei modelli generativi odierni è la loro **propensione all’errore e al bias**. Poiché apprendono dai dati esistenti (che possono contenere pregiudizi, imprecisioni o lacune), queste AI possono **ereditare distorsioni e preconcetti**. Ad esempio, se nei dati di addestramento ci fossero più testi prodotti da uomini che da donne, il modello potrebbe usare un linguaggio con bias di genere; oppure, come evidenziato da studi recenti, i grandi modelli linguistici tendono a riprodurre stereotipi: in un esperimento del 2023, dei ricercatori hanno chiesto a un’AI di generare lettere di referenze per candidati maschi e femmine, scoprendo che nel testo venivano descritti gli uomini con termini come “leader” o “eccezionale”, mentre per le donne comparivano aggettivi come “deliziosa” o “compassionevole”, suggerendo un bias implicito nelle descrizioni. Questo accade perché nei dati di addestramento (ad esempio, un ampio insieme di lettere di referenze reali) esistono certi bias sociali che la macchina assorbe passivamente. Analogamente, **dati distorti o incompleti** possono portare l’AI a fornire risposte tendenziose o errate: i modelli “assorbono i pregiudizi della società” se questi sono presenti nell’enorme quantità di dati con cui vengono addestrati. Un set di dati che nel tempo ha risentito di disuguaglianze sociali può finire per danneggiare gruppi storicamente emarginati quando l’AI viene usata per decisioni su assunzioni, concessione di crediti, valutazioni di rischio e altri ambiti delicati. In breve: **“spazzatura in ingresso, spazzatura in uscita”** – se i dati di partenza sono viziati, anche l’intelligenza artificiale produrrà risultati distorti.

Un altro limite importante da comprendere è la tendenza di questi modelli a *inventare informazioni* quando non sanno qualcosa, pur di fornire una risposta plausibile. Sono le cosiddette **“allucinazioni”** dell’AI. Immaginiamo di chiedere a ChatGPT una breve recensione di un romanzo poco conosciuto. Se il modello non ha realmente informazioni su quel libro, potrebbe comunque generare una recensione di sana pianta: scriverà un riassunto di trama (inventato), commenterà lo stile, e concluderà magari con “un romanzo consigliato a chi ama le storie di formazione” – il tutto con tono convincente, ma completamente falso! Questo non perché l’AI voglia deliberatamente mentire, ma perché il suo obiettivo primario è **generare testo coerente e credibile**. Se dire “non lo so” suonerebbe come una risposta fuori luogo alla richiesta dell’utente,

il modello tenderà a improvvisare piuttosto che ammettere la propria mancanza di conoscenza. Nei fatti, ChatGPT e simili sono *addestrati* anche sul feedback umano: spesso gli utenti premiano (con un giudizio positivo) una risposta che *sembra* corretta anziché una sincera ammissione di ignoranza. Col tempo, quindi, il modello “impara” che è preferibile riempire i vuoti con una risposta fittizia purché plausibile. Il risultato è che, su argomenti che conosce bene, l’AI può produrre testi eccellenti; ma su temi poco noti o domande specifiche, **tende a inventare con sicurezza**, mescolando informazioni vere e false in modo difficile da distinguere.

Questo significa che i modelli generativi **non sono fonti infallibili** di verità, e il loro output va sempre preso con cautela e **verificato** da un umano, soprattutto se usato in un contesto serio (dalla stesura di un articolo scolastico, all’elaborato universitario, fino a decisioni aziendali). Non a caso, gli stessi sviluppatori di queste AI stanno introducendo filtri e avvertenze. Ad esempio, oggi i chatbot evoluti evitano di affrontare *argomenti sensibili* (politica, religione, ecc.) perché si è visto che senza filtri potevano assumere toni offensivi o estremisti: il caso famoso è **Tay**, una primitiva AI di Microsoft lanciata su Twitter nel 2016, che in meno di 24 ore, interagendo liberamente con gli utenti, iniziò a pubblicare messaggi razzisti e aggressivi, al punto da costringere la società a disattivarla immediatamente. Questo episodio ha insegnato quanto sia importante incorporare limiti e controlli etici nei modelli. Oggi, dunque, le AI “chiacchierone” come ChatGPT hanno vari *guardrail*: rifiutano insulti, incitamenti all’odio, richieste illegali, e tendono a dare risposte neutre su temi controversi. Ciò le rende più sicure, ma anche **meno “creative” in certi contesti rispetto a un umano**, specie su questioni di opinione o di sensibilità. Insomma, i modelli generalisti odierni sono strumenti potentissimi, ma vanno conosciuti nei loro **pregi e difetti**. È anzi un ottimo esercizio, per noi utenti, “mettere alla prova” queste AI con domande di cui già sappiamo la risposta o con quesiti complessi, per capire dove eccellono e dove falliscono. Una volta compresi i limiti (allucinazioni, bias, incapacità di vero ragionamento astratto), possiamo usare queste tecnologie in modo più consapevole, senza farci ingannare dal loro linguaggio scorrevole che può dare l’illusione di un’intelligenza *quasi umana*.

Quali rischi e opportunità nell’utilizzo dell’AI?

Come ogni innovazione tecnologica, l’intelligenza artificiale porta con sé **grandi opportunità** ma anche **rischi e sfide**. È fondamentale che i giovani, futuri cittadini e lavoratori, abbiano una visione equilibrata: l’AI non è né la soluzione magica a tutti i problemi, né un male da cui rifuggire, ma uno **strumento potente**. Vediamo allora, in modo articolato, alcuni dei principali **benefici** e **pericoli** legati alla diffusione dell’AI nella società.

Opportunità dell’AI

- **Automazione e aumento dell’efficienza:** L’AI può assumere compiti ripetitivi, noiosi o gravosi, liberando gli esseri umani da queste incombenze e rendendo i processi più rapidi. Nelle fabbriche i robot dotati di AI possono lavorare 24/7 senza errori di distrazione; negli uffici, algoritmi intelligenti possono sbrigare in pochi secondi attività amministrative che a una persona richiederebbero ore (come analizzare migliaia di pagine di documenti, smistare e-mail, inserire dati). Questa automazione intelligente si traduce in **produttività maggiore e costi minori** per molte aziende e servizi. Ad esempio, secondo uno studio di

McKinsey, nelle banche l'adozione diffusa dell'AI potrebbe ridurre i costi operativi fino al **25%** grazie all'ottimizzazione dei processi. In generale, svolgendo certe mansioni in modo più efficiente, l'AI permette agli umani di concentrarsi su compiti a **maggior valore aggiunto** (creativi, strategici, di relazione).

- **Nuove scoperte scientifiche e progresso medico:** L'AI è un potente alleato nella ricerca. Può analizzare set di dati giganteschi (ad esempio tutte le pubblicazioni mediche esistenti su una malattia) molto più velocemente di un team umano, scoprendo correlazioni e pattern. In campo biomedico, ci sono AI che hanno identificato nuovi possibili farmaci analizzando milioni di molecole, o che riescono a prevedere come una proteina si piegherà nello spazio (una sfida complessa risolta dal programma **AlphaFold**). In medicina, sistemi di AI supportano i medici nella **diagnosi**: un algoritmo ben addestrato può individuare segnali di un tumore in una radiografia con sensibilità pari o talvolta superiore a quella di un radiologo umano, specialmente se il medico utilizza l'AI come secondo parere. L'AI viene usata anche per monitorare epidemie, ottimizzare piani di trattamento personalizzati, e persino per assistere in interventi chirurgici (robot chirurgici controllati da AI). Tutto questo può portare a cure più veloci, precise ed efficaci, salvando vite e migliorando la qualità dell'assistenza sanitaria.
- **Personalizzazione dei servizi e miglior esperienza per gli utenti:** Avrete notato come piattaforme come Netflix o Spotify sembrano “conoscere” i vostri gusti, proponendovi film o canzoni che vi piacciono: è merito degli algoritmi di AI che analizzano le vostre preferenze e quelle di milioni di altri utenti per **personalizzare** l'offerta. Questa personalizzazione resa possibile dall'AI riguarda tanti ambiti: dallo shopping online (prodotti consigliati su Amazon su misura per noi) alla pubblicità (annunci mirati ai nostri interessi), dai percorsi di apprendimento (*e-learning* adattivo che calibra le difficoltà degli esercizi in base ai progressi dello studente) fino all'intrattenimento (videogiochi che adattano il livello di sfida al giocatore). Anche i servizi finanziari oggi sfruttano l'AI per tagliare su misura prodotti per i clienti: ad esempio, analizzando il comportamento di spesa di un utente, la banca può suggerirgli un piano di risparmio personalizzato. Tutto questo aumenta la **soddisfazione** dell'utente, che si sente compreso nei propri bisogni. Inoltre, chatbot e assistenti virtuali (basati su AI) forniscono supporto immediato h24: pensate a quando interagite con il servizio clienti di un e-commerce o del vostro operatore telefonico e vi risponde un assistente automatico che capisce la vostra domanda e vi guida nella soluzione – magari senza dover attendere minuti al telefono per parlare con un umano. Questa **immediatezza** nel servizio è un vantaggio reso possibile dall'intelligenza artificiale.
- **Supporto alle decisioni e problem solving complessi:** L'AI è bravissima nel macinare numeri e valutare tante opzioni; quindi, è ideale come “consigliere” in decisioni complesse. In ambito aziendale, vengono usati sistemi di *AI analytics* che aiutano i manager a prendere decisioni basate sui dati (il famoso approccio *data-driven*): se devi decidere quanta merce ordinare per il prossimo mese, un'AI può prevedere la domanda con modelli predittivi, riducendo gli sprechi di magazzino. Nel settore finanziario, *robo-advisor* basati su AI forniscono consulenza automatica sugli investimenti, analizzando i trend di mercato e il profilo di rischio del cliente: già oggi, i **robo-consulenti** gestiscono miliardi di dollari di patrimonio in tutto il mondo, e si stima che entro il 2025 potranno

arrivare a gestire oltre 4,6 *bilioni* (migliaia di miliardi) di dollari di asset globali. Anche nella pubblica amministrazione, l'AI può aiutare a prendere decisioni migliori: ad esempio ottimizzando i trasporti pubblici, individuando le zone della città che necessitano priorità d'intervento, o analizzando l'impatto di politiche sociali. Se usata bene, l'AI può renderci **società più efficienti e data-driven**, dove le scelte non sono prese “a naso” ma su solide analisi.

- **Affiancamento in lavori pericolosi o pesanti:** Un'opportunità da non dimenticare è l'uso dell'AI per proteggere gli esseri umani dai pericoli. I robot intelligenti possono operare in ambienti ostili dove l'uomo rischierebbe la vita: pensiamo ai droni autonomi che ispezionano l'interno di edifici pericolanti dopo un terremoto, o ai rover che puliscono siti contaminati da radiazioni, o ancora ai macchinari minerari automatici che estraggono materiali in profondità. Anche i veicoli a guida autonoma, se un domani saranno perfezionati, potrebbero ridurre drasticamente gli incidenti stradali eliminando l'errore umano (causa principale degli incidenti) – già oggi molte auto hanno sistemi di assistenza alla guida basati su AI che frenano da sole in caso di ostacolo improvviso, mantenendo la distanza di sicurezza e avvisando il guidatore di potenziali collisioni. In generale, delegare alle macchine i compiti fisicamente usuranti o rischiosi può rendere molti lavori più **sicuri** e confortevoli per le persone.

Naturalmente, l'elenco delle opportunità potrebbe continuare: dall'AI usata per combattere i cambiamenti climatici (simulando modelli climatici per prevedere eventi estremi) all'educazione (tutor intelligenti per ogni studente), dall'arte e creatività (AI che aiutano a comporre musica o disegnare nuove forme) fino alla possibilità, un domani, di avere veri e propri **“collaboratori artificiali”** che lavorino al nostro fianco in ufficio, alleggerendoci di compiti ingrati. L'importante è capire che l'AI, come strumento, **amplifica** le capacità umane: può renderci più produttivi, farci esplorare territori prima impensabili e risolvere problemi complessi, a patto di essere guidata e controllata dall'ingegno e dai valori umani.

Rischi e sfide dell'AI

Accanto ai benefici, è fondamentale riconoscere i **rischi** associati all'uso (o abuso) dell'intelligenza artificiale. Alcuni di questi pericoli sono già realtà concreta, altri sono possibili derive future da tenere sotto controllo. Ecco i principali:

- **Rischio di disinformazione e deepfake:** L'AI generativa può creare contenuti falsi *estremamente realistici*. Questo apre la porta a **truffe e manipolazioni** difficili da smascherare. In particolare, i **deepfake** – video o audio in cui l'AI sostituisce il volto o la voce di una persona con quella di un'altra in modo credibile – rappresentano una minaccia attuale. Con un deepfake, si può far apparire in video un politico mentre dice cose che non ha mai detto, o clonare la voce di una persona per ingannare un suo familiare o collega al telefono. Un sondaggio Ipsos del 2024 ha rivelato che il **46% degli italiani** crede che l'AI aumenti molto il rischio di disinformazione nella società. Paradossalmente, quasi la metà degli intervistati italiani non conosceva nemmeno il termine “deepfake”, pur essendo potenzialmente vulnerabile a questo fenomeno. È un dato che fa riflettere: la

consapevolezza del pubblico su questi rischi è ancora limitata. I deepfake possono causare seri danni: **diffamazione**, ricatti, truffe finanziarie, fino a interferenze nei processi democratici se usati per diffondere fake news su larga scala. Un esempio concreto e recente in Italia: pochi mesi fa un manager della Ferrari ha ricevuto una telefonata dal *finto* amministratore delegato Benedetto Vigna – in realtà una persona che utilizzava una voce *clonata dall'AI* identica a quella del CEO, annunciando una falsa acquisizione riservata e chiedendo di firmare documenti urgenti. Fortunatamente il manager, insospettito, ha fatto una domanda di sicurezza a cui il truffatore non ha saputo rispondere, sventando l'inganno. L'episodio, raccontato anche da Bloomberg, mostra chiaramente come i deepfake possano essere usati per *ingannare anche operatori esperti* all'interno di aziende o banche, con possibili perdite milionarie. Su scala più ampia, immagini o video deepfake diffusi online possono seminare caos: basti pensare a un falso video di un leader mondiale che annuncia un attacco militare – potrebbe scatenare reazioni incontrollate prima di essere smentito. **Distinguere il vero dal falso** diventa sempre più difficile: una recente indagine ha rilevato che il 74% degli italiani è consapevole che l'AI può generare immagini e narrazioni totalmente fasulle ma realistiche, e circa **8 persone su 10** (in Italia e nel mondo) dichiarano di essere preoccupate per questo fenomeno. Il rischio maggiore indicato è la **“crisi dei criteri di realtà”**: non sapere più di cosa fidarsi (e infatti il 50% degli intervistati teme che la democrazia stessa sia minacciata da queste manipolazioni digitali). Sarà dunque cruciale sviluppare **contromisure tecnologiche** (software in grado di rilevare contenuti fake) e normative (leggi che puniscano severamente chi li crea per nuocere), oltre che educare tutti – specialmente i più giovani – a verificare le fonti e a non credere ciecamente a tutto ciò che circola online.

Un esempio di presentazione che mette in guardia dal fenomeno dei “deepfake” (contenuti multimediali falsificati dall'AI). Immagini e video generati dall'intelligenza artificiale possono risultare così realistici da rendere difficile distinguere il falso dal vero, con rischi per la disinformazione e le truffe.

- **Bias, discriminazione e decisioni ingiuste:** Come accennato prima, i sistemi di AI possono ereditare (o addirittura amplificare) i **pregiudizi** presenti nei dati su cui sono addestrati. Questo è un pericolo serio quando l'AI viene usata in ambiti sensibili: pensiamo ai software che aiutano a selezionare candidati per un lavoro, o agli algoritmi che valutano la solvibilità di una persona per concederle un prestito. Se i dati storici mostrano – per bias sociali – che in passato poche donne sono state assunte in ruoli tecnici, un'AI di reclutamento potrebbe ingiustamente penalizzare i curricula femminili (*selection bias*). Se un algoritmo di credito viene addestrato su dati in cui un certo quartiere (magari abitato in prevalenza da minoranze) aveva più insolvenze, potrebbe **negare un prestito** a un richiedente di quel quartiere *a causa* del suo indirizzo, indipendentemente dalla sua effettiva affidabilità, replicando così pratiche discriminatorie note come *redlining*. Casi simili sono stati documentati: ad esempio negli Stati Uniti un software usato per decidere la libertà vigilata (algoritmo COMPAS) si è scoperto trattare in modo più severo i detenuti afroamericani rispetto ai bianchi, perché i dati su cui era addestrato riflettevano disparità del sistema giudiziario. Un altro esempio inquietante è VioGén, un algoritmo usato in Spagna per valutare il rischio che una donna subisca violenza dal partner: dal 2007 almeno 247 donne sono state uccise nonostante fossero state valutate a **basso rischio** dal

sistema, perché in ben il 95% dei casi gli agenti si fidavano ciecamente del punteggio dell'AI senza approfondire oltre. Questo mostra come *affidarsi ciecamente* a un algoritmo possa portare a sottovalutare situazioni gravi, specialmente se l'AI non è ben calibrata o se incorpora bias (ad esempio, forse i dati di VioGén sottorappresentavano certi fattori di rischio). Il rischio dei bias dell'AI, dunque, è duplice: da un lato **ingiustizie verso individui o gruppi** (discriminazioni di genere, razza, età, ecc.), dall'altro **perdita di fiducia** verso la tecnologia quando questi casi vengono alla luce. Immaginate il danno reputazionale per un'azienda che lancia un chatbot e scopre che dà risposte sessiste, o per una banca accusata di razzismo perché il suo algoritmo di concessione mutui penalizza sistematicamente certe etnie. Per questo motivo, oggi c'è grande attenzione sul tema dell'**AI Ethics** (etica dell'intelligenza artificiale): sviluppatori e legislatori lavorano per introdurre verifiche, audit e regolamentazioni che assicurino che i sistemi di AI siano **equi, trasparenti e privi di pregiudizi ingiustificati**. Anche noi utenti dobbiamo essere consapevoli che l'AI può sbagliare e va tenuta sotto controllo: l'ultimo giudizio, soprattutto in decisioni critiche, dev'essere umano.

- **Privacy e uso improprio dei dati personali:** L'AI si nutre di dati, tanti dati. Molte applicazioni raccolgono informazioni personali (dalle abitudini di navigazione online, ai dati sanitari, alla posizione GPS del telefono) per offrire servizi "intelligenti". Ciò però solleva forti **preoccupazioni sulla privacy**: dove finiscono tutti questi dati? Chi ne garantisce la sicurezza? C'è il rischio che finiscano in mani sbagliate? In Europa esiste il GDPR, una severa normativa sulla protezione dei dati, e presto arriveranno leggi specifiche sull'AI, ma, nonostante ciò, gli incidenti capitano: ad esempio, sono avvenute fughe di dati sensibili utilizzati per addestrare AI senza il consenso degli interessati. Un caso recente che ha fatto scalpore: l'uso di immagini pubbliche di volti (prese da social network) per addestrare algoritmi di riconoscimento facciale senza informare milioni di persone ritratte. Questo è un *uso improprio* che ha portato a cause legali. Un altro esempio: alcune app "divertenti" che trasformano i selfie in ritratti artistici grazie all'AI hanno nei Termini di Servizio clausole in cui l'utente concede all'azienda il diritto di usare a piacimento quelle foto – di fatto regalando la propria immagine. I dati personali sono definiti "il petrolio del XXI secolo", e l'AI è affamata di questo petrolio. **Proteggerli** diventa dunque vitale: il rischio è di creare una società dove ogni nostra azione è profilata e predetta da qualche algoritmo, magari senza che ce ne rendiamo conto. Pensa se un domani un'assicurazione sanitaria potesse (illegalmente) utilizzare un'AI per setacciare i tuoi post sui social e scoprire che hai abitudini poco salutari, alzandoti il premio assicurativo; o se un datore di lavoro potesse conoscere, tramite data mining, aspetti privati della tua vita. Sembrano scenari distopici ma, tecnicamente, l'AI li potrebbe rendere possibili se non poniamo paletti. Ecco perché oltre alle leggi serve **consapevolezza individuale**: leggere le informative, usare con cautela i servizi gratuiti che chiedono dati in cambio, e sostenere una cultura che privilegi la tutela della persona rispetto alla sorveglianza. Le stesse aziende dovranno investire molto in *cybersecurity* per proteggere i database su cui addestrano le AI: un sistema intelligente che custodisce dati sensibili dev'essere al sicuro da hacker e violazioni, altrimenti il danno per gli utenti sarebbe doppio (oltre alla violazione della privacy, i dati rubati potrebbero essere usati per frodi, furti d'identità, ecc.). La **fiducia del pubblico**

verso l'AI dipenderà molto da come sapremo gestire in modo trasparente e sicuro le informazioni personali.

- **Impatto sul lavoro e necessità di adattamento:** Uno dei temi più dibattuti è: *l'AI ci ruberà il lavoro?* La risposta non è semplice. Sicuramente l'automazione intelligente **trasformerà** il mondo del lavoro. Alcune professioni verranno *in parte o del tutto sostituite* dalle macchine, mentre ne nasceranno di nuove legate alla gestione dell'AI. Secondo l'Osservatorio del Politecnico di Milano, in Italia entro 10 anni potrebbero essere automatizzati (quindi *eliminati* in forma attuale) circa **3,8 milioni di posti di lavoro**. Si tratta soprattutto di lavori manuali routinari o mansioni d'ufficio ripetitive. Questo numero può spaventare, ma va letto insieme ad altri dati: a causa dell'invecchiamento demografico, in Italia si prevedono comunque 5,6 milioni di posti vacanti entro il 2033 per mancanza di lavoratori giovani. Paradossalmente, l'AI potrebbe colmare proprio parte di questo gap, automatizzando quei 3,8 milioni di posizioni e compensando la carenza di manodopera. Il vero punto è che **cambieranno le competenze richieste**: i lavoratori dovranno sapersi adattare, imparare nuove abilità complementari all'AI, spostarsi su mansioni dove l'umano aggiunge valore (creatività, empatia, pensiero critico). Molti esperti parlano di **“cooperazione uomo-macchina”**: l'AI non rimpiazzerà completamente l'uomo, ma chi saprà usarla *bene* rimpiazzerà chi non la sa usare. Ad esempio, in uno studio legale l'AI potrà preparare bozze di contratti standard in pochi secondi (facendo il lavoro di 5 praticanti), ma quei praticanti potrebbero dedicarsi ad altro – ricerca giuridica più avanzata, consulenza personalizzata ai clienti, ecc. Emergeranno nuovi lavori: dal **“prompt engineer”** (l'esperto nel dialogare con i modelli generativi per ottenere il meglio), agli specialisti di AI governance (chi tarerà e monitorerà gli algoritmi perché siano equi e conformi alle norme). La **sfida per la società** è gestire la transizione: formare i lavoratori la cui mansione è a rischio, sostenere chi perderà il posto in quella fase e facilitare la creazione di nuovi ruoli. Storicamente, ogni rivoluzione tecnologica (macchine industriali, computer, internet) ha sì tolto lavoro in certi settori ma ne ha creato di nuovi spesso migliori. Serve ottimismo **ma anche realismo**: alcuni lavori spariranno (es. centralinisti, operatori di catena di montaggio, forse autisti con i veicoli autonomi), e non tutti i lavoratori colpiti riusciranno a riqualificarsi facilmente. Perciò, la questione dell'AI e lavoro è anche una questione di **politiche sociali**: se ne discute a livello di governi di tutto il mondo, valutando idee come la riduzione dell'orario di lavoro (se le macchine ci rendessero più produttivi, potremmo lavorare meno ore a parità di output) o nuovi modelli di welfare per sostenere chi è in transizione occupazionale. In sintesi, l'AI è un rischio per l'occupazione tradizionale, ma può essere un'opportunità per ripensare il lavoro, renderlo meno alienante e più sicuro. Starà alla nostra generazione prepararsi con le giuste **competenze digitali e flessibilità mentale** per cogliere i nuovi lavori che verranno.
- **Dipendenza tecnologica e perdita di controllo:** Un rischio più sottile è quello di diventare **troppo dipendenti dalle decisioni delle AI**, perdendo abilità o delegando eccessivo potere alle macchine. Se ci abituiamo a un'auto che parcheggia da sola, rischiamo di disimparare completamente a farlo noi. Se lasciamo che sia l'algoritmo dei social a decidere quali notizie leggiamo, potremmo chiuderci in bolle informative senza spirito critico. In contesti critici, delegare all'AI può essere pericoloso: immaginiamo un futuro in

cui la maggior parte delle operazioni finanziarie in borsa siano automatizzate – se un algoritmo impazzisse potrebbe causare crolli lampo (flash crash) prima che un umano intervenga. Oppure pensiamo a sistemi militari autonomi: droni armati con AI, se non adeguatamente controllati, aprono dilemmi etici enormi (chi è responsabile di un errore fatale commesso dalla macchina?). Anche senza andare agli estremi, già oggi assistiamo a piccoli problemi: GPS che fanno finire gli automobilisti su strade improbabili perché si fidano ciecamente delle indicazioni, o utenti che prendono per oro colato i consigli di una app sulla salute senza consultare un medico. **Mantenere il controllo umano** sull'AI è essenziale. Questo è riconosciuto anche a livello normativo: l'Unione Europea sta in parte ancora lavorando sull'**AI Act**, una legge che tra le altre cose richiede che in ambiti ad alto rischio (per esempio, decisioni su assunzioni, giustizia, sorveglianza pubblica) l'AI sia sempre trasparente e sotto supervisione umana. Bisogna evitare di scivolare in un mondo dove le macchine prendono decisioni incomprensibili e noi le subiamo. Ciò significa anche sviluppare AI dotate di **“spiegabilità”** (explainable AI): algoritmi i cui criteri di decisione siano interpretabili, così da poter contestare o correggere un risultato algoritmico se appare sbagliato.

- **Timori esistenziali (superintelligenza fuori controllo):** Infine, c'è un dibattito, per ora teorico ma molto vivo, sui rischi di una futura **AI superintelligente** che possa sfuggire al controllo umano. Figure come l'astrofisico Stephen Hawking o imprenditori come Elon Musk negli anni scorsi hanno avvertito che un'AI evoluta potrebbe rappresentare una minaccia per la nostra stessa sopravvivenza. È uno scenario da fantascienza? Al momento sì: non esiste nulla di paragonabile a un'AI autocosciente. Ma alcuni filosofi e scienziati sottolineano che, se mai dovessimo riuscire a crearne una, dovremmo assicurarci che i **suoi obiettivi siano allineati ai nostri**. C'è chi propone fin d'ora di inserire *“principi etici”* nelle AI, un po' come le leggi della robotica di Asimov (prima fra tutte: un robot non può arrecare danno agli esseri umani). Questo è un campo ancora filosofico, ma che forse i ragazzi della vostra generazione si troveranno a dover affrontare concretamente se l'AI continuerà ad avanzare. Vale la pena menzionarlo per completezza: **avere coscienza dei potenziali rischi a lungo termine** ci permette di sviluppare da subito un'AI in modo più sicuro e responsabile.

Riassumendo, l'AI offre possibilità straordinarie ma comporta una serie di sfide non banali: dovremo imparare a **sfruttare i vantaggi minimizzando i rischi**. Questo richiede competenze tecniche (per costruire e controllare le AI), ma anche dialogo tra diverse discipline: giuristi per le leggi, psicologi per capire l'impatto sociale, filosofi per guidare le scelte etiche, economisti per gestire le ricadute sul lavoro. E richiede, soprattutto, **cittadini consapevoli** – come voi ragazzi dovrete diventare – capaci di comprendere questi strumenti e di prendere posizione sul loro utilizzo nella società.

Come viene usata l'AI nel contesto bancario?

Uno dei settori in cui l'intelligenza artificiale sta avendo un impatto significativo è proprio quello **bancario e finanziario**. Le banche, tradizionalmente viste come istituzioni lente e conservative,

negli ultimi anni hanno invece **abbracciato con decisione l'innovazione AI** per migliorare servizi e operatività. Sia a livello internazionale che in Italia, l'AI sta rivoluzionando il modo in cui operano gli istituti di credito: dalle interazioni con la clientela, alla gestione dei rischi, fino ai processi interni. Esaminiamo i principali ambiti di applicazione.

1. Ottimizzazione dei processi interni e riduzione dei costi: Le banche gestiscono enormi volumi di operazioni ripetitive ogni giorno (pagamenti, bonifici, verifiche di identità, compilazione di documenti, controlli normativi). L'AI consente di **automatizzare molti di questi processi** aumentando efficienza e velocità. Ad esempio, compiti di *back-office* come il controllo dei documenti per l'apertura di un conto, la verifica dell'identità del cliente (**KYC – Know Your Customer**) o il monitoraggio della conformità alle normative possono essere svolti da sistemi intelligenti, riducendo errori e tempi di attesa. Tecnologie come la **RPA (Robotic Process Automation)**, arricchite con AI, permettono di eseguire transazioni e aggiornamenti su vari sistemi come farebbe un impiegato, ma in maniera instancabile e rapidissima. In Italia, per esempio, si sta sperimentando un sistema basato su *machine learning* che legge e processa automaticamente migliaia di nuove norme di vigilanza bancaria, per aiutare i dipartimenti legali a tenere il passo con i continui aggiornamenti regolamentari. Così, un lavoro che prima richiedeva decine di analisti e molto tempo ora viene svolto in pochi istanti, con gli esperti umani che possono concentrarsi sulle valutazioni finali di alto livello. In generale, **efficienza** è la parola chiave: l'AI in banca può servire ad automatizzare ciò che prima veniva fatto manualmente.

2. Miglioramento dell'esperienza cliente (Customer Experience): Avrete probabilmente interagito almeno una volta con un **chatbot** bancario. Molti istituti hanno introdotto assistenti virtuali online o via app che rispondono alle domande frequenti dei correntisti, forniscono informazioni su saldo e movimenti, aiutano a svolgere operazioni semplici. Questi **assistenti AI** permettono ai clienti di ricevere risposta immediata senza dover aspettare al telefono o recarsi in filiale: si può chiedere “Come faccio a modificare il massimale della carta di credito?” o “Qual è l'IBAN del mio conto?” e ottenere subito la guida passo-passo. Per le banche significa **alleggerire i call center** (riducendo i costi) e allo stesso tempo offrire un servizio più rapido. Naturalmente per questioni complesse o personalizzate c'è sempre l'operatore umano di secondo livello, ma intanto il bot smaltisce le richieste base. L'AI nel customer service non si limita alle chat: esistono anche sistemi di **IVR intelligenti** (risponditori telefonici automatici potenziati dall'AI) che riconoscono la voce del cliente e l'intento della chiamata, smistando le telefonate in modo più accurato di prima. Inoltre, grazie all'analisi avanzata dei dati, le banche stanno passando da un approccio uguale per tutti a un approccio **personalizzato per ciascun cliente**: algoritmi analizzano lo storico delle transazioni e delle interazioni e possono così proporre al cliente prodotti o servizi più adatti alle sue esigenze o abitudini. Questa **proattività** guidata dall'AI migliora la soddisfazione e spesso porta anche benefici alla banca (il cliente scopre servizi che altrimenti non avrebbe utilizzato). In sintesi, l'AI aiuta le banche a essere **più vicine e reattive** verso i loro utenti, pur mantenendo un rapporto scalabile (un singolo assistente AI può gestire contemporaneamente migliaia di conversazioni, cosa impossibile per degli operatori umani).

3. Sicurezza, antitruffa e gestione del rischio: La sicurezza è un ambito cruciale per le banche e qui l'AI è diventata un alleato imprescindibile. Oggi i sistemi bancari processano milioni di transazioni al giorno: intercettare manualmente quelle fraudolente (es. clonazioni di carte, bonifici illeciti) sarebbe come cercare un ago in un pagliaio. **Algoritmi di AI per il fraud detection**

invece possono monitorare in tempo reale tutte le operazioni e **segnalare quelle anomale** con un alto livello di accuratezza. Ad esempio, se la tua carta di credito di solito viene usata solo in Italia per spese sotto i 100€, e improvvisamente appare un addebito di 1000€ in un altro continente, l'AI lo noterà immediatamente come fuori dal tuo "profilo" e potrà bloccare transitoriamente l'operazione, in attesa di verifica. Questi sistemi imparano dai dati passati quali schemi tipicamente indicano una frode (importi strani, luoghi inusuali, frequenza elevata di transazioni ravvicinate, ecc.) e **prevenire i furti** prima che causino grosse perdite. Lo stesso vale per i tentativi di attacco informatico alle piattaforme di *home banking*: l'AI può rilevare accessi sospetti (per esempio, decine di tentativi di login falliti che potrebbero indicare un attacco brute-force) e far scattare allarmi. In pratica, l'AI agisce come una sorta di **guardia del corpo digitale** che veglia sui conti e i sistemi informatici bancari 24 ore su 24. Questo ha un impatto enorme: non solo riduce le frodi (tutelando sia la banca sia i clienti), ma aumenta anche la **fiducia** degli utenti nel canale digitale. Sapere che c'è un cervello elettronico sempre all'erta dietro le quinte dà maggiore tranquillità nel fare operazioni online. Oltre alle frodi, l'AI è impiegata nella **gestione del rischio finanziario**: ad esempio modelli predittivi valutano la probabilità che un certo prestito diventi insolvente, analizzando centinaia di variabili sul cliente e sul contesto economico; oppure stimano l'esposizione della banca a possibili oscillazioni di mercato (rischio di tasso, di cambio, ecc.) così che la banca possa coprirsi per tempo. Tuttavia, come discusso prima, bisogna vigilare affinché questi sistemi non incorporino bias discriminatori. Per questo molte banche affiancano alle decisioni automatiche comunque un **controllo umano**, specie se si tratta di rifiutare una richiesta, così da verificare che non ci siano ingiustizie. In ogni caso, è chiaro che l'AI sta rendendo i servizi finanziari **più sicuri e controllati**: le frodi vengono bloccate sul nascere, le perdite per default di crediti si riducono perché si anticipano i problemi, e perfino il rischio di crisi viene monitorato con modelli che analizzano indicatori macroeconomici e segnalano possibili stress all'orizzonte.

4. Adozione dell'AI generativa e nuove frontiere: La più recente ondata tecnologica – quella dell'**AI generativa** come ChatGPT – non ha risparmiato il settore finanziario. Molte banche stanno esplorando l'uso di modelli linguistici avanzati per **migliorare ulteriormente l'interazione con i clienti** e l'efficienza interna. Un rapporto del 2025 evidenzia che a livello globale il **96% delle banche** sta già implementando sistemi basati sull'AI generativa e il 71% sta aumentando gli investimenti in questo campo. L'AI generativa può essere usata, ad esempio, per *rispondere in linguaggio naturale* a quesiti complessi dei clienti (come fa ChatGPT, ma addestrato sui dati specifici della banca), oppure per generare automaticamente **report e documenti**. Immaginiamo un gestore che deve preparare una relazione trimestrale per un cliente premium: potrebbe chiedere a un'AI interna "Riassumi gli andamenti del portafoglio di Mario Rossi negli ultimi 3 mesi e spiega le variazioni principali" e ottenere una bozza ben formattata in pochi secondi, da ritoccare e consegnare al cliente. Oppure, nel lavoro di **programmazione IT** in banca, l'AI generativa può aiutare i tecnici suggerendo codice, documentando procedure o rispondendo a domande sul funzionamento di vecchi sistemi (diventando un assistente per gli sviluppatori e riducendo il tempo di sviluppo di nuovi software bancari). Un'altra applicazione è la **formazione del personale**: chatbot addestrati sulle policy interne che rispondono alle domande dei nuovi assunti ("Come faccio ad avviare una procedura X?") fornendo subito la risposta tratta dai manuali aziendali. Insomma, le possibilità sono moltissime. Le banche però, specie in Europa, stanno procedendo con **prudenza e gradualità** nell'adottare queste novità. Ciò è dovuto sia a motivi

regolamentari (la normativa richiede cautela quando si introducono nuovi strumenti che potrebbero impattare i clienti, e c'è sempre il tema della privacy dei dati), sia culturali (il settore bancario europeo è meno incline a movimenti rivoluzionari e più orientato a miglioramenti incrementali). Diversi istituti stanno conducendo **progetti pilota** e test per capire come integrare l'AI generativa in modo sicuro.

5. Sfide e aspetti da monitorare: L'entusiasmo per l'AI nelle banche, come nella maggior parte delle altre aziende, va di pari passo con nuove **sfide**. Uno dei timori principali è l'**impatto occupazionale**: se l'AI automatizza molte funzioni, cosa succede ai dipendenti? Negli ultimi anni, in effetti, si è registrata una diminuzione degli organici tradizionali e la chiusura di molte filiali fisiche a favore dei canali online. Le banche italiane hanno però avviato programmi di riqualificazione del personale: ad esempio, formare gli impiegati di sportello in competenze digitali e spostarli su ruoli di consulenza personalizzata, lasciando alle macchine le operazioni base. Un'altra questione è quella **etica e regolatoria**: come assicurarsi che le decisioni prese dall'AI in banca siano trasparenti e non discriminatorie? Le autorità di vigilanza finanziaria (come Banca d'Italia, BCE, etc.) stanno emanando linee guida sull'uso responsabile dell'AI. Ad esempio, se un cliente si vede rifiutare un prestito deve aver diritto a una spiegazione chiara: “perché il modello mi ha bocciato?”. Questo spinge le banche a usare modelli interpretabili o a fornire comunque un *contatto umano* per discutere le decisioni. Inoltre, c'è il tema della **sicurezza informatica**: affidandosi all'AI si aumenta anche la “superficie d'attacco” per i criminali informatici (immaginate se un hacker manipolasse il sistema di rilevamento frodi per far passare transazioni illecite). Serve quindi un impegno fortissimo sulla *cybersecurity* per proteggere questi sistemi. Infine, non va trascurato l'aspetto **culturale**: l'adozione efficace dell'AI richiede un cambio di mentalità dentro la banca. Non è solo una sfida tecnologica ma un **cambiamento organizzativo**. Occorre che dipendenti e manager si fidino dei nuovi strumenti, li comprendano e collaborino con essi, invece di vederli come “scatole nere” incomprensibili. Le banche più lungimiranti stanno investendo nel *training* interno per spiegare ai loro team come funziona l'AI e come usarla al meglio, creando una cultura aziendale “data-driven”. Quelle che resteranno indietro rischiano di perdere competitività. A livello internazionale, le differenze regionali sono già visibili: secondo NTT Data, banche di Stati Uniti e Asia vedono l'AI soprattutto come leva per **vantaggio competitivo** e innovazione spinta, mentre le banche europee la approcciano più per migliorare produttività e processi interni, riflettendo un contesto normativo più prudente e un atteggiamento conservativo. In ogni caso, il **trend è chiaro**: l'AI diventerà parte integrante di tutte le aree operative bancarie, rendendo il settore più **predittivo, interattivo e personalizzato** per il cliente. Quello che vediamo ora è probabilmente solo l'inizio di una trasformazione profonda.

Altre considerazioni e conclusione

Abbiamo attraversato un panorama ampio: dalla definizione di intelligenza artificiale ai suoi vari tipi, dalle opportunità ai rischi, fino a uno sguardo concreto sul mondo bancario. Prima di concludere, vale la pena ribadire alcuni punti chiave e aggiungere un paio di riflessioni finali.

Consapevolezza, non magia: Spesso l'AI viene presentata nei media in modo sensazionalistico, quasi fosse una magia incomprensibile. In realtà, come abbiamo spiegato, dietro c'è tanta

matematica e informatica, non trucchi sovrannaturali. È importante che vi avviciniate all'AI con curiosità ma anche con senso critico. Dovete sapere che, quando uno smartphone riconosce il vostro volto, sta semplicemente confrontando l'immagine con un modello matematico ricavato da milioni di altri volti; quando ChatGPT vi risponde, sta pescando tra miliardi di frasi apprese quale è la più probabile prosecuzione del vostro prompt – non c'è un *pensiero* nel senso umano dietro le sue parole. Questa lucidità vi aiuterà a non farvi abbagliare e a **usare l'AI per quello che è: uno strumento**. Un martello amplifica la forza del braccio, un microscopio amplifica la vista: ecco, un'AI amplifica la nostra capacità di analizzare dati e prendere decisioni rapide. Ma resta un *mezzo*, mentre il **fine** e la responsabilità restano nostri.

Opportunità per la vostra generazione: Voi adolescenti di oggi crescete in un'epoca in cui l'AI sarà sempre più presente. Questo può essere uno stimolo enorme: avrete a disposizione strumenti che i vostri genitori nemmeno immaginavano alla vostra età. Potrete imparare più in fretta (pensate ai tutor AI personalizzati nello studio), esprimere creatività in nuovi modi (magari collaborando con un'AI per creare arte o musica), o inventare lavori che oggi non esistono. L'Italia ha tanto bisogno di giovani formati sulle nuove tecnologie: il mercato dell'AI nel nostro Paese sta crescendo a ritmo sostenuto (+52% nel 2023, per 760 milioni di € di valore), e mancano figure professionali adeguate. Significa che se vi appassionarete a queste tematiche e sviluppate competenze in campo AI, avrete **opportunità di carriera** significative, contribuendo allo sviluppo tecnologico nazionale. Non è necessario diventare tutti programmatori di AI; serviranno anche giuristi esperti di AI, psicologi del lavoro per l'AI, comunicatori scientifici per spiegare l'AI al pubblico, imprenditori che sappiano usare l'AI per innovare settori tradizionali, ecc. La **alfabetizzazione AI** diventerà un elemento importante del vostro bagaglio, qualunque strada professionale intraprendiate. Quindi il consiglio è: sperimentate, informatevi, non abbiate paura di *mettere le mani* su queste tecnologie (sempre in modo etico). Più capirete come funzionano, più ne coglierete il potenziale e anche i limiti, diventando cittadini digitali consapevoli.

Responsabilità e approccio etico: L'AI, come abbiamo visto, solleva anche interrogativi morali e richiede un utilizzo responsabile. La vostra generazione sarà chiamata a prendere decisioni importanti: ad esempio, come regolamentare l'uso dell'AI nelle armi? Come assicurare che un'AI usata in tribunale non sia di parte? Come evitare che le AI generative diffondano odio o fake news? Sono dilemmi che non possiamo lasciare solo agli "esperti tecnici": riguardano tutti noi, perché l'AI tocca aspetti sociali, giuridici, umani. Dovremo tenere al centro i **valori umani**: l'AI dovrà servire a migliorare la vita delle persone, non a peggiorarla o a creare nuove ingiustizie. Questo richiede vigilanza, leggi adeguate e soprattutto una forte base etica in chi progetta e utilizza queste tecnologie. Come citato prima, già oggi l'UE è capofila nel cercare di normare l'intelligenza artificiale con l'AI Act, classificando i sistemi in base al rischio e imponendo requisiti di trasparenza, sicurezza e non discriminazione. È probabile che sentirete parlare sempre più di **"AI affidabile"**, **"AI trasparente"**, **"AI spiegabile"**: sono concetti chiave su cui si baserà la fiducia pubblica. Voi in quanto utenti potete fare la vostra parte premiando con la preferenza servizi e prodotti AI che rispettano la vostra privacy e i vostri diritti, ed evitando quelli opachi o scorretti. In futuro, alcuni di voi potranno trovarsi a progettare AI: ricordatevi in quel caso di inserire nel processo di sviluppo la domanda "questo a chi giova? potenziali danni collaterali per qualcuno? come possiamo mitigarli?". Una mentalità etica è il miglior antidoto agli effetti negativi involontari.

L'avventura dell'intelligenza artificiale si prospetta come uno degli elementi più caratterizzanti del XXI secolo. L'AI sta entrando in tutte le sfere, dalla quotidianità (con gli assistenti vocali, le traduzioni automatiche, le fotocamere intelligenti dei telefoni) ai settori strategici come la finanza, la medicina, l'educazione, la sicurezza nazionale. Per voi giovani, parlarne non significa solo imparare definizioni o casi d'uso, ma **prepararsi a un mondo in cui collaborerete con queste intelligenze artificiali**. Il messaggio da portare a casa è duplice: da un lato, *non abbiate paura* dell'AI, esploratela con mente aperta perché sarà parte della vostra vita e può arricchirla; dall'altro, *non mitizzatela acriticamente*, perché è frutto dell'ingegno umano e come tale può sbagliare o essere usata male, quindi richiede sempre il vostro giudizio. Come disse un ricercatore, "l'intelligenza artificiale non è né artificiale (perché richiede enormi risorse materiali, energia, lavoro umano dietro le quinte) né realmente intelligente (perché non ha comprensione): è **calcolo al servizio dell'intelligenza umana**". Usata bene, può aiutarci a realizzare cose prima impensabili; usata male o senza controllo, può creare nuovi problemi. Spetterà anche a voi, con la vostra **consapevolezza e creatività**, guidare l'AI nella direzione giusta – quella che pone la tecnologia al servizio dell'uomo e non viceversa.

Speriamo che questo saggio vi abbia fornito le basi per capire e discutere di AI in modo informato. Il mondo dell'AI è in rapido movimento: restate aggiornati, ponetevi domande, e chissà che qualcuno di voi non sia un futuro scienziato o imprenditore in questo campo! In ogni caso, come cittadini, sarete chiamati a navigare un futuro dove l'intelligenza naturale (la vostra) e quella artificiale dovranno convivere. **Prepariamoci con conoscenza e senso critico**, così che questa convivenza sia proficua e armoniosa. L'AI è un potente strumento: il *come* verrà usato dipende soprattutto da noi.