



**Politecnico
di Torino**

POLITECNICO DI TORINO

**Collegio di Ingegneria Gestionale – Classe L/9
Corso di Laurea in Ingegneria Gestionale**

Tesi di Laurea di II livello

**Intelligenza artificiale e mercati digitali:
trasformazioni economiche, sfide regolatorie e
nuove prospettive degli AI agents**

**Relatore:
Prof. Carlo Cambini**

**Candidata:
Elena Zuccolo**

Anno accademico 2024-2025

Indice

1 Introduzione.....	3
2 Struttura mercato dell'AI stack	5
2.1 Introduzione al concetto di AI stack	6
2.1.1 L'AI stack nell'ecosistema dell'intelligenza artificiale	6
2.1.2 Funzioni concorrenziali e punti nevralgici dello stack	7
2.2 Struttura dell'AI Stack	8
2.2.1 Hardware e infrastrutture di calcolo	8
2.2.2 Servizi cloud.....	11
2.2.3 Dati	13
2.2.4 Foundation Models	17
2.2.5 Applicazioni	23
2.3 Dinamiche competitive	27
2.3.1 Concentrazione per livello e interdipendenze	27
2.3.2 Meccanismi economici comuni nell'AI stack	28
2.3.3 Integrazione verticale e leve cross-layer.....	29
3 Effetti economici ed organizzativi dell'AI.....	32
3.1 Occupazione.....	33
3.1.1. L'AI ed il futuro dell'occupazione	33
3.1.2 L'approccio <i>task-based</i> e le implicazioni organizzative.....	33
3.1.3 Evidenze empiriche	34
3.1.4 Salari e qualifiche	35
3.1.5 La polarizzazione occupazionale	37
3.2 Produttività e crescita.....	38
3.2.1 Il paradosso di Solow e la J-curve	38
3.2.2 Evidenze empiriche	40
3.2.3 Co-invenzione e disomogeneità tra imprese	42
3.2.4 Implicazioni macroeconomiche	43
3.2.5 Stime e previsioni quantitative sulla produttività	43
3.3 Organizzazione aziendale	47
3.3.1 Trasformazione delle strutture organizzative ed approccio <i>data-driven</i>	47
3.3.2 Modularità e interdipendenze	48
3.3.3 Capitale umano e complementarità	50
3.3.4 Cambiamento dei processi decisionali e nuovi ruoli organizzativi.....	50
3.3.5 Sfide organizzative dell'AI nelle imprese	51
3.4 Innovazione e modelli di business	53
3.4.1 Piattaforme digitali	53
3.4.2 AI come fonte di personalizzazione	53
3.4.3 Il futuro dell'AI generativa	54
3.5 Educazione.....	55
3.5.1 Trasformazione della domanda di competenze e AI literacy.....	55
3.5.2 Competenze umanistiche e nuove figure professionali.....	56
3.6 Concorrenza e struttura dei mercati	58
3.6.1 Effetti pro-competitivi	58
3.6.2 Effetti anti-competitivi	59
3.6.3 Discriminazione di prezzo e potere di mercato legato ai dati	60
3.6.4 Rischi di collusione nei mercati digitali	62
3.6.5 Fusioni ed acquisizioni: il caso americano	63
3.6.6 Venture Capital, imprenditorialità e barriere per i nuovi entranti.....	64
4 Regolamentazione e policy dell'AI.....	67

4.1 I <i>driver</i> della regolamentazione	67
4.2 Strumenti regolatori e governance	70
4.2.1 AI Act.....	70
4.2.2 Oltre l'AI Act: le altre norme europee di riferimento	71
4.2.3 Sandboxes.....	72
4.3 Contesto internazionale e cooperazione multilaterale.....	73
4.4 AI continent action plan EU.....	75
4.4.1 Obiettivo e scopo.....	75
4.4.2 Iniziative sulle risorse computazionali e le infrastrutture	76
4.4.3 Verso una governance del calcolo	79
4.4.4 Data Labs ed European data union strategy.....	79
4.5 Sfide e prospettive future	80
5 AI Agents	82
5.1 Caratteristiche degli AI agents.....	82
5.2 Architettura e struttura organizzativa dell'agency	85
5.3 I vari tipi di AI agents	87
5.3.1. Categorie di alto livello	87
5.3.2 Configurazioni specializzate.....	88
5.4 Settori d'impatto ed applicazioni	89
5.5 Posizionamento degli AI agents nello stack.....	90
5.6 Impatto sulle piattaforme e sui modelli di business	91
5.7 Impatto sulla concorrenza	93
5.7.1 Premesse generali	93
5.7.2 Esclusione degli AI agents	94
5.7.2.1 Esclusione a monte	94
5.7.2.2 Esclusione a valle.....	96
5.7.3 Esclusione da parte degli AI agents.....	96
5.8 Applicabilità del DMA.....	98
5.8.1 Criteri per la designazione.....	98
5.8.2 Categorizzazione dei CPS.....	100
5.8.3 Obblighi e divieti previsti per i CPS	101
5.9 Politiche europee di sostegno e regolazione	102
5.10 Situazione attuale: fallimenti, hype ed aspettative per il futuro	103
6 Conclusioni.....	107
Bibliografia.....	109
Sitografia	116

1 Introduzione

Negli ultimi anni l'intelligenza artificiale si è evoluta ad un ritmo senza precedenti, affermandosi come una tecnologia pervasiva e flessibile. Essa è infatti diventata una priorità strategica per imprese, governi ed istituzioni internazionali, attirando l'attenzione anche dei semplici utenti o consumatori, che l'hanno fatta entrare nella propria quotidianità attraverso l'utilizzo di chatbot, sistemi di raccomandazione o piattaforme intelligenti. Non si tratta più di una tecnologia "per pochi", in quanto permea settori molto differenti tra loro, come ad esempio la finanza, la sanità, la difesa, l'educazione, la manifattura ed i servizi.

I progressi nell'AI generativa hanno permesso la ridefinizione di processi economici e sociali, nonché di modelli organizzativi e decisionali, anche grazie all'ampio spettro di settori su cui essa è in grado di impattare. Si tratta, infatti, di una *general purpose technology* (GPT), che se legata ad altre tecnologie, quali *cloud computing*, *data science*, *business intelligence*, *data mining*, *big data* e *natural language processing* (NLP), permette di sprigionare il massimo del suo potenziale, determinando trasformazioni rivoluzionarie in vari ambiti. Proprio per queste caratteristiche viene sempre più associata e paragonata alle precedenti ondate tecnologiche, come la rivoluzione industriale o l'ICT, le quali hanno segnato una discontinuità rispetto ai paradigmi precedenti.

I passi in avanti che sono stati fatti recentemente sono correlati all'aumento della potenza di calcolo, alla disponibilità di dati e all'evoluzione dei *foundation models* (FMs), cuore pulsante di tutti i sistemi di intelligenza artificiale. I FMs si collocano ad uno stadio successivo rispetto al *deep learning*, branca del *machine learning* che si basa su reti neurali profonde in modo da simulare il funzionamento del cervello umano. Tale architettura permette ai modelli di imparare dai dati in input e di estrarne rappresentazioni complesse, portando i dati stessi a diventare un asset cruciale per l'addestramento e la competitività dei sistemi di AI.

In questo contesto si sviluppa una catena del valore verticalmente integrata, l'*AI stack*, in cui hardware, cloud, dati, modelli ed applicazioni interagiscono generando economie di scala e di scopo. Questo pone interrogativi riguardanti la contestabilità dei mercati e la sovranità tecnologica in quanto la struttura attuale, e presumibilmente quella futura, vede il potere tecnico ed economico concentrato in pochi attori globali, richiedendo interventi per un accesso più equo alle risorse fondamentali.

Un altro aspetto cruciale è quello economico, in quanto gli effetti previsti sono ambivalenti: da una parte l'intelligenza artificiale ha tutte le caratteristiche necessarie per portare ad un aumento della produttività, alla creazione di nuove figure professionali e, più in generale, ad una

ridefinizione del mercato e delle strutture organizzative. D'altro canto, però, emergono rischi per quanto riguarda la disoccupazione o la polarizzazione della ricchezza con conseguente inasprimento della disuguaglianza già presente tra i vari Paesi. Tali dinamiche potrebbero incidere profondamente sulla concorrenza globale, alterando gli equilibri consolidati. Accanto a tali problematiche, l'AI solleva anche questioni di natura sociale e culturale, legate al confine entro cui ci si potrà spingere nella sostituzione dell'uomo.

In risposta a tali sfide, le autorità preposte al controllo stanno agendo in modo tale da proporre un quadro regolatorio adeguato, affiancato da policy ed iniziative che incentivino allo sviluppo ed all'innovazione rispettando, però, la contestabilità del mercato. In questi termini, diversi Paesi assumono diversi atteggiamenti, indicando che ancora non c'è una visione unificata di come ci si debba muovere nel contesto creatosi.

Di recente, inoltre, si stanno affermando gli *AI agents*, sistemi autonomi e proattivi capaci di prendere decisioni ed agire in nome dell'utente. Essi sono un'implementazione avanzata dei FMs, che rimangono il nucleo cognitivo, ma che vengono affiancati da memoria a lungo termine, *tools* e capacità di ragionamento. Guardando al futuro, potrebbero diventare le nuove interfacce tra uomo e macchina, ridefinendo in maniera radicale il modo in cui oggi l'umano interagisce con i dispositivi digitali.

La presente tesi nasce con l'obiettivo di approfondire tali argomenti attraverso un approccio qualitativo e comparativo, basato sull'analisi della letteratura economica e giuridica, dei documenti istituzionali e dei regolamenti attuali.

La tesi è strutturata in cinque capitoli, di cui uno introduttivo. Il secondo, il terzo ed il quarto, i quali trattano rispettivamente l'AI stack, gli effetti attesi della recente ondata tecnologica e la regolamentazione e policy, sono stati scritti in collaborazione col collega Vincenzo Salvatore, mentre l'ultimo, riguardante gli AI agents, è stato redatto individualmente.

2 Struttura mercato dell'AI stack

L'avvento dell'intelligenza artificiale ha portato alla creazione di una propria catena del valore, comunemente definita *AI stack*, che si innesta e interagisce con mercati tecnologici preesistenti, come quello delle infrastrutture cloud, dei chip e dei servizi digitali avanzati.

Nell'analizzare l'*AI stack* non bisogna soffermarsi soltanto sul descrivere una sequenza tecnica di componenti, ma è opportuno comprendere l'ecosistema nel suo complesso, poiché l'evoluzione di uno stadio genera cambiamenti negli altri livelli della catena, creando o circoli virtuosi o, al contrario, colli di bottiglia che possono rallentare l'innovazione. Infatti, rispetto alle catene del valore tradizionali, questa struttura si distingue per un'elevata interdipendenza tra i diversi livelli a causa della presenza di forti retroazioni positive o negative, economie di scopo e di scala e per gli effetti di rete che possono essere diretti o indiretti.

Per comprendere a fondo la struttura dello *stack* è quindi essenziale analizzare le dinamiche del mercato dell'AI e cogliere le interdipendenze tra i diversi stadi della catena, così da individuare i punti critici in cui si concentra il valore e il potere di mercato, nonché le aree dove si formano le principali barriere all'ingresso. Questi elementi, se non bilanciati da adeguati meccanismi competitivi, possono portare a situazioni di monopolio o a una riduzione della contestabilità del mercato sia a monte sia a valle.

In questo capitolo viene dapprima fornita una visione complessiva dello *stack*, utile a definire il quadro concettuale e le principali categorie che lo compongono; segue l'analisi di ciascun segmento, dall'hardware e dalle infrastrutture di calcolo fino ai *foundation models* e alle applicazioni finali, evidenziandone le caratteristiche tecniche, le dinamiche competitive e le implicazioni per la catena del valore.

2.1 Introduzione al concetto di AI stack

2.1.1 L'AI stack nell'ecosistema dell'intelligenza artificiale

L'*AI stack* rappresenta un insieme strutturato di segmenti tecnologici interdipendenti che comprende sia il processo di creazione che di distribuzione dei contenuti creati dall'intelligenza artificiale, rendendo possibile lo sviluppo, l'addestramento, la distribuzione e l'utilizzo di sistemi AI. Questa configurazione attuale dello *stack* è il risultato di un'evoluzione recente, dovuta alle innovazioni nella potenza di calcolo, negli algoritmi di *machine learning*, nei servizi cloud e alla disponibilità su larga scala di dati. Tali fattori hanno contribuito anche a ridefinire i rapporti di forza tra i diversi livelli che la compongono.

La catena del valore dell'intelligenza artificiale si sviluppa verticalmente in una serie di livelli interdipendenti come riportato in *figura 1*. A monte si trovano le componenti hardware, costituite dai chip e software specializzati. Su questo strato si innesta il livello dei servizi cloud, che forniscono potenza di calcolo, capacità di archiviazione e strumenti software accessibili su larga scala. Nel livello successivo ci sono i dati, che comprendono sorgenti da cui si estrapolano dati grezzi o strutturati per addestrare i modelli. I due livelli più a valle sono rispettivamente i *foundation models*, che sono modelli addestrati su grandi quantità di dati per svolgere compiti più o meno generici, e le applicazioni, che traducono tali capacità in servizi concreti per gli utenti.

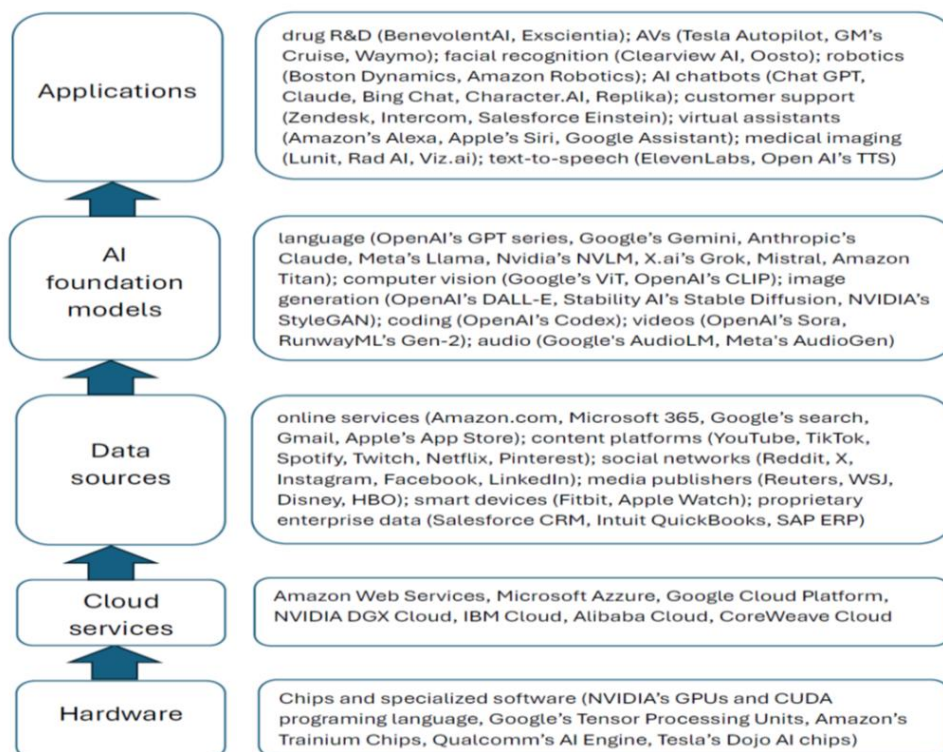


Figura 1: Rappresentazione dei livelli dell'AI stack.

I livelli non devono essere considerati come compartimenti stagni; infatti, lo *stack* dell'AI si distingue dalle filiere produttive tradizionali perché le performance di ciascun livello dipendono strettamente dall'evoluzione degli altri.

Questa dinamica rende evidente una differenza sostanziale rispetto alle catene di valore lineari. Se si considera un settore industriale classico, infatti, il processo produttivo segue una sequenza relativamente rigida, in cui il miglioramento a monte non sempre si traduce in benefici immediati a valle. Nell'*AI stack*, invece, i diversi livelli sono connessi da *feedback loops*, che permettono che miglioramenti nei chip rendano possibile addestrare modelli più complessi, ma allo stesso tempo modelli migliori generino applicazioni più performanti, che stimolano ulteriormente la domanda di calcolo e la raccolta di dati. Si innesca così un circolo di retroazioni che può accelerare l'innovazione e la crescita, ma anche rafforzare la concentrazione del potere economico esistente.

Inoltre, l'interdipendenza emerge chiaramente con la presenza di attori trasversali come gli *hyperscaler*, grandi imprese che operano nei servizi cloud tramite vaste reti di data center distribuiti, come Amazon Web Services, Microsoft Azure e Google Cloud capaci di operare contemporaneamente su più livelli dello *stack*. Infatti, oltre a fornire infrastrutture di calcolo hanno progressivamente ampliato la loro presenza anche in segmenti, come la gestione dei dati, lo sviluppo di *foundation models* e la creazione di applicazioni.

Questo dimostra come questa struttura non sia semplicemente una rappresentazione tecnologica, ma anche una mappa delle interdipendenze economiche e concorrenziali in cui si concentrano risorse e leve strategiche.

2.1.2 Funzioni concorrenziali e punti nevralgici dello stack

Lo *stack*, oltre a rappresentare un'architettura tecnologica, costituisce un vero e proprio campo di potere competitivo, dove chi controlla gli input critici condiziona l'accesso, l'innovazione e la possibilità stessa di competere a valle.

I colli di bottiglia principali si trovano nella potenza di calcolo, nei dati necessari ad addestrare i modelli e nelle API, che regolano l'interazione tra *foundation models* (FM) e le applicazioni finali.

Questi elementi possono essere considerati veri e propri *control points*, presidiati in larga misura dagli *hyperscaler*. Grazie alle economie di scala e di scopo derivanti dalle loro infrastrutture cloud, questi attori non offrono soltanto infrastrutture con capacità computazionale (IaaS), ma anche piattaforme che presentano ambienti e strumenti per addestrare modelli (PaaS) e servizi e modelli pre-addestrati (SaaS). In questo modo riescono a

inserirsi trasversalmente nello *stack*, sia sviluppando modelli proprietari, sia controllando indirettamente l'accesso di terzi.

I *foundation models* svolgono un ruolo centrale in questa dinamica, poiché trasformano gli input critici, calcolo e dati, in competenze trasversali in grado di svolgere compiti differenti, che alimentano le applicazioni a valle.

Questo genera forti effetti di rete, poiché più applicazioni vengono sviluppate, più dati vengono prodotti e reinseriti nel ciclo di addestramento dei modelli, alimentando così un circolo di rafforzamento, che consolida ulteriormente il vantaggio degli *incumbent*. Di conseguenza, gli FM assumono la funzione di piattaforme intermedie e di vere e proprie *general-purpose technologies*, con rischi di *lock-in* e di una riduzione della contestabilità del mercato.

2.2 Struttura dell'AI Stack

2.2.1 Hardware e infrastrutture di calcolo

Il primo livello della catena è rappresentato dall'hardware, costituito dai server che, grazie ai chip specializzati, rendono disponibile l'elevata potenza di calcolo necessaria per lo sviluppo dell'intelligenza artificiale e in particolare dei *foundation models*. Questo livello rappresenta un collo di bottiglia, a causa della scarsità di componenti e della localizzazione geografica della produzione.

I chip più utilizzati che permettono di velocizzare i calcoli, addestrare i *foundation models* e operare il *fine-tuning* e l'inferenza degli stessi, sono le GPU (*Graphics Processing Units*), nate per l'elaborazione grafica, ma diventate cruciali per l'addestramento nel *deep learning*, grazie alla loro architettura parallela che le rende due volte più efficienti delle CPU.

Nvidia è l'azienda leader in questo campo, e in particolare si stima che nel 2024 detenesse oltre il 90% del mercato delle GPU per AI. Il vantaggio competitivo che ha creato negli anni non deriva soltanto dall'hardware, ma anche dall'ecosistema software; infatti, dà la possibilità di usufruire del linguaggio di programmazione e della piattaforma CUDA (*Compute Unified Device Architecture*), che è diventato, ormai, lo standard industriale per sfruttare al meglio le GPU e i carichi di lavoro dovuti ai compiti svolti dall'intelligenza artificiale. Questo genera forti effetti di rete poiché gli sviluppatori tendono ad utilizzare con maggiore frequenza CUDA, che a sua volta interagisce con librerie e framework. Per esempio, PyTorch, che è un framework di *deep learning* open-source sviluppato inizialmente da Facebook (Meta), e TensorFlow, un framework di *machine learning* open-source sviluppato da Google. Questo meccanismo comporta che i framework citati vengano ottimizzati sempre di più per le GPU Nvidia, riducendo così gli incentivi a passare a soluzioni alternative. Inoltre, Nvidia aggiorna

costantemente CUDA per integrarlo con il proprio hardware, rafforzando il *lock-in* e riducendo la compatibilità con chip rivali.

Un altro elemento cruciale è che i principali acquirenti delle GPU di Nvidia sono proprio gli *hyperscaler* come Amazon, Microsoft e Google che utilizzano queste schede per i loro data center. Però nel lungo periodo questa dipendenza da Nvidia potrebbe essere destinata a ridursi; infatti, stanno sviluppando parallelamente chip proprietari, come si nota nella *Tabella 1*, per ridurre i costi e aumentare l'autonomia. Nello specifico Google ha introdotto le proprie TPU (*Tensor Processing Units*), progettate appositamente per il *deep learning* e impiegate per addestrare i suoi *foundation models*, tra cui Gemini. Amazon ha sviluppato Trainium e Inferentia, oltre a Neuron, che permette di passare da chip Nvidia a chip di terze parti. Mentre Microsoft ha presentato il chip Maia. Anche produttori come Intel e AMD o Big Tech come Meta e Apple hanno annunciato nuovi acceleratori dedicati all'AI, oltre ai chip AI Engine di Qualcomm, Dojo AI chips di Tesla e quelli di Open AI. Alcune startup, invece, stanno puntando su architetture innovative, puntando su costi minori e risparmio energetico, come Cerebras Systems, SambaNova Systems, negli acceleratori specializzati; Lightmatter, Luminous Computing nel *computing* ottico; BrainChip, nel neuromorphic computing; Groq nel *language Processing Units* fino al *quantum computing* con PsiQuantum e IonQ. Una o più di queste tecnologie in futuro potrebbero rappresentare una rottura radicale rispetto alle architetture tradizionali.

Un altro elemento da considerare è il recente spostamento di enfasi dal training iniziale, fortemente intensivo in termini di calcolo, verso un maggiore impiego di risorse computazionali nella fase di inferenza. Questa novità ha aperto nuove opportunità per la progettazione di chip potenzialmente più efficienti di quelli Nvidia in tale aspetto, come dimostrano le iniziative di Amazon, Google e altri attori.

Nonostante queste alternative, una parte significativa della domanda di calcolo accelerato continua a confluire verso le GPU dell'azienda leader del settore, tanto che i suoi modelli più avanzati hanno raggiunto prezzi elevatissimi e tempi di attesa fino a dodici mesi, poi ridotti a tre mesi, mettendo in risalto il collo di bottiglia rappresentato da questo livello della catena nell'ecosistema AI.

Questi colli di bottiglia sono aggravati dal fatto che la catena di fornitura dei semiconduttori è fortemente concentrata, dal momento che la quasi totalità dei chip di fascia alta di Nvidia viene prodotta da TSMC (*Taiwan Semiconductor Manufacturing Company*), situato a Taiwan, e da Samsung Electronics, che ha sede in Corea del Sud e ricopre un ruolo minore. Mentre ASML,

situata nei Paesi Bassi è l'unico produttore al mondo delle macchine per litografia EUV necessarie per la fabbricazione dei chip.

AZIENDE	CHIP
	GPU 90% del mercato
	TPU per training AI / integrazione in Google Cloud
	Trainium e Inferentia per AWS Bedrock e servizi cloud AI
	Meta Training and Inference Accelerator (MTIA) per addestramento modelli interni
	Maia per Azure cloud

Tabella 1: Stato attuale e dinamico della concorrenza sui chip.

Questa concentrazione rende l'accesso al calcolo avanzato vulnerabile a shock geopolitici o restrizioni commerciali. Per questo motivo governi come Stati Uniti e Unione Europea hanno avviato politiche industriali per sostenere la produzione domestica di semiconduttori e per ridurre la dipendenza da paesi terzi. Al contempo sono stati introdotti controlli alle esportazioni, per esempio a partire dal 2022 gli Stati Uniti hanno imposto limiti alla vendita di chip avanzati verso la Cina, nel tentativo di rallentare l'avanzamento tecnologico di un potenziale rivale strategico.

Nel complesso, l'hardware costituisce quindi il primo grande control point dell'*AI stack*, dove Nvidia detiene una posizione dominante e potrebbe ostacolare i nuovi entranti tramite pratiche commerciali, come contratti di esclusiva o sconti fedeltà chiamati anche *loyalty rebates*, mentre gli *hyperscaler* cercano di ridurre la propria dipendenza sviluppando chip proprietari.

Il futuro del settore dipenderà dall'equilibrio fra la capacità dei concorrenti di erodere il primato di Nvidia e il mantenimento della concentrazione attuale, che rappresenta una barriera all'ingresso significativa per la crescita di un ecosistema AI realmente competitivo.

2.2.2 Servizi cloud

Il secondo livello dell'*AI stack* è rappresentato dal settore dei servizi cloud, che forniscono la capacità computazionale, lo storage e gli strumenti software indispensabili per l'addestramento, il *fine-tuning* e l'inferenza dei *foundation models*. Questi servizi permettono alle imprese di accedere a risorse scalabili on-demand, senza dover sostenere gli ingenti investimenti necessari per costruire e mantenere data center dedicati e privati, riducendo così i costi fissi e le barriere tecnologiche all'ingresso negli stadi più a valle.

Il cloud rappresenta a tutti gli effetti un input critico per lo sviluppo dell'AI, al pari dell'hardware e dei dati. Infatti, l'accesso ai servizi dei grandi provider costituisce un collo di bottiglia per le imprese che intendono sviluppare *foundation models*.

L'offerta di servizi cloud si articola tipicamente su tre modalità. La prima è l'*Infrastructure-as-a-Service* (IaaS), che consente di noleggiare la capacità computazionale ad alte prestazioni, come GPU o TPU, essenziali per il training dei modelli. La seconda modalità è *Platform-as-a-Service* (PaaS), che offre ambienti integrati e strumenti per preparare i dati, addestrare e distribuire modelli, come Amazon SageMaker o Google Vertex AI. Infine, l'ultima è *Software-as-a-Service* (SaaS), che permette di accedere a modelli già addestrati o a interfacce di programmazione, le cosiddette API, che le imprese possono integrare direttamente nelle proprie applicazioni, senza il bisogno di svilupparle da zero. Questa diversità di servizi offerti mostra come il cloud non sia soltanto l'infrastruttura in sé, ma un ecosistema completo, in grado di coprire e interagire in diversi stadi della catena del valore dell'intelligenza artificiale.

In questo senso, i provider cloud fungono da vero e proprio ponte tra i produttori di hardware, come Nvidia e Google, e gli sviluppatori di modelli, mediando l'accesso alle risorse computazionali attraverso le modalità di servizio appena viste.

Il mercato dei servizi cloud risulta altamente concentrato attorno a tre grandi *hyperscaler* che comprendono Amazon Web Services (AWS), Microsoft Azure e Google Cloud, che insieme detengono circa i due terzi della capacità mondiale di cloud computing. Ognuno di essi integra i servizi cloud con altri asset strategici nello *stack*. Infatti, AWS, leader globale con circa un terzo del mercato, presenta chip proprietari come Trainium e Inferentia per training e inferenza. Allo stesso tempo Microsoft ha consolidato la propria posizione grazie alla partnership con OpenAI, di cui distribuisce i modelli in esclusiva tramite il proprio servizio cloud Azure. Google, invece, utilizza le proprie TPU e integra i *foundation models* Gemini nella propria piattaforma Vertex AI.

Accanto ai tre leader globali, come indicati nella *tabella 2*, si muovono altri importanti provider, tra cui Alibaba Cloud, che è il principale attore asiatico, Oracle Cloud, che si è posizionato

come alternativa per le startup del settore grazie alle GPU fornite da Nvidia a seguito di una partnership tra le aziende. IBM Cloud, invece, è un attore più piccolo, che non compete direttamente nel training dei *foundation models* su larga scala, ma è orientato alle applicazioni *enterprise*, fornendo soluzioni mirate alle imprese in settori regolati come sanità, finanza e pubblica amministrazione, aggiungendo alla sua offerta alti standard di sicurezza e affidabilità. Parallelamente, si sono affermati operatori specializzati come CoreWeave, Lambda Labs e Denvr AI, che offrono infrastrutture dedicate al calcolo per l’AI. Questi provider possono garantire un miglior rapporto qualità-prezzo, offrire maggiore autonomia agli sviluppatori di AI e permettere di gestire direttamente i propri carichi di lavoro, senza dover dipendere dai servizi costosi e più generici dei grandi fornitori di cloud. Inoltre, queste aziende, pur dipendenti dalle tecnologie dei produttori dominanti, hanno attratto investimenti diretti dalle stesse Big Tech, come per esempio Nvidia che ha investito in CoreWeave, a conferma della forte interconnessione tra i diversi livelli dello *stack*. Anche attori non tradizionali sfruttano le proprie conoscenze sviluppate in altri settori per entrare in questo segmento, come Tesla con il supercomputer Dojo AI.

Livello	Concorrenza attuale	Concorrenza dinamica
Cloud AI / IaaS	AWS, Azure, Google Cloud (2/3 del mercato combinato)	CoreWeave, Lambda, Denvr AI, Alibaba Cloud, Oracle Cloud, IBM Cloud

Tabella 2: Concorrenza attuale e dinamica nel settore cloud.

L’accesso al calcolo tramite cloud ha importanti implicazioni concorrenziali. Infatti, gli *hyperscaler* non si limitano a fornire solo l’infrastruttura, ma integrano nei propri ecosistemi anche i *foundation models*. Microsoft distribuisce GPT tramite Azure, Google Cloud rende disponibile il proprio modello Gemini, Amazon ospita Anthropic. Questa integrazione verticale rafforza i rischi di auto-preferenza, poiché i fornitori di cloud possono favorire i propri modelli o partner strategici a scapito di concorrenti esterni. Inoltre, la scelta di sviluppare e addestrare modelli su un determinato cloud comporta costi di *switching* elevati dovuti alla migrazione dei dati, alla riconfigurazione delle pipeline e alla dipendenza da strumenti proprietari, che generano effetti di *lock-in* che riducono la contestabilità del mercato. Infatti, oltre a migrare

enormi volumi di dati, bisognerebbe riscrivere parte del codice e rinunciare a ottimizzazioni fatte su misura. In questo contesto, assume una rilevanza importante la strategia degli *hyperscaler*, che offrono spesso crediti cloud a start-up e sviluppatori di AI, ma impongono vincoli contrattuali, come le elevate *egress fees*, costi per l'estrazione e il trasferimento dei dati, e il *bundling* di servizi PaaS. Questi meccanismi rafforzano, ancora di più gli effetti di *lock-in*, rendendo difficile per gli utenti migrare verso fornitori alternativi.

Un'ulteriore sfida riguarda la sostenibilità. I data center richiesti dall'intelligenza artificiale consumano enormi quantità di energia: nel 2024 gli investimenti globali in data center hanno superato i 465 miliardi di dollari, e parte consistente è legata alla crescita dell'AI. Tuttavia, barriere regolatorie, limiti infrastrutturali e vincoli energetici rischiano di rallentare l'espansione, avvantaggiando i player già presenti. Questo rafforza ulteriormente il potere dei principali attori, che dispongono delle risorse necessarie per investire in nuove infrastrutture e in energie rinnovabili.

Infine, il cloud è diventato anche un tema di politica industriale e geopolitica. L'Unione Europea ha promosso iniziative come GAIA-X per sviluppare un cloud europeo, mentre Stati Uniti e Cina sostengono le proprie imprese nazionali, AWS e Azure da un lato, Alibaba e Baidu dall'altro. La concentrazione del mercato globale, unita alla rilevanza strategica dei data center, trasforma quindi i servizi cloud in un *control point* cruciale per l'intero ecosistema dell'IA.

I cloud stanno diventando anche dei veri e propri canali di distribuzione per i *foundation models*, tramite marketplace integrati come i cosiddetti Model Gardens, per esempio Google Model Garden, Azure AI Model Catalog e AWS Bedrock. Questo rafforza il loro ruolo strategico, poiché controllano non solo l'infrastruttura, ma anche l'accesso ai modelli per gli sviluppatori e le imprese.

In netto contrasto con i livelli dei dati o dei modelli, il livello dei servizi cloud è invece destinato a una continua concentrazione attorno ai tre attori dominanti. La ragione principale è che i servizi cloud comportano sostanziali economie di scala, soprattutto nel contesto dell'AI. Infatti, i principali provider cloud hanno investito miliardi di dollari nella costruzione di enormi cluster di GPU, che gli sviluppatori di modelli AI possono noleggiare.

2.2.3 Dati

Il livello successivo della catena dal valore è rappresentato dai dati, che sono fondamentali per l'addestramento, il *fine-tuning* e l'impiego dei *foundation models*. Nel loro complesso i dati, oltre ad essere usati da un punto di vista tecnico, assumono rilevanza anche dal punto di vista economico, giuridico e regolatorio, poiché il loro utilizzo è soggetto ai vincoli di proprietà intellettuale e alle norme sulla privacy. Questo input può determinare chi può e chi non può

sviluppare le tecnologie necessarie allo sviluppo dei sistemi avanzati da utilizzare a supporto dell'intelligenza artificiale, diventando di conseguenza un punto fondamentale di controllo dell'intera catena.

Le tipologie di dati utilizzati dall'IA individuate dalla Competition and Markets Authority (CMA, 2023) e dall'Italian Competition Authority, AGCM, (Discussion at the G7 Competition Summit, 2024) sono suddivise in quattro grandi categorie.

La prima categoria è quella dei dati pubblici, che comprendono quei contenuti disponibili sul web e che vengono raccolti tramite *web scraping*, pratica che consiste nell'estrarre automaticamente i contenuti disponibili online. Alcuni esempi di questi dataset comprendono Wikipedia, Common Crawl e GitHub, *repository* di codice sorgente. Ulteriori esempi includono dataset open-source come The Pile o LAION, usato per l'addestramento multimodale, l'Internet Archive e forum tecnici come Stack Exchange, che rappresentano risorse significative, ma soggette a crescenti restrizioni di accesso. Questa tipologia di dati è fondamentale in fase di *pre-training*, in cui la varietà delle fonti è cruciale.

La seconda categoria include i dati proprietari o licenziati, che permettono di accedere a contenuti di alta qualità o specializzati, andando a rappresentare una risorsa centrale e di valore aggiunto in fase di addestramento dei modelli. Molti di questi comprendono grandi archivi controllati direttamente dalle Big Tech, come le *query* di ricerca di Google Search, e i contenuti video di YouTube per Google, o i dati aziendali di Microsoft 365 e le reti professionali di LinkedIn per Microsoft, oltre agli app store e ai *social graph* di piattaforme come Facebook e Instagram. Questi rappresentano dataset unici normalmente inaccessibili ai nuovi entranti se non tramite costosi accordi di licenza. Inoltre, negli ultimi anni si sono instaurate alcune sinergie commerciali fra sviluppatori di modelli e titolari di grandi archivi digitali. Per esempio, la partnership tra OpenAI e Reddit consente di utilizzare i contenuti generati dalle community online, invece quella tra Google e Stack Overflow, è finalizzata a migliorare i modelli Gemini con dati tecnici e conversazioni di programmazione. Un altro accordo in questo campo è tra Meta e Shutterstock, che permette di integrare immagini e video di alta qualità nell'addestramento dei modelli multimodali. I dataset delle due categorie sono riassunti nella *tabella 3*.

Tabella 3: Stato e accordi dei dataset disponibili per l'addestramento degli FM.

Tipologia di dato	Dataset / Fonte	Pubblico / Privato / Accordo Big Tech	Note / Utilizzo principale
Testo	Wikipedia	Pubblico	Pre-training, ampia varietà di contenuti generali
	Common Crawl	Pubblico	Pre-training, dati web generici
	GitHub	Pubblico	Codice sorgente per modelli di programmazione
	The Pile	Pubblico / open-source	Pre-training multimodale, dati di codice, articoli scientifici, libri
	LAION	Pubblico / open-source	Pre-training multimodale
	Stack Exchange	Pubblico	Dati tecnici, Q&A per modelli di programmazione
	Reddit	Privato / accordo con OpenAI	Miglioramento conversazioni e interazioni social AI
	Stack Overflow	Privato / accordo con Google	Miglioramento modelli Gemini con dati tecnici
Video	YouTube	Privato / dati proprietari di Google	Addestramento di modelli video multimodali
	Shutterstock	Privato / accordo con Meta	Integrazione immagini e video di alta qualità per modelli multimodali
Immagini	Internet Archive	Pubblico	Pre-training, immagini storiche e artistiche
	Shutterstock	Privato / accordo con Meta	Dataset di immagini per modelli multimodali
	LAION	Pubblico / open-source	Immagini per training multimodale

Le ultime due categorie sono i dati sintetici e quelli da feedback. I primi sono generati artificialmente anche dallo stesso FM, per arricchire la profondità dei dataset disponibili, soprattutto in contesti in cui i dati di qualità sono scarsi. Tuttavia, questa pratica rischia di portare all'effetto opposto a causa della perdita di qualità nei dati e all'aumento dei rischi legati al *model collapse*: un addestramento concentrato prevalentemente su output sintetici porta a una riduzione della diversità informativa e amplifica eventuali errori o bias, compromettendo sia la generalizzazione del sistema, sia la rappresentazione della realtà da parte del modello.

I dati da feedback sono dinamici e vengono raccolti a valle della catena, infatti, derivano direttamente dalle interazioni degli utenti con le applicazioni di intelligenza artificiale. Questo meccanismo porta a migliorare le performance future del prodotto, creando un *feedback loop*, nel quale più utenti interagiscono con il modello, più dati vengono raccolti, più si migliora il modello, maggiore è la base di utenti che lo utilizzano. Tale dinamica favorisce le aziende *incumbent* che hanno già acquisito questi vantaggi cumulativi e dispongono di basi utente ampie.

I dati generano asimmetrie e colli di bottiglia, poiché sono abbondanti, ma esiste un problema di qualità del dato. Infatti, le grandi imprese tecnologiche possiedono le risorse finanziarie e il potere contrattuale necessario per accedere a dataset esclusivi o stipulare accordi con grandi editori o piattaforme digitali, invece i nuovi entranti o le startup si accontentano spesso di usufruire di dataset pubblici che sono meno curati o incompleti, generando asimmetrie informative.

A questa dinamica si aggiungono le questioni legali che possono generare colli di bottiglia all'interno della catena: numerosi processi giudiziari sono in atto a causa della pratica dello *scraping* che violerebbe il copyright sull'utilizzo di alcuni contenuti. Alla luce di questo sempre più piattaforme iniziano ad adottare il blocco dei *crawler*, che consiste nel limitare l'accesso automatizzato ai contenuti di un sito web, impedendo che vengano indicizzati o utilizzati da bot. Per mettere in pratica questo freno ci si serve di barriere tecniche o legali, con la conseguente riduzione della disponibilità di dati per i nuovi attori. In aggiunta l'applicazione europea del GDPR e in futuro del AI Act pone dei limiti stringenti per quanto riguarda l'uso dei dati personali.

Queste evoluzioni porterebbero un vantaggio alle imprese già presenti nell'utilizzo di archivi esclusivi tramite licenza, andando a minare la competitività sull'AI, poiché per i nuovi entranti diventerebbe difficile replicarne la qualità. Anche a causa dei *feedback loop* che rafforzano ulteriormente la posizione di dominanza delle principali aziende. Questo processo accentua gli effetti *winner-takes-all* e riduce la contestabilità del mercato. Quindi bisognerà valutare se i dataset aperti e i cosiddetti *commons* saranno sufficienti per addestrare *foundation models* competitivi, o se l'ecosistema si sposterà sempre di più verso l'uso esclusivo di dati proprietari, con conseguente rafforzamento del potere degli incumbent.

L'utilizzo dei dati sintetici potrebbe rappresentare una soluzione a queste barriere, senza però trascurare i rischi che potrebbero esserci a livello di errori pregressi e bias, con perdita di qualità del modello addestrato artificialmente.

Dal punto di vista geopolitico si hanno tre situazioni differenti in tre macroaree geografiche e politiche.

Negli Stati Uniti si nota una propensione alle partnership commerciali tra Big Tech e fornitori di contenuti. In Cina vengono sfruttati la larga quantità di dati nazionali a disposizione, congiunta a regole meno restrittive sulla privacy che offrono un vantaggio competitivo significativo. In Europa, invece, le norme in atto sui dati potrebbero frenare l'utilizzo degli stessi per l'addestramento dei modelli, portando a una perdita di competitività a livello globale. Quindi sarà fondamentale la capacità di garantire un accesso equo e sostenibile ai dataset poiché diventerà una delle variabili decisive per la contestabilità dei mercati e per la possibilità di sviluppare un ecosistema AI realmente competitivo.

2.2.4 Foundation Models

I *foundation models* (FM) rappresentano il cuore dell'*AI stack*, diventando di fatto il principale nodo della catena del valore. Rispetto ai sistemi di intelligenza artificiale tradizionali sono progettati non per un compito specifico, ma per apprendere rappresentazioni generali, partendo da dataset vasti e multimodali, che comprendono immagini, audio, video e codici. Questo addestramento è stato reso possibile grazie alla tecnica del *deep learning*, affinando successivamente l'ambito dell'applicazione del modello con procedure di *fine-tuning*. La loro natura "generalista", li rende assimilabili a una tecnologia di uso generico, *general purpose technology*, con effetti trasversali che toccano molti settori dell'economia e della società.

Il ciclo vita dei FM comprende una pipeline composta da tre stadi. La prima fase è di pre-training, dove il modello è addestrato su una quantità vasta di dati eterogenei, con rappresentazioni di base del linguaggio, delle immagini e di altre modalità. La seconda è il fine-tuning, nella quale le conoscenze acquisite nello stadio precedente vengono raffinate attraverso dataset più curati e specifici per un certo dominio, che consente la specializzazione del modello in un ambito ristretto, per esempio nel settore sanitario, legale o finanziario. L'ultimo stadio è il deployment, in cui viene reso disponibile all'utilizzo, tramite API fornite da un provider esterno, oppure mantenendolo su un'infrastruttura controllata direttamente tramite self-hosting. Quest'ultima è un'opzione percorribile solo da attori dotati di ingenti risorse.

La pipeline, appena descritta, evidenzia ancora di più la centralità dei FM nello *stack*: trasformano gli input critici, potenza di calcolo e dati, in competenze trasversali utilizzate nelle applicazioni e servizi a valle, diventando un punto di controllo della catena.

Gli FM si distinguono per scopo e modalità d'uso in diverse tipologie, rendendo il panorama di questo stadio della supply chain non uniforme, ma articolato ed eterogeneo. I primi sono i general FM che sono addestrati per essere usati in contesti molteplici, con rendimenti crescenti

a causa della loro versatilità, alcuni esempi sono la serie GPT di OpenAI che nel 2023 deteneva il 76% del market share dei principali chatbot commerciali, Gemini di Google, LLaMA di Meta. Al contrario i network FM sono pensati su particolari domini o piattaforme, ad esempio modelli linguistici specializzati per il coding o per la sicurezza informatica. Infine, stanno emergendo i personal FM, progettati per adattarsi ad un singolo individuo o a una singola organizzazione, integrando al suo interno dati personali e informazioni che variano in base al contesto di applicazione.

Per sviluppare un FM è necessaria una grande quantità di potenza calcolo: centinaia o migliaia, a seconda dei casi, di GPU o TPU utilizzate in parallelo per settimane o mesi, comportando milioni di costi per realizzarlo. Inoltre, si ha una spesa considerevole anche per quanto riguarda l'inferenza, che consiste nell'uso quotidiano del modello, relativo alle risorse computazionali richieste. Per ovviare a questi vincoli, si stanno diffondendo strategie e tecniche di ottimizzazione che hanno l'obiettivo di migliorare l'efficienza computazionale dei modelli, riducendo il consumo di memoria e i costi associati all'inferenza, così da rendere più sostenibile l'impiego dei *foundation models* su larga scala. Allo stesso tempo, la tendenza è privilegiare la qualità più che la quantità dei dati. Infatti, non è più sufficiente avere dataset ampi, ma è necessario disporre di dati curati e *domain-specific*, combinati con tecniche di *reinforcement learning* basate sul feedback: RLHF (*Reinforcement Learning with Human Feedback*) progettata per affinare i modelli usando valutazioni umane per orientare le risposte verso quelle più utili e sicure, e RLAIIF (*Reinforcement Learning from AI Feedback*) che rappresenta una variante più scalabile, in cui sono altri modelli AI a fornire i riscontri al posto degli umani.

Un'ulteriore considerazione sui modelli è la multimodalità, essendo che non si limitano solo al testo scritto, ma integrano diverse modalità di input e output, come linguaggio, immagini, audio, video e codice, ampliando le capacità di generalizzazione e rappresentazione. Questa caratteristica consente di affrontare compiti complessi e di combinare diverse tipologie di informazione in un unico modello. Esempi rilevanti sono Gemini di Google, progettato fin dall'origine come architettura multimodale, e GPT-5 di OpenAI, che integra testo e immagini. Questi progressi rafforzano ulteriormente il ruolo dei FM come piattaforme generali, in grado di abilitare applicazioni sempre più diversificate nei livelli a valle dello *stack*.

Un aspetto cruciale per analizzare la catena è il grado di apertura dei FM. La distinzione tra aperti o chiusi non è una distinzione binaria, ma uno spettro di possibilità. Sono presenti modelli chiusi come GPT-5 (OpenAI), Gemini (Google) e Claude (Anthropic) che sono accessibili solo tramite API e mantengono un controllo stretto sull'uso e sulla distribuzione. Invece all'estremo opposto si colloca BLOOM, progetto guidato e coordinato da BigScience e Hugging Face con

la collaborazione internazionale di oltre mille ricercatori, che ha reso disponibili pesi, dataset e documentazione completa, rappresentando a tutti gli effetti un modello open che favorisce la ricerca anche in ambito accademico e la concorrenza nel mercato. Nel mezzo si trovano numerose soluzioni ibride, ognuno con un proprio grado di apertura. Il modello LLaMA di Meta presenta pesi disponibili ed è usabile da ricercatori e sviluppatori, però i dataset su cui è addestrato non sono pubblici e la licenza che concede limita gli usi commerciali. Anche di Falcon, sviluppato da TII ad Abu Dhabi, sono stati rilasciati i pesi con l'aggiunta della documentazione tecnica, ma anche in questo caso i dataset sono parziali o disponibili solo a condizioni di licenza restrittive. Mistral, sviluppata da una start-up francese, rilascia pesi e parte del codice, ma non sempre i dataset, e spesso utilizza licenze che limitano l'uso commerciale. Queste aperture parziali favoriscono la ricerca, ma non eliminano le barriere all'ingresso.

Dal punto di vista della concorrenza di mercato i modelli chiusi tendono a rafforzare il *lock-in* e l'integrazione verticale, mentre i modelli aperti favoriscono la democratizzazione dell'accesso, ma espongono a rischi di sicurezza. Inoltre, negli ultimi anni si è osservata la tendenza alla traiettoria “*open early, closed late*”: nelle prime fasi dell'implementazione del modello è presente un'apertura per stimolare la diffusione e formare una comunità di sviluppatori, seguita nelle fasi successive da restrizioni sempre più severe.

Accanto ai *large language models* (LLM) di cui finora si è parlato, stanno emergendo modelli più piccoli e leggeri dal punto di vista dell'addestramento, i cosiddetti *small language models* (SLM). Questi hanno requisiti computazionali ridotti e possono essere impiegati in contesti specifici o integrati in dispositivi mobili. L'esistenza di questi modelli “*good enough*” dimostra che si può competere anche con strategie differenti rispetto alla corsa della dimensione e varietà. Infatti, l'efficientamento e l'uso di dataset molto curati consentono a modelli più piccoli di competere con quelli più grandi in domini verticali, cioè in contesti specialistici. La coesistenza di modelli grandi e piccoli aumenta la varietà dell'offerta e riduce il rischio di un mercato dominato da un unico paradigma.

Un altro elemento fondamentale in questo segmento sono i talenti. L'addestramento e la gestione dei FM richiedono competenze avanzate in *machine learning*, *data science*, ingegneria del software e *AI safety*. Questi profili sono rari e si concentrano soprattutto presso i grandi *incumbent*, che possono attrarli con salari elevati e progetti di ricerca all'avanguardia. Questa scarsità di capitale umano qualificato rappresenta una barriera all'ingresso tanto quanto la mancanza di dati o di capacità computazionale. Allo stesso tempo, la diffusione di strumenti più accessibili e modelli open source riduce, almeno in parte, la dipendenza da grandi team altamente specializzati, abbassando le barriere per lo sviluppo di modelli più piccoli o verticali.

Un ulteriore punto critico che coinvolge anche i modelli sono le API attraverso cui i FM vengono resi disponibili, diventando non solo un'interfaccia tecnica, ma un vero e proprio strumento di controllo economico e strategico. Sfruttando le API i provider possono stabilire chi accede, a quali condizioni e con quali restrizioni. Inoltre, raccolgono i dati di utilizzo che possono essere reimpiegati per migliorare ulteriormente i modelli, alimentando il *feedback loop*.

Quest'ultime possono anche fungere da leva competitiva, consentendo strategie di auto-preferenza o di esclusione.

Nella catena del valore degli FM si possono distinguere due figure principali: gli sviluppatori *upstream*, che addestrano i modelli di base, e i *deployer downstream*, che li adattano, li distribuiscono e li integrano in applicazioni. Questo rapporto porta a una competizione tra i *deployer* sulla personalizzazione e monetizzazione di un asset comune fornito dagli sviluppatori, chiamata *fine-tuning game*.

Il mercato dei FM ha caratteristiche assimilabili a un monopolio naturale. Gli altissimi costi fissi di addestramento e i costi marginali decrescenti dell'inferenza producono rendimenti di scala crescenti: più un modello viene utilizzato, più diventa profittevole. Questo rafforza la posizione dei leader e scoraggia l'ingresso di nuovi attori.

Per questo motivo il segmento rappresentato dagli FM è ad oggi altamente concentrato ed è caratterizzato da ingenti investimenti privati come dimostra la *figura 2* tratta dall'AI Index Report 2024.

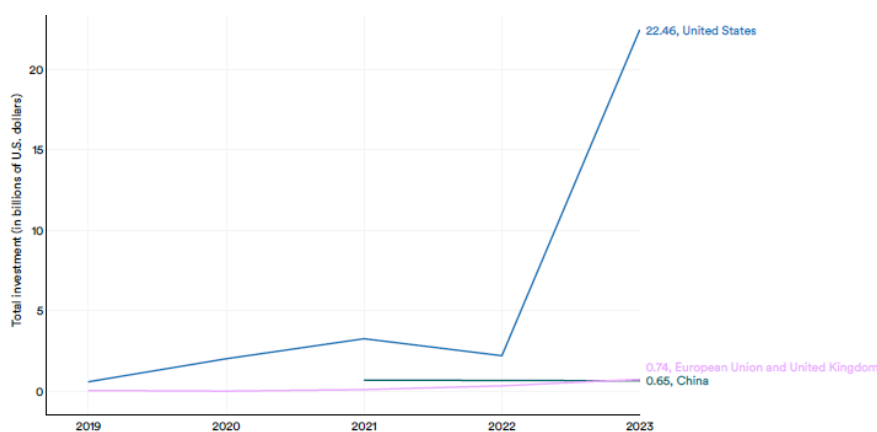


Figure 4.3.11

Figura 2: Investimenti privati in AI generativa divisa per aree geografiche

Negli Stati Uniti, i principali attori sono cinque: OpenAI, Google, Anthropic, Meta e Amazon. Open AI, sostenuta finanziariamente e strategicamente da Microsoft, ha aperto la strada con la serie GPT. La partnership con Microsoft ha garantito non solo capitali, ma anche l'integrazione

esclusiva dei modelli nei servizi cloud di Azure e nelle applicazioni downstream, come Copilot, il motore di ricerca Bing e il browser Edge, creando un ecosistema fortemente integrato.

Google ha sviluppato Gemini, grazie alla capacità di combinare modelli multimodali e infrastruttura proprietaria con il supporto dei propri chip. Inoltre, l'azienda fa leva sull'integrazione del proprio modello nei servizi di largo consumo a valle di sua proprietà, come il suo famoso motore di ricerca e i propri Workspace (Gmail, Meet, Drive ecc.), con un impatto immediato su miliardi di utenti.

Anthropic, fondata da ex ricercatori di OpenAI, si è distinta per l'attenzione alla sicurezza e all'allineamento dei modelli anche dal punto di vista etico, con la serie Claude, distribuita attraverso Amazon Web Services (AWS), che ha a sua volta investito in modo importante nell'azienda.

Meta ha intrapreso, come anticipato precedentemente, una strategia più aperta con LLaMA rilasciata in forma accessibile alla comunità scientifica e agli sviluppatori, anche se l'impresa mantiene il controllo interno sulle versioni più avanzate.

Infine, Amazon oltre a distribuire modelli di terzi tramite la piattaforma Bedrock, ha investito, come già visto, in chip, cloud, infrastruttura e servizi di intelligenza artificiale, rafforzando la propria posizione grazie anche agli accordi con Anthropic.

In Cina, lo sviluppo è sostenuto da una combinazione di politiche industriali e di accesso privilegiato a vasti bacini di dati nazionali. In questo contesto Baidu, il motore di ricerca per eccellenza cinese, ha sviluppato la serie Ernie, Alibaba ha introdotto il modello Tongyi Qianwen, mentre Huawei e Tencent hanno investito sia nello sviluppo di FM sia nella creazione di infrastrutture cloud e chip proprietari. In aggiunta ByteDance, la società madre di TikTok, possiede la capacità di sfruttare dati comportamentali e modelli multimodali in applicazioni legate ai contenuti digitali e all'intrattenimento, potendo avere in futuro la possibilità di sviluppare FM proprietari o in accordo con altri attori.

Un caso rilevante all'interno dello sviluppo dei FM cinesi è DeepSeek, che dimostra come l'innovazione non debba necessariamente passare da un incremento illimitato delle risorse computazionali, ma utilizzando dataset molto curati e tecniche di ragionamento strutturate, come la *chain-of-thought*, è possibile raggiungere performance competitive con un fabbisogno di calcolo ridotto rispetto ai concorrenti occidentali. Questo approccio suggerisce che la competizione tra USA e Cina non si giochi soltanto sulla disponibilità di GPU e infrastrutture, ma anche sull'efficienza metodologica e sulla qualità dei dati.

L'Europa, al contrario, si presenta più frammentata e con risorse relativamente limitate. Tuttavia, iniziative come BLOOM, hanno portato a un tentativo significativo di costruire un

modello aperto e collaborativo. Parallelamente, programmi pubblici, come Eurostack, puntano a rafforzare la sovranità digitale europea, sebbene con forti ritardi rispetto alle strategie statunitensi e cinesi.

Tabella 4: FM principali ed emergenti e livello globale.

Regione	Modello	Strategia	Accesso	Dataset	Licenza
USA	OpenAI – GPT-5	Partnership con Microsoft, integrazione in Azure, Copilot, Bing, Edge	API	Proprietari	Restrizioni commerciali
USA	Google – Gemini	Integrazione in Google Search, Workspace, modelli multimodali	API	Proprietari	Restrizioni commerciali
USA	Anthropic – Claude	Distribuzione tramite AWS, focus su sicurezza e allineamento etico	API	Proprietari	Restrizioni commerciali
USA	Meta – LLaMA	Rilasci scientifici parziali, FM utilizzabile dai ricercatori	Ricercatori e sviluppatori	Parziale	Uso commerciale limitato
Cina	Baidu – Ernie	FM proprietario integrato nei servizi consumer e cloud			Controllo statale forte
Cina	Alibaba – Tongyi Qianwen	FM proprietario per applicazioni verticali			Controllo statale forte
Cina	Huawei e Tencent	FM proprietari e infrastruttura cloud			Controllo statale forte
Cina	ByteDance	FM futuri con dati comportamentali sfruttando TikTok			
Cina	DeepSeek	Dataset curati e tecniche di reasoning, ottimizzazione dell'efficienza computazionale			
EU	BLOOM	Open-source collaborativo, comunità scientifica internazionale	Pubblico	Completo	Open-source
EU	Mistral	Rilascio pesi e parte del codice rilasciati	Ricercatori e startup	Parziale	Uso commerciale limitato

UAE	Falcon (TII)	Rilascio pesi e documentazione tecnica	Ricercatori	Parziale	Uso commerciale limitato
-----	--------------	--	-------------	----------	--------------------------

Questa panoramica di attori e strategie mostra chiaramente come lo sviluppo dei *foundation models* non dipenda soltanto da competenze tecniche, ma sia il risultato di una combinazione di risorse finanziarie, infrastrutture cloud, disponibilità di dati e capacità di integrazione verticale lungo l'intero *stack*.

Il funzionamento dei mercati legati all'AI mostra dunque una crescente tendenza alla concentrazione, che nei FM assume la forma di un vero e proprio oligopolio ristretto. A livello globale, il settore si configura sempre più come un duopolio di fatto fra Stati Uniti e Cina.

Questa configurazione è rafforzata dai *network effects*: più utenti adottano un modello, più dati vengono generati, più migliorano le performance del modello stesso, creando un vantaggio cumulativo difficile da colmare. In parallelo, l'integrazione verticale lungo la catena del valore fa sì che i grandi player non si limitino a sviluppare FM, ma controllino anche altri segmenti della catena.

Il rischio è che la concorrenza si trasformi in un mercato di tipo *winner-takes-most*, dove pochi attori accumulano dati, calcolo e risorse finanziarie, per sviluppare gli FM, limitando l'ingresso di nuovi operatori e condizionando l'innovazione a valle.

2.2.5 Applicazioni

Il livello più a valle della catena è rappresentato dalle applicazioni, che traducono le capacità dei FM in servizi e prodotti accessibili agli utenti finali. È in questo segmento che il valore generato dall'intelligenza artificiale si monetizza, poiché i modelli addestrati e raffinati a monte trovano impiego in settori verticali e servizi trasversali, trasformando processi produttivi e modelli di business.

A differenza degli strati superiori fortemente concentrati e dominati da pochi attori globali, il livello applicativo si presenta più frammentato e dinamico. Questo succede grazie a start-up innovative, grandi imprese tecnologiche, fornitori di software *enterprise* e istituzioni pubbliche, le quali creano competizione e alimentano la varietà delle soluzioni disponibili. Nonostante la frammentazione, la capacità di distribuire le applicazioni, però, resta fortemente condizionata dal controllo esercitato a monte dai provider di *foundation models* e di servizi cloud, che detengono le principali leve di accesso.

Le applicazioni coprono diversi domini specializzati. In ambito medico, i modelli vengono impiegati nella diagnostica per immagini, nella ricerca di nuove molecole e nella personalizzazione delle cure. In finanza, supportano l'analisi dei rischi, il trading algoritmico e

la compliance regolatoria. Nella sfera dell'educazione sono alla base di tutor personalizzati e strumenti di valutazione e traduzione automatica. Nel settore dei media sono a supporto per la generazione di testi, immagini, musica e video. Infine, nel ramo della cybersecurity consentono di rilevare minacce e potenziali usi offensivi.

Le applicazioni oltre ad essere *specific-domain*, trovano spazio in settori trasversali. Nel campo della creatività vengono utilizzate per generare testi, immagini, musica e video. Nei servizi di customer care sono impiegate attraverso chatbot avanzati e assistenti virtuali multilingua. Nel contesto della produttività aziendale, strumenti come l'assistente di programmazione GitHub Copilot, Duet AI in Workspace Google e Microsoft Copilot in Office 365 rappresentano integrazioni strategiche in suite SaaS e sistemi operativi. Un ulteriore esempio di un servizio che potrà avere un utilizzo trasversale è Lara, il sistema di traduzione automatica sviluppato da Google: un'applicazione verticale nel dominio linguistico, costruita sopra i modelli multimodali, e adatta all'utilizzo in diversi contesti comunicativi.

Questi esempi mostrano come l'AI non agisca soltanto come tecnologia di supporto, ma ridisegni le modalità con cui le imprese interagiscono con clienti, dipendenti e mercati e anche nella vita quotidiana.

Secondo la *figura 3* tratta dall'AI Index Report 2024, si evince che nel 2023 circa il 33% delle organizzazioni a livello globale ha adottato strumenti di AI generativa, e se si guarda al Nord America la percentuale sale al 40%. Le funzioni che hanno riscontrato un'incidenza maggiore sono state il marketing e le vendite, lo sviluppo di prodotti e/o servizi, e il *service operations*. L'adozione, invece, è più limitata in settori regolati, come sanità e pubblica amministrazione, anche se la tendenza è in crescita.

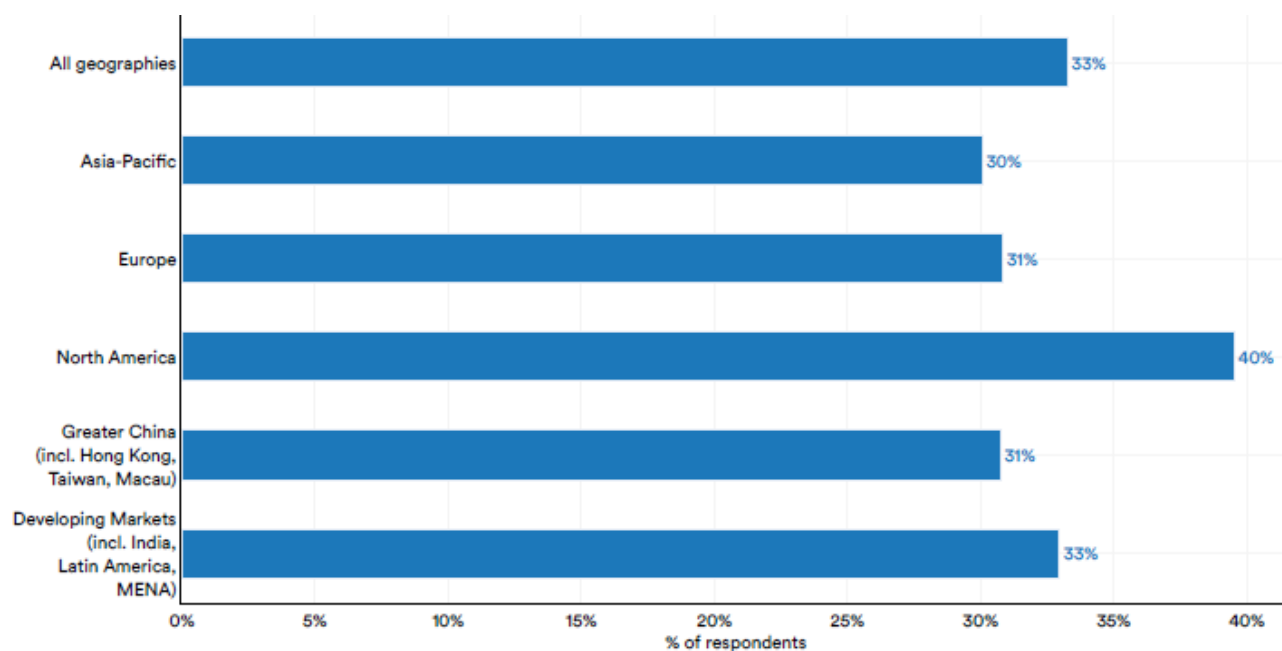


Figura 3: Tasso di adozione dell'AI generativa dalle organizzazioni nel mondo nel 2023

Un aspetto centrale sono i canali di distribuzione delle applicazioni. Le principali modalità sono tre: API, plugin e integrazioni. Le API consentono agli sviluppatori di collegare i modelli ai propri software; i plugin estendono le funzionalità di applicazioni già esistenti, rendendo l'AI fruibile senza necessità di programmazione; le integrazioni incorporano l'intelligenza artificiale direttamente in suite SaaS o in sistemi operativi.

Inoltre, i cloud provider hanno iniziato a proporre marketplace integrati, come Model Garden di Google, che consentono alle imprese di scegliere e implementare modelli già pronti, propri o di terzi. Questa infrastruttura di distribuzione rafforza il ruolo degli *hyperscaler*, che controllano anche i punti di accesso delle applicazioni downstream.

La struttura della concorrenza nelle applicazioni presenta diversi attori in gioco. Microsoft si posiziona come leader nell'integrazione esclusiva di GPT nei propri servizi cloud Azure e nelle applicazioni Office come Copilot e Bing; in aggiunta presenta un proprio marketplace con Azure AI Model Catalog di Microsoft. Google ha integrato Gemini nelle proprie piattaforme, da Workspace a Vertex AI, che è sia un ambiente per lo sviluppo tecnologico e la gestione di modelli di *machine learning* sia un vero e proprio canale di distribuzione dei FM, poiché attraverso il Model Garden consente di accedere a un ampio catalogo di modelli proprietari e open-source, disponibili via API e integrabili in applicazioni e servizi aziendali. Amazon ospita Anthropic tramite AWS e offre un canale di distribuzione per modelli di terzi tramite la

piattaforma Bedrock. Anthropic e Meta applicano rispettivamente i modelli Claude e LLaMA al settore dei social e dell'advertising. Il primo viene utilizzato per applicazioni legate alla moderazione dei contenuti online come il filtraggio, la sicurezza e la prevenzione di disinformazione, diventando una risorsa di moderazione del contenuto automatica e di gestione delle interazioni dell'utente. Il secondo svolge compiti pratici di miglioramento del targeting pubblicitario, tramite analisi del linguaggio e profilazione più precisa, di generazione di contenuti personalizzati, di supporto alla creazione automatica di annunci pubblicitari e di ottimizzazione delle interazioni degli utenti con chatbot e customer service. Entrambi rappresentano concorrenti rilevanti nel panorama applicativo dei modelli.

Questa concentrazione di player mostra come lo sviluppo applicativo sia strettamente legato alle scelte dei grandi attori *upstream*, rafforzando rischi di auto-preferenza e *lock-in*.

La memoria storica accumulata degli utenti costituisce un ulteriore fattore di *lock-in* a valle. Con questa espressione non si intende la memoria informatica in senso stretto, bensì l'insieme dei dati, delle interazioni passate, delle preferenze e delle personalizzazioni che si consolidano progressivamente all'interno di una piattaforma o di un servizio di intelligenza artificiale. Gli *incumbent* sfruttano questo fattore per aumentare il proprio vantaggio competitivo. Infatti, la migrazione verso un servizio alternativo implicherebbe per l'utente la perdita di questo patrimonio informativo e di tutte le personalizzazioni costruite nel tempo, con la conseguente necessità di ricominciare da zero, o nel caso di un'organizzazione di ricostruire i dataset proprietari, reimpostare le pipeline di lavoro e di rinunciare a ottimizzazioni sviluppate nel tempo. Questi cambiamenti si traducono in *switching costs* che scoraggiano l'utente o l'impresa ad abbandonare la piattaforma di riferimento, rafforzando così le barriere al cambiamento e il potere di mercato degli operatori dominanti.

Un esempio concreto di questo fenomeno è dato dagli assistenti AI generativi come ChatGPT, Claude o Copilot che memorizzano le conversazioni pregresse e lo stile di interazione, realizzando una forma di personalizzazione dinamica che non è trasferibile ad altre piattaforme. In sintesi, il livello applicativo non è solo il punto di contatto con l'utente finale, ma anche uno spazio strategico in cui si definisce la concorrenza tra grandi ecosistemi tecnologici. La capacità di controllare la distribuzione, integrare i modelli nei propri servizi e consolidare basi utenti ampie rende le applicazioni un campo cruciale per comprendere il futuro dell'AI e le sue implicazioni competitive.

Attore	Modello / Applicazione	Distribuzione	Strategia competitiva
Microsoft	Copilot/ChatGPT	Integrato in Office 365 e Azure, marketplace Azure AI Model Catalog	Memorizza conversazioni e preferenze utente, elevati switching costs
Google	Gemini/ Duet AI	Integrato in Workspace, Vertex AI e Model Garden	Personalizzazione attraverso dati d'uso e storici, difficoltà di migrazione
Amazon	Anthropic su AWS Bedrock	Distribuzione tramite piattaforma cloud, canale per modelli di terzi	Dipendenza dalla piattaforma cloud e dai servizi associati
Meta	LLaMA	Applicazioni social, advertising	Personalizzazione degli annunci e profilazione utenti, lock-in legato ai social graph
Altri/Startup	Applicazioni verticali specifiche	Settori specialistici, modelli SLM	Meno lock-in, ma risorse limitate e accesso ai dati più difficile

Tabella 5: Principali applicazioni dei maggiori attori nell'ecosistema AI.

2.3 Dinamiche competitive

2.3.1 Concentrazione per livello e interdipendenze

L'analisi della catena del valore dell'intelligenza artificiale non può limitarsi a osservare i singoli segmenti in maniera isolata. Le dinamiche competitive emergono infatti dalle interdipendenze tra hardware, cloud, dati, modelli fondamentali e applicazioni, che si rafforzano a vicenda e rendono l'intero ecosistema caratterizzato da un'elevata concentrazione. Per analizzare le diverse connessioni degli stadi della catena è opportuno riportare in sintesi le concentrazioni di potere spiegate esaurientemente nei paragrafi precedenti del capitolo 2.

Nel livello hardware Nvidia possiede una posizione dominante con il controllo dell'80% del mercato dei processori grafici per AI, in particolare delle GPU, e ha rafforzato il proprio vantaggio competitivo attraverso le economie di scala e all'ecosistema CUDA, che ha creato un *lock-in* software attorno alla propria architettura. Tale dominio è ulteriormente rafforzato dalla dipendenza strutturale da attori come TSMC e ASML nella catena fisica di produzione dei semiconduttori, evidenziando una concentrazione che non si limita al lato tecnologico, ma coinvolge anche la supply chain.

Il livello cloud è a sua volta altamente concentrato: AWS, Microsoft Azure e Google Cloud raccolgono insieme circa i due terzi del mercato globale dei servizi infrastrutturale. Ai principali attori si affiancano provider emergenti come CoreWeave e Lambda, che però sopravvivono

soprattutto grazie a partnership e ingenti investimenti degli incumbent. In questo senso, la concorrenza riguarda i tre cloud di riferimento che controllano l'accesso alla capacità computazionale necessaria allo sviluppo dei *foundation model*.

La concentrazione nel livello dati è caratterizzata da due aspetti. Dall'asimmetria marcata tra i dataset proprietari detenuti dalle Big Tech e le risorse “*commons*” aperte alla comunità, e dalla crescente pratica del blocco dei crawler e delle restrizioni legate al copyright che stanno irrigidendo ulteriormente l'accesso a questo input essenziale.

Il livello dei FM presenta pochi attori (OpenAI, Google, Anthropic, Meta e Amazon), che agiscono da leader globali. Tuttavia, si ha una competizione dinamica per la presenza di modelli open source come LLaMA, Falcon e Mistral.

Infine, il livello delle applicazioni appare più frammentato, ma con tendenze al consolidamento. Infatti, la diffusione di Copilot di Microsoft, Duet AI di Google e Bedrock di AWS indica come l'integrazione nei grandi ecosistemi digitali rischi di creare posizioni dominanti anche a valle. Oltre alle singole concentrazioni, è fondamentale evidenziare le interdipendenze tra i diversi livelli dello *stack*, che rafforzano ulteriormente la tendenza alla diminuzione di competizione. L'hardware è strettamente legato al cloud, infatti Nvidia fornisce le GPU in quantità limitate, e i principali acquirenti di queste risorse sono gli *hyperscaler*, rafforzando così un oligopolio bilaterale. Il cloud a sua volta costituisce la piattaforma necessaria per lo sviluppo dei FM, che allo stesso tempo trainano la domanda di capacità computazionale. La relazione tra dati e modelli è bidirezionale: i FM necessitano di enormi quantità di dati per essere addestrati, ma allo stesso tempo le applicazioni generate dai modelli producono nuovi dataset che alimentano il ciclo successivo di training. Infine, le applicazioni dipendono strettamente dai FM, ma al contempo ne rafforzano la diffusione e il consolidamento. In questo modo, ogni livello non solo presenta proprie barriere e concentrazioni, ma alimenta direttamente quelle degli altri, generando un sistema chiuso e fortemente interconnesso.

2.3.2 Meccanismi economici comuni nell'AI stack

Al di là delle concentrazioni per livello, le dinamiche competitive dell'*AI stack* sono guidate da alcuni meccanismi economici ricorrenti.

Sono presenti economie di scala che emergono soprattutto nel training dei modelli: il costo fisso per addestrare un FM è altissimo, mentre i costi marginali di inferenza tendono a decrescere. Dinamiche simili si osservano anche nel livello hardware con costi miliardari per gli impianti produttivi dei semiconduttori e gli investimenti in R&D, nel cloud con la costruzione e manutenzione dei data center che vengono ripagati solo con una rete globale vasta. Questa

configurazione spinge verso strutture oligopolistiche e porta a condizioni di monopolio naturale.

Le economie di scopo derivano invece dall'integrazione tra diversi livelli dello *stack*: chi controlla compute, dati e modelli può sfruttare sinergie e *cross-subsidization* difficilmente replicabili dai nuovi entranti; infatti, i dati raccolti da un servizio migliorano i modelli, questi ultimi alimentano nuove applicazioni ed esse generano ulteriori dati. I casi di Microsoft–OpenAI, Google–Gemini e Amazon–Anthropic mostrano come la cooperazione strategica consenta di combinare competenze e risorse per rafforzare la posizione di mercato. Anche sul piano infrastrutturale si può notare il verificarsi di tale dinamica con gli *hyperscaler* che usano gli stessi data center per servire clienti cloud generici, addestrare FM proprietari e distribuire applicazioni integrate, massimizzando la redditività degli investimenti.

Gli effetti di rete pur assumendo forme diverse, contribuiscono anch'essi a rafforzare i leader. Nel cloud computing sono presenti effetti indiretti: il consolidamento della base clienti genera vantaggi reputazionali e finanziari che rafforzano i leader grazie alla possibilità di espansione globale. Nel caso dei modelli, le interazioni con gli utenti rendono progressivamente più accurati i sistemi, creando vantaggi cumulativi. Lo stesso meccanismo si instaura nelle applicazioni, dove l'adozione diffusa tende a consolidare standard de facto, per esempio GitHub Copilot per gli sviluppatori, attirando sempre nuovi utenti e sviluppatori di plug-in.

Infine, i *feedback loops* tra livelli alimentano circuiti di retroazione che rafforzano i vantaggi acquisiti. L'hardware più avanzato consente di addestrare modelli più potenti, i modelli migliori attirano più applicazioni e queste ultime generano nuovi dati che alimentano ulteriori cicli di training per i modelli. A differenza del *search* tradizionale, dove i dati di *clickstream* alimentano in modo massiccio i motori, qui i feedback sono più selettivi, ma non per questo meno strategici. Un esempio di questo flusso è rappresentato da Copilot, ChatGPT o altri sistemi che generano interazioni utente di alto valore qualitativo, le quali diventano input esclusivi per raffinare i modelli stessi. Ne deriva un meccanismo di accumulazione inter-livello che rafforza i leader e alza ulteriormente le barriere all'ingresso.

2.3.3 Integrazione verticale e leve cross-layer

Uno degli aspetti più rilevanti riguarda l'integrazione verticale che permette a un singolo attore di presidiare più livelli della catena, come mostrato nella *tabella 6*, grazie alle strategie di *leveraging* e *self-preferencing*, compromettendo così la concorrenza.

Hyperscaler	Livelli dello stack	Asset / Tecnologie principali
Microsoft	Cloud, FM, Applicazioni	Azure, GPU NVIDIA, modelli OpenAI (GPT, Copilot), Office 365, Bing, Edge, Azure AI Model Catalog
Google	Hardware, Cloud, FM, Applicazioni	TPU, Google Cloud, modelli Gemini, Duet AI, Workspace, Model Garden, Search AI integration
Amazon	Hardware, Cloud, FM, Applicazioni	Trainium e Inferentia, AWS Cloud, marketplace Bedrock, integrazione enterprise
Meta	Hardware, Cloud, FM, Applicazioni	Chip proprietari per LLaMA, infrastruttura data center, integrazione advertising e targeting

Tabella 6: Gli hyperscaler nell'AI stack.

Gli esempi più evidenti di questa dinamica sono gli *hyperscaler*. Infatti, non solo dominano i servizi cloud, ma hanno progressivamente esteso la loro attività alla progettazione di chip proprietari, come le TPU di Google o il Trainium di Amazon, nella messa a disposizione di modelli sviluppati in proprio o tramite partnership, e alla distribuzione di applicazioni integrate nei propri ecosistemi, come Microsoft Copilot o Google Duet AI. In questo modo, riescono a esercitare un controllo che non si esaurisce in un singolo segmento, ma che si estende lungo più stadi dello *stack*, condizionando anche la competizione a valle.

Un aspetto che contribuisce a questo fenomeno è l'integrazione *cross-layer*, che estende il potere di mercato da segmenti infrastrutturali come calcolo e dati fino alle applicazioni finali, alterando gli incentivi competitivi. Non si tratta solo di un'integrazione verticale classica, ma di una strategia che collega vari livelli dello *stack* e consente agli *incumbent* in un punto della catena di condizionare gli altri. Un'impresa che possiede chip proprietari e cloud, ad esempio, può riservare risorse preferenziali ai propri prodotti, riducendo la contestabilità da parte di terzi. Analogamente chi detiene dataset esclusivi può sfruttarli per addestrare FM più potenti per poi trasformare questo vantaggio anche nelle applicazioni a valle.

Una strategia di leva verticale tipica della parte finale della catena è il *bundling*: i grandi player all'interno dei propri servizi integrano le proprie tecnologie, più nello specifico i propri *foundation models* e chatbot, per creare un ecosistema chiuso. Alcuni esempi concreti di questa pratica si osservano nelle app di messaggistica come Whatsapp, dove Meta, proprietaria della piattaforma, ha integrato l'accesso a Meta AI consentendo agli utenti di utilizzare direttamente le funzioni della propria intelligenza artificiale generativa. Un meccanismo analogo si riscontra nei social di sua proprietà, come Instagram o Facebook dove accedendo rispettivamente ai Dm

e a Messenger si può usufruire di tale servizio. In questo modo chi detiene il controllo di social e piattaforme può promuovere e sponsorizzare direttamente le proprie soluzioni AI, rafforzando la dominanza e il potere di controllo del mercato.

Infine, la possibilità di controllare i canali di distribuzione alimenta fenomeni di *foreclosure* indiretta, in cui i concorrenti non vengono esclusi formalmente, ma subiscono una perdita di visibilità e accesso. Questo meccanismo avviene poiché controllando le API si può decidere chi può accedere ad un determinato servizio e, possedendo i cosiddetti *model gardens*, gli *incumbent* possono decidere il livello di visibilità dei vari modelli, compresi i propri, creando una forma di controllo sull'accesso al mercato. In questo modo, il potere acquisito a monte tende a insidiarsi lungo tutta la catena, riducendo la contestabilità e orientando gli incentivi competitivi verso strategie di mantenimento del potere piuttosto che verso il miglioramento della qualità o l'innovazione.

In sintesi, le dinamiche competitive dell'*AI stack* mostrano un duplice volto: da un lato, forti tendenze alla concentrazione e all'integrazione verticale; dall'altro, segnali di innovazione, apertura e nuovi ingressi che suggeriscono come il gioco competitivo sia ancora in corso, e non sia ancora avvenuto un definitivo punto di svolta verso il "*winner-takes-all*".

3 Effetti economici ed organizzativi dell'AI

L'adozione dell'intelligenza artificiale da parte delle imprese e, più in generale, dei sistemi economici, sta generando effetti macroeconomici su un'ampia varietà di aree: occupazione, disuguaglianza, produttività e crescita, organizzazione aziendale, innovazione e modelli di business, concorrenza e struttura dei mercati. Tuttavia, gli impatti osservabili oggi sono ancora limitati in quanto la diffusione su ampia scala è ancora in una fase iniziale. Di conseguenza, ci si aspetta che le trasformazioni più radicali e profonde si manifesteranno nel lungo periodo.

Inoltre, i dati disponibili sono pochi e spesso contraddittori. È difficile accedere a informazioni a livello aziendale e spesso i risultati delle analisi empiriche appaiono incoerenti, non permettendo di generalizzare a contesti più ampi. Questo succede in quanto gli effetti sono disomogenei: dipendono dal settore, dalla base tecnologica di partenza delle imprese, dalle competenze del capitale umano, dalla geografia e da molti altri fattori strettamente legati alla situazione analizzata.

Tali complessità portano ad avere alcune discrepanze fra le ricerche teoriche e quelle empiriche, a partire dal concetto stesso di intelligenza artificiale. Le prime tendono ad adottare una visione molto più estesa di AI, coerente con quella di scienziati ed ingegneri. Come anticipato nell'introduzione, si tratta di una *general purpose technology*, quindi una tecnologia abilitante e pervasiva, che sarà in grado di rivoluzionare ogni ambito dell'economia e della società in modo paragonabile all'elettricità o alla macchina a vapore. I ricercatori empirici, invece, devono fare i conti con la scarsità di dati. La mancanza di informazioni affidabili, dettagliate, recenti ed aggiornate si riflette nei loro studi: essi, infatti, spesso utilizzano dati relativi alla precedente ondata tecnologica, ovvero l'automazione, limitando così di molto il campo di applicazione.

In aggiunta, gli effetti parziali e sfumati dell'adozione attuale non permettono di costruire una visione unificata, emergono pareri divergenti: da un lato, l'AI è vista come una leva per aumentare l'efficienza, innovare e rilanciare la produttività; dall'altro, suscita timori legati alla sostituzione del lavoro, all'aumento delle disuguaglianze e alla concentrazione del potere economico.

Alla luce di queste premesse, il capitolo prosegue analizzando le aree principali in cui l'intelligenza artificiale sta già producendo, o potrà produrre, gli effetti più rilevanti. Al termine di ogni paragrafo verrà proposta una tabella di sintesi in cui sono riportati i paper citati e le rispettive conclusioni.

3.1 Occupazione

3.1.1. L'AI ed il futuro dell'occupazione

Tra le molteplici aree in cui l'intelligenza artificiale sta generando trasformazioni, il mondo del lavoro è senza dubbio uno degli ambiti più discussi e studiati. Il dibattito, come anticipato in precedenza, oscilla tra visioni pessimistiche ed ottimistiche. Le prime sostengono che l'AI sarà uno strumento in grado di aumentare l'efficienza, creare nuove occupazioni e valorizzare le competenze umane, liberando il loro massimo potenziale; le seconde associano la nuova tecnologia a rischi elevati di disoccupazione, riduzione dei salari e, più in generale, ad un deterioramento del benessere collettivo. Questi poli di pensiero sono la semplificazione di una realtà molto più ampia e complessa. I progressi del *deep learning* e del *machine learning* sono sorprendenti e stanno già avendo ripercussioni sulla forza lavoro, ma si è ancora lontani dall'avere un'*artificial general intelligence* (AGI) in grado di eguagliare l'uomo in tutte le aree cognitive. Diventa, quindi, fondamentale distinguere tra occupazioni e compiti, in quanto l'intelligenza artificiale impatterà sui compiti costituenti le occupazioni e, solo in via indiretta, su queste ultime nel loro complesso. In secondo luogo è necessario capire se l'AI agirà come fattore sostitutivo o complementare delle competenze umane e, se in caso di spiazzamento, sarà in grado di generare nuove occupazioni necessarie a sostenere la società e l'equilibrio economico.

3.1.2 L'approccio *task-based* e le implicazioni organizzative

Una delle questioni principali che gli studiosi stanno affrontando riguarda l'effetto netto che l'AI avrà sul lavoro umano: porterà a disoccupazione o creerà nuovi compiti? Prevarrà l'effetto sostituzione o complementarità? Per rispondere a questi interrogativi si adotta l'approccio *task-based* proposto dagli autori Autor, Acemoglu e Restrepo (2020), secondo cui le occupazioni possono essere scomposte in diversi compiti elementari. L'obiettivo è quello di individuare quali task sono più propensi ad essere automatizzati e quali invece resteranno prettamente di competenza umana. Ci si basa sullo schema introdotto da Brynjolfsson e Mitchell (2017)¹. Gli studiosi citati per svolgere la loro analisi si sono concentrati su 23 domande valutative, che accettano risposte comprese in un range fra 1 e 5, in cui 1 è "per niente automatizzabile" e 5 è "perfettamente automatizzabile". Un'esposizione neutra corrisponde a un punteggio di 3. La scala ordinale così formata permette di introdurre il concetto di *suitability for machine learning* (SML), che indica il grado di esposizione all'automazione per ogni tipo di task. Lo studio si

¹ Si tenga presente che il metodo adottato si basa esclusivamente sulla fattibilità tecnica e non su valutazioni di tipo sociale, etico, organizzativo o culturale.

basa su 2.059 attività lavorative dettagliate, 950 occupazioni e 18.112 compiti comuni a più mansioni diverse (dati provenienti dal database O*NET).

I risultati suggeriscono che ci sono poche occupazioni esposte all'AI nel loro complesso, ma che ognuna di esse ha almeno qualche compito altamente automatizzabile. Inoltre, il fatto che non tutti i compiti siano ugualmente suscettibili di automazione, porta a rilevare una delle prime origini della disuguaglianza degli effetti visibili: non tutte le occupazioni sono soggette allo stesso impatto dell'AI e, più nello specifico, del ML. La varianza all'interno dei compiti, infatti, è maggiore rispetto a quella rilevata all'interno delle occupazioni, che invece si mostra minore in quanto l'aggregazione ha un potere diversificante. Si evince, quindi, che la soluzione non sta nel concentrarsi sulla completa automazione di alcune occupazioni, con conseguente sostituzione radicale di interi ruoli, ma nel riorganizzare il lavoro in modo che esso sia composto da un insieme coerente ed armonizzato di compiti svolti da una parte dall'essere umano (non SML) e dell'altro dalla macchina (SML).

Per effettuare la riorganizzazione del lavoro, con l'opportuna separazione e ricombinazione dei compiti, saranno richieste competenze umane. Nello specifico occorre individuare le attività suscettibili di automazione, analizzandole per capire se esiste una mappatura precisa tra input e output, in modo che l'algoritmo di ML possa apprendere, generare o prevedere in modo preciso e accurato il risultato richiesto. Inoltre, viene esaminata la presenza di dati digitali, di una routine e di una struttura prevedibile. L'aumento delle capacità delle macchine o la riduzione dei costi del capitale per un determinato compito aumentano gli incentivi a sostituire il capitale al lavoro, quindi hanno una grande influenza sulla riorganizzazione generale. Un'aggregazione inefficiente dei task riduce inevitabilmente i potenziali guadagni di produttività; basti pensare alla classica funzione di produzione in stile Leontief, in cui gli input sono vincolati dall'input presente in quantità minima: se, ad esempio, le occupazioni vengono strutturate in pacchetti standard, in cui l'uomo svolge anche compiti che in realtà potrebbero essere automatizzati, si ha una perdita di efficienza in quanto l'umano potrebbe essere reindirizzato verso altre attività e mansioni che invece sono non SML e quindi richiederebbero il suo contributo.

3.1.3 Evidenze empiriche

Sebbene la teoria *task-based* proposta sia molto chiara ed esplicativa, le evidenze empiriche sono spesso contraddittorie. Come anticipato, alcune mettono in luce gli effetti sostitutivi mentre altre quelli di complementarità. L'eterogeneità rilevata deriva innanzitutto dal fatto che compiti diversi sono soggetti ad un grado differente di esposizione all'automazione. In secondo luogo è causata dalla grande variabilità degli effetti legata al settore, alla dimensione ed alla geografia.

È stato riscontrato, infatti, che nella manifattura l'azione è più lenta e soprattutto rimane strettamente legata alla robotica ed alla automazione, mentre nel settore dei servizi c'è una maggiore spinta, con la propensione all'adozione dell'AI specialmente nelle attività cognitive, come ad esempio funzioni HR o amministrative. Un esempio concreto è rappresentato dagli algoritmi che supportano le risorse umane nello screening dei curriculum dei candidati ad opportunità lavorative.

Per quanto riguarda, invece, la dimensione dell'impresa, le aziende di maggiore dimensione tendono ad adottare con più facilità la nuova tecnologia, mentre le PMI rischiano di rimanere indietro. Ciò è dovuto principalmente a risorse limitate, barriere organizzative e carenza di competenze digitali. Approfondiremo queste tematiche nei paragrafi successivi, dando luce al concetto di ITIC (*IT-related intangible capital*).

Infine, si nota una disuguaglianza sostanziale nell'adozione dell'AI a livello geografico. Tenendo in considerazione America, Cina ed Europa, l'evidenza empirica mette in luce il fatto che le prime due abbiano tassi di adozione molto più elevati, mentre la terza sia più conservativa. Il fenomeno riscontrato è sicuramente legato al contesto normativo di riferimento, che verrà approfondito nel capitolo successivo.

3.1.4 Salari e qualifiche

Per quanto riguarda l'effetto a livello occupazionale in correlazione con il livello salariale, ci sono diversi studi che mostrano risultati divergenti. Alcuni di essi sostengono che l'ondata tecnologica attuale probabilmente si differenzierà da quella legata all'automazione ed ai robot in quanto mentre quest'ultima ha portato ad aumenti di produttività legati alla sostituzione di esseri umani a bassa qualifica con macchine, l'AI avrà un effetto più generalizzato e, di conseguenza, potrà colpire anche i lavoratori altamente qualificati. Altri studiosi, invece, continuano a sostenere che i lavoratori più colpiti saranno quelli a bassa qualifica, i quali svolgono attività altamente ripetitive.

Brynjolfsson, Mitchell e Rock (2018) affermano che non si nota una correlazione significativa fra l'esposizione all'AI e il salario: prendendo in considerazione i grafici scatterplot riportati in *figure 4 e 5*, infatti, si osserva come per ciascun livello salariale i dati riguardanti la *suitability for machine learning* siano molto diversificati, suggerendo che tutti i lavoratori saranno soggetti agli impatti dell'AI, senza una reale distinzione legata alla paga.

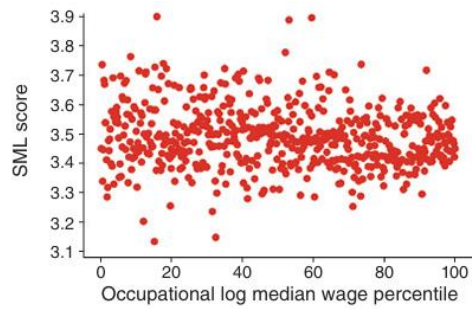


Figura 4: punteggio SML rispetto al percentile del salario mediano logaritmico occupazionale

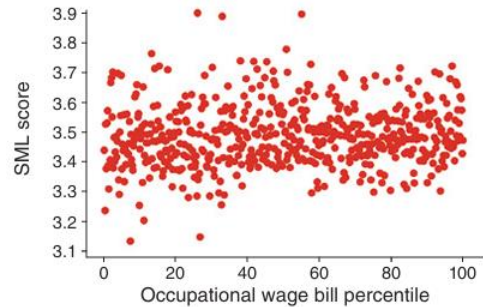


Figura 5: punteggio SML rispetto al percentile della massa salariale professionale

In un successivo studio Brynjolfsson, Rock, Tambe (2019) propongono una letteratura leggermente diversa. Analizzando dati più recenti essi riscontrano un lieve pattern a sostegno del fatto che invece ci sia una propensione all'automazione per i lavoratori soggetti a salari più bassi. In *figura 6* è mostrata la retta di regressione stimata. Essa si presenta leggermente inclinata negativamente, ma comunque i dati mostrano grande differenziazione interna, con ampia varietà per ciascun livello salariale.

È utile analizzare anche l'impatto dell'AI e del ML sui salari stessi. Molti studiosi sostengono che i salari in media non saranno ridotti, ma ci sarà uno spostamento dai lavoratori meno qualificati, che vedranno una contrazione delle paghe a causa dell'effetto sostituzione, a coloro con titoli di studio più avanzati. Essi potranno beneficiare della nuova ondata tecnologica grazie all'effetto di complementarità che potrà verificarsi. Se l'intelligenza artificiale agirà per potenziare le capacità umane, essi saranno pagati maggiormente in quanto con l'ampia diffusione e adozione di algoritmi ci sarà sempre più bisogno di personale con ampie competenze, specialmente in ambito STEM. In questo senso, l'AI non ridisegna soltanto le dinamiche salariali, ma contribuisce a ridefinire la gerarchia delle competenze richieste nel mercato del lavoro.

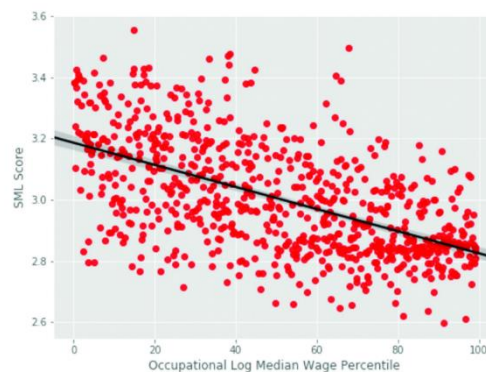


Figura 6: punteggio SML rispetto al percentile del salario mediano nel 2016 (coefficiente di regressione: -0.0034)

3.1.5 La polarizzazione occupazionale

Un fenomeno molto discusso è quello della possibile polarizzazione del lavoro. Alcuni studi (Lu 2021, seguendo il filone di Autor e Acemoglu), indicano che le nuove tecnologie digitali, tra cui l'intelligenza artificiale, porteranno ad un progressivo svuotamento (*hollowing out*) delle occupazioni a media qualifica. Ciò risulta possibile in quanto tali mansioni sono altamente codificabili, non totalmente soggette ad una routine precisa, ma comunque suscettibili ad un elevato grado di automazione. Questo, invece, non accade così frequentemente per alcuni lavori a bassa qualifica in quanto l'uomo è ancora necessario in molti compiti che devono essere svolti manualmente (ad esempio meccanici ed idraulici) e perché il *machine learning* aiuterà a richiedere nuove figure di riferimento a bassa qualifica per svolgere determinate attività legate all'AI, come ad esempio personale addetto alla gestione e pulizia dati per quanto riguarda l'addestramento dei modelli.

Per far fronte alla situazione prevista, quindi, bisognerà avviare una ristrutturazione delle competenze possedute da coloro che attualmente si occupano di svolgere mansioni riferenti alle competenze di medio livello, in modo tale da proporre una soluzione concreta al possibile spiazzamento. Parallelamente, col sorgere di nuovi bisogni, nasceranno nuove figure professionali, presumibilmente orientate verso competenze di alto profilo: si prevede, infatti, una richiesta sempre maggiore di informatici e matematici, fino a quando, qualora mai accadesse, l'AI diventerà così potente da superare l'uomo e sarà in grado di generare macchine in autonomia. In tal caso si parlerà di singolarità e le figure richieste saranno orientate verso uno stampo umanistico, come verrà approfondito nel paragrafo inerente all'educazione.

In sintesi, l'impatto dell'AI sull'occupazione va letto in maniera multidimensionale: essa non sostituisce completamente le occupazioni ma piuttosto richiede una loro profonda ristrutturazione, a livello di compiti, competenze e organizzazione.

Tabella 7: Effetti dell'AI a livello occupazionale

Paper	Tipologia di dato	Effetti trovati
Autor, Acemoglu, Restrepo (2020)	Modello teorico <i>task based</i>	Le occupazioni possono essere scomposte in compiti, di cui alcuni più automatizzabili ed altri meno. L'AI non sostituisce le occupazioni nella loro interezza ma piuttosto richiede

		una ricombinazione dei compiti tra uomo e macchina.
Brynjolfsson & Mitchell (2017)	Analisi su database O*NET	Introduzione del concetto di SML. Poche occupazioni sono totalmente automatizzabili ma la maggior parte hanno almeno alcuni compiti che lo sono.
Brynjolfsson, Mitchell, Rock (2018)	Analisi empirica basata su O*NET con aggiunta di dati salariali BLS	Tutti i lavoratori sono potenzialmente esposti all'impatto dell'AI. Non si trova una correlazione significativa col livello dei salari.
Brynjolfsson, Rock, Tambe (2019)	Analisi empirica aggiornata basata su O*NET con aggiunta di dati salariali BLS	I lavoratori a basso salario hanno probabilità maggiore di essere automatizzati, tuttavia la dispersione interna resta elevata.
Lu e Zhou (2021)	Analisi teorico-empirica	Rischio di erosione delle occupazioni a media qualifica, richiesta in crescita di competenze di alto livello (in particolare STEM). I lavori manuali a bassa qualifica potrebbero resistere in quanto richiedono ancora presenza umana.

3.2 Produttività e crescita

3.2.1 Il paradosso di Solow e la J-curve

Molti studiosi considerano l'intelligenza artificiale una potenziale GPT, quindi come una tecnologia rivoluzionaria, abilitante e pervasiva, in grado di massimizzare la produttività, ottimizzare i processi produttivi e, più in generale, dare una svolta all'intera economia e generare effetti trasformativi su larga scala. Già nel decennio scorso Brynjolfsson e McAfee

(2014) in *The Second Machine Age: Work, Progress, and Prosperity in a Time of Brilliant Technologies* descrivevano la situazione dell'epoca come la fase iniziale di un nuovo cambiamento profondo e radicale come quello della Rivoluzione Industriale. Saniee et al. (2017) sostenevano che ci sarebbe stato un balzo di produttività nell'economia degli Stati Uniti fra il 2028 ed il 2033. Queste idee, però, non sono ancora state associate ad evidenze pratiche che ne confermino il potenziale ed i risultati. Nonostante gli ingenti e continui investimenti tecnologici fatti dalle imprese, la *total factor productivity* (TFP) non rivela alcun miglioramento e, anzi, è calata dall'1.5% all'1% annuo negli ultimi 50 anni. Il fenomeno ha risentito della crisi finanziaria globale del 2008 (GFC). In conseguenza al momento di forte turbolenza, si è verificato un rallentamento della crescita della produttività del lavoro. Questo è stato riscontrato sia nei paesi sviluppati che in quelli emergenti o in via di sviluppo, che agli inizi degli anni 2000 avevano visto una crescita rapida, fino a raggiungimento di un picco proprio in prossimità della crisi.

Si configura così un paradosso, in quanto da un lato c'è la spinta all'automazione ed all'efficienza, mentre dall'altro si nota una stagnazione dell'economia. Questa evidenza è nota come paradosso di Solow, e molti studiosi se ne sono occupati negli ultimi anni, fra cui Gordon (2016) e Brynjolfsson et al. (2019). Questi ultimi attribuiscono il calo di produttività al ritardo nell'ampia diffusione delle tecnologie avanzate di AI, come già era accaduto per i progressi precedenti, quali l'ICT. Prima che le nuove scoperte portassero ad un'effettiva crescita della produttività ci fu un periodo di stallo. Solamente dopo circa 25 anni di crescita molto lenta si ebbe il boom di produttività, il quale durò circa un decennio.

Negli studi di Acemoglu et al. (2014) si trovano alcuni risultati sorprendenti: nei settori manifatturieri statunitensi a maggiore intensità di IT i guadagni in termini di produttività sono evidenti più nei produttori che nei distributori. Inoltre, essi sono dovuti non ad un aumento della produzione ma ad una diminuzione dell'output associata ad una maggiore concentrazione dell'occupazione. Il paradosso, quindi, non è risolto nemmeno con l'ultima tecnologia precedente all'AI, e resta aperta la questione se quest'ultima sarà in grado di superarlo.

Le evidenze empiriche sono ancora scarse e poco significative e, in un contesto così incerto e complesso, Brynjolfsson, Rock e Syverson (2021) propongono la teoria della *Productivity J-Curve* (si veda la *figura 7*). Essa afferma che le tecnologie ad uso generale, come l'intelligenza artificiale, seguono 3 fasi specifiche:

1. Inizialmente incontrano frizioni: i costi da sostenere sono elevati ed i ritorni molto bassi se non nulli

2. Dopo gli investimenti iniziali c'è una fase di rallentamento in cui si apprende il funzionamento della tecnologia, si assorbono i complementi e si ristrutturano i processi e l'organizzazione stessa.
3. Infine nel lungo periodo si traggono i benefici e i vantaggi: si incrementano i ritorni e finalmente c'è il tanto atteso guadagno di produttività.

Quando si discute di investimenti, non si tratta solo di investimenti diretti tangibili, ma anche e soprattutto di investimenti intangibili (ITIC), costi associati alla riorganizzazione aziendale e delle mansioni, re-invenzione di processi aziendali e formazione del capitale umano. La nuova ondata tecnologica, infatti, si distingue profondamente da quella precedente riguardante l'ICT, in quanto attualmente, grazie ai servizi cloud ed ai modelli open-source gli investimenti richiesti in capitale sono molto inferiori, si tratta più di conoscenza intangibile e know-how. Solo una volta che tutti questi elementi saranno sviluppati ed implementati l'AI inizierà a manifestare i suoi più profondi benefici e guadagni.

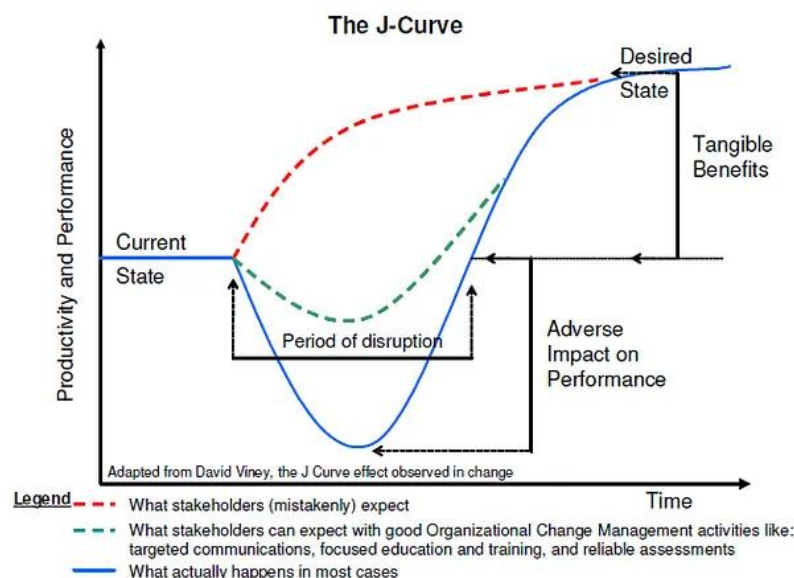


Figura 7: Andamento J-curve

3.2.2 Evidenze empiriche

Molte ricerche empiriche si concentrano sulla misura dell'impatto dell'AI in termini di produttività e innovazione. In particolare, lo studio di McElheran et al. (2024) analizza l'adozione dell'intelligenza artificiale negli Stati Uniti, evidenziando che, nonostante solo il 6% delle imprese dichiarino di utilizzare tecnologie AI, tale quota sale al 18% quando ponderata per il numero di lavoratori coinvolti. Questo risultato dimostra che attualmente solo poche imprese stanno sfruttando appieno il potenziale della nuova tecnologia e, inoltre, esse sono le imprese di più grandi dimensioni. Ne deriva un rischio rilevante: quello di amplificare il divario di performance tra imprese di dimensioni medio-grandi e piccole. Infatti, alcune realtà aziendali

potrebbero realizzare aumenti di produttività e emergere per quanto riguarda l'innovazione, mentre molte altre potrebbero rimanere indietro.

Gli effetti riscontrati sono strettamente correlati alle conoscenze ed al know-how che caratterizza ogni impresa. Proprio in questo ambito gioca un ruolo cruciale l'ITIC (*IT-related Intangible Capital*), ovvero il capitale intangibile legato a tutto ciò che è il mondo IT. Esso comprende, ad esempio: cultura organizzativa, competenze manageriali e governance dei dati. Le imprese, che già nella scorsa ondata tecnologica avevano sviluppato competenze chiave, avranno un vantaggio nel sapersi destreggiare meglio nell'utilizzo dell'AI.

In aggiunta, anche il settore di appartenenza incide sulla propensione all'adozione: infatti, l'evidenza empirica mette in luce che la manifattura, i laboratori diagnostici e l'ambito dell'informazione attualmente hanno un tasso di adozione superiore alla media e, di conseguenza, un vantaggio comparato in termini di competenze rispetto ad altri settori.

Inoltre, l'AI e, nello specifico, il ML, non agiscono in modo isolato. Gli studiosi Goldfarb, Taska e Teodoridis (2020) analizzano un grande quantitativo di dati dal dataset di Burning Glass Technologies. Gli annunci lavorativi portano ad affermare che, nonostante si tratti di *general purpose technology*, le nuove tecnologie sono strettamente legate ad altre discipline, quali: robotica, business intelligence, big data, data mining, data science e NLP (*natural language processing*). Questi legami, riportati in *figura 8*, indicano che l'AI non è una tecnologia autonoma, ma parte integrante di un ecosistema tecnologico complesso, in cui il ruolo dei dati è centrale. Infatti, quando l'AI viene combinata alle tecnologie sopra citate, porta ad avere la massimizzazione degli impatti potenziali e del progresso in termini di innovazione di prodotto e processo.

Nel complesso, i risultati ottenuti empiricamente confermano l'ipotesi della *productivity J-curve*: per avere un incremento di produttività a livello aggregato sarà necessario un processo graduale di diffusione dell'AI, supportato da investimenti in capitale intangibile. Per ora, però, rimane evidente che l'intelligenza artificiale porti con sé un grande potenziale trasformativo, soprattutto nelle imprese e nei contesti in cui le condizioni complementari sono già presenti.

	ML	BI	Big Data	Data Mining	Data Science	NLP	Cloud	Telecom	VM	GIS	Quantum	Robotics	Nanotech	IoT	CRISPR	VR	3D Print	Polymer	Blockchain	Web2.0	SOA	RFID
ML	100	9.8	32.5	12.8	44.8	14.1	19.2	1.9	4.2	0.6	0.6	6.3	0.0	5.7	0.1	1.5	0.3	0.0	2.2	0.0	0.7	0.1
BI	4.4	100	10.7	5.9	9.0	0.6	10.3	2.0	2.2	0.5	0.1	0.7	0.0	0.9	0.0	0.1	0.0	0.0	0.3	0.0	0.7	0.0
Big Data	19.1	14.0	100	6.5	21.3	3.8	27.6	2.8	10.7	0.5	0.2	0.9	0.0	3.2	0.0	0.2	0.0	0.0	0.7	0.1	1.0	0.2
Data Mining	20.7	21.1	17.8	100	26.6	6.2	7.0	2.2	1.6	1.1	0.3	1.1	0.0	1.1	0.1	0.2	0.0	0.0	0.4	0.1	0.2	0.1
Data Science	36.2	16.1	29.3	13.3	100	7.6	13.5	1.9	3.3	1.0	0.2	1.5	0.0	2.7	0.1	0.3	0.1	0.0	0.7	0.0	0.3	0.2
NLP	54.3	5.4	25.1	14.7	36.3	100	14.2	3.9	2.8	0.4	0.5	4.4	0.0	3.1	0.0	0.4	0.0	0.0	2.1	0.1	0.5	0.0
Cloud	4.9	5.9	12.1	1.1	4.3	1.0	100	4.2	17.0	0.3	0.2	0.6	0.0	2.7	0.0	0.1	0.0	0.0	0.6	0.1	1.2	0.0
Telecom	0.7	1.6	1.8	0.5	0.9	0.4	6.1	100	5.0	0.7	0.1	0.4	0.0	1.3	0.0	0.1	0.0	0.0	0.1	0.1	0.1	0.2
VM	2.9	3.3	12.5	0.7	2.8	0.5	45.3	9.4	100	0.3	0.2	0.4	0.0	2.0	0.0	0.1	0.0	0.0	0.3	0.1	0.8	0.1
GIS	2.0	3.5	2.5	2.2	3.8	0.3	3.9	6.2	1.5	100	0.0	0.6	0.0	0.5	0.0	0.1	0.1	0.0	0.0	0.0	0.6	0.2
Quantum	14.9	4.3	8.0	4.3	5.6	2.9	22.4	3.0	5.1	0.2	100	4.3	0.1	74.3	0.1	0.5	0.0	0.0	13.2	0.0	0.6	0.0
Robotics	9.1	2.1	2.1	1.0	2.8	1.7	3.4	1.5	0.8	0.3	0.3	100	0.1	1.8	0.1	1.1	1.9	0.0	1.0	0.0	0.1	0.2
Nanotech	5.6	0.9	2.0	0.2	3.3	0.2	1.2	6.1	0.9	0.3	0.5	9.3	100	3.6	1.1	0.7	3.2	1.9	0.0	0.2	0.0	0.0
IoT	13.7	4.9	13.3	1.7	8.1	1.9	24.9	8.4	6.9	0.4	7.6	2.9	0.1	100	0.0	1.1	0.3	0.0	6.2	0.0	0.8	0.7
CRISPR	2.0	0.2	0.7	2.3	2.5	0.1	0.5	0.1	0.1	0.0	0.2	3.0	0.3	0.0	100	0.2	0.1	0.0	0.2	0.0	0.0	0.0
VR	16.6	1.7	3.2	1.1	3.9	1.3	4.7	2.0	1.4	0.4	0.2	8.8	0.1	5.2	0.1	100	2.7	0.0	2.4	0.1	0.2	0.1
3Dprint	3.4	0.4	0.5	0.3	1.0	0.0	0.9	0.8	0.1	0.3	0.0	16.3	0.3	1.7	0.0	2.9	100	0.4	0.2	0.0	0.0	0.3
Polymer	0.5	0.2	0.0	0.7	0.4	0.0	0.1	0.2	0.0	0.0	0.0	0.7	1.3	0.0	0.0	0.2	2.9	100	0.0	0.0	0.0	0.1
Blockchain	21.4	5.5	11.2	2.3	8.0	5.1	23.1	2.6	4.3	0.1	5.4	6.8	0.0	24.5	0.0	2.0	0.2	0.0	100	0.0	0.8	0.3
Web2.0	1.4	4.0	8.1	2.0	0.8	0.8	21.1	6.5	5.7	0.4	0.0	0.3	0.1	0.2	0.0	0.3	0.1	0.0	0.1	100	2.0	0.0
SOA	3.9	9.1	9.4	0.7	1.8	0.8	26.3	2.1	6.6	1.1	0.1	0.3	0.0	1.8	0.0	0.1	0.0	0.0	0.4	0.3	100	0.0
RFID	0.6	1.0	4.1	0.4	3.5	0.1	2.3	6.0	1.3	0.7	0.0	2.1	0.0	3.9	0.0	0.1	0.3	0.0	0.3	0.0	0.0	100

Figura 8: Sovrapposizione delle tecnologie negli annunci di lavoro (2019)

3.2.3 Co-invenzione e disomogeneità tra imprese

Come anticipato dalla teoria della *Productivity J-curve*, l'AI necessita di investimenti in complementi intangibili per vedere massimizzati i suoi impatti. In particolare, si tratta di competenze, infrastrutture digitali, governance, riorganizzazione aziendale, di processo e di prodotto (argomenti che verranno affrontati nel paragrafo 3.3). La tecnologia, quindi, non genera aumenti di produttività immediati: è necessario del tempo per permettere alle organizzazioni di assorbire i cambiamenti, integrarli e, infine, sfruttarli al meglio per ottenere i benefici sperati.

L'andamento della produttività è inizialmente piatto, se non decrescente. Successivamente vede un incremento, dovuto all'assorbimento e integrazione dei cambiamenti. Si tratta quindi di un fenomeno ciclico, analogo a quello che ha caratterizzato innovazioni nel passato, come quella dell'elettricità, ma anche più recenti, come l'ICT. Tutte queste tecnologie hanno iniziato a dare benefici dopo anni, nel momento in cui l'intero sistema economico ha interiorizzato i complementi necessari per un utilizzo efficace.

In questo quadro, si inserisce il concetto chiave di co-invenzione, introdotto da Bresnahan e Greenstein (1996). Quest'ultima consiste nella combinazione e nell'adattamento di varie GPT in base al contesto, in modo da innescare innovazione continua. Per fare ciò è necessario capitale umano altamente qualificato, dotato di capacità analitiche e strategiche. La velocità di individuazione e l'efficacia dello sviluppo delle co-invenzioni è determinante nella differenziazione dei livelli di produttività tra imprese e, più in generale, sistemi economici. In

definitiva, la co-invenzione è l'ennesimo elemento che si allinea alla teoria sulla produttività ampiamente discussa nel paragrafo. Essa aiuta a spiegare la disomogeneità ed il disallineamento degli effetti dell'intelligenza artificiale, fortemente dipendenti dalle capacità di assorbimento ed integrazione.

3.2.4 Implicazioni macroeconomiche

Le dinamiche osservate a livello aziendale hanno forti ripercussioni su quello che è il sistema economico nel suo complesso. Si evidenziano effetti aggregati che vanno ad incidere su più fronti: crescita, occupazione, distribuzione della ricchezza. Nonostante l'AI abbia un ampio potenziale per rilanciare la produttività, esistono rischi rilevanti riguardanti la polarizzazione dei benefici derivanti: solo una piccola parte di imprese possiede gli strumenti e le competenze per integrare in modo efficace la nuova tecnologia nelle proprie realtà. Secondo l'analisi di McElheran e altri, si tratta per la maggior parte di imprese di grandi dimensioni, dotate di capitale umano e know-how pregresso, che grazie alle loro capacità possono diventare *superstar firms*. Questo fenomeno incide sulle realtà medio-piccole, che purtroppo rischiano di rimanere indietro, facendosi surclassare da imprese più avanzate.

Il meccanismo studiato porta ad un potenziale aumento della concentrazione dei mercati, con una concorrenza sempre più debole ed una disuguaglianza, invece, amplificata. Tutto ciò ricade sul consumatore, che vedrà il suo benessere diminuito.

Tuttavia, se l'AI verrà diffusa a livello globale in modo piuttosto omogeneo, si potranno avere benefici macroeconomici rilevanti, senza andare ad inasprire le disuguaglianze già presenti. Come verrà approfondito nel capitolo successivo, in questi termini è fondamentale l'azione regolatoria degli enti competenti e giocano un ruolo cruciale gli strumenti di policy su cui si può far leva per garantire una diffusione efficace.

In conclusione, l'AI potrà generare un nuovo ciclo di crescita solo se accompagnata da investimenti nei complementi intangibili e da una diffusione ampia e omogenea tra imprese e settori.

3.2.5 Stime e previsioni quantitative sulla produttività

Accanto alle previsioni teoriche, alcuni studiosi hanno provato a fornire stime quantitative in modo da rafforzare le tesi già espresse in letteratura. I numeri forniti dai vari centri di ricerca raramente coincidono: c'è molta eterogeneità nei dati in quanto la tecnologia è ancora in via di diffusione, ma proprio la mancanza di omogeneità può essere utile ad individuare l'intervallo percentuale su ciò che potrà essere l'aumento di produttività globale. Nel paragrafo si analizzano diversi studi, andando in ordine da quelli più ottimistici a quelli più conservativi.

McKinsey (2023) pone particolare attenzione al concetto di automazione, affermando che quest'ultima, se supportata dalla nascente AI generativa, potrebbe far crescere la produttività globale dallo 0,2% al 3,3% annuo. Eliminando il contributo più legato alla robotica e concentrandosi sull'AI in senso stretto, le percentuali calano un po', passando ad un range annuo compreso fra 0,1% e 0,6%.

Baily, Brynjolfsson e Korinek (2023) si focalizzano sul contesto statunitense e si basano sul fatto che l'intelligenza artificiale agirà come GPT. Essi propongono stime in termini di crescita della produttività generale e della produttività dei lavoratori cognitivi. In particolare, se la prima crescesse del 2% ed i lavoratori fossero più produttivi del 20%, la crescita di produttività totale salirebbe al 2,4%. Questo considerando un periodo annuale: se ci si spostasse su un orizzonte temporale più lungo la crescita della produttività godrebbe di meccanismi cumulativi, senza contare che l'AI potrebbe potenzialmente auto-migliorarsi, generando ulteriori benefici.

Goldman Sachs (2023) propone una visione più tradizionale, sempre riferendosi al contesto statunitense: si prevede una crescita della produttività di 1,5 punti percentuali, ipotizzando al contempo un'adozione diffusa dell'AI nell'arco di 10 anni. Negli altri Paesi si prevede un andamento simile, ad eccezione delle realtà più emergenti a causa della maggiore quota di occupazione in settori a bassa esposizione AI, come l'edilizia e l'agricoltura.

Cambiando scenario geografico, il Fondo Monetario Internazionale (IMF, 2024), prendendo in analisi il Regno Unito, suggerisce che l'impatto sulla produttività dipenderà dal futuro evolversi della tecnologia e dal tasso di adozione ma che, in linea di massima, si potrà beneficiare di guadagni in termini di produttività dell'1,5% annui circa.

Spostandosi in Francia, la Commissione AI istituita dal governo francese fa riferimento alle precedenti ondate tecnologiche, quali l'elettricità e l'ICT, e afferma che la crescita in termini di produttività potrà aumentare di 1,3 punti percentuali annui.

Aghion e Bunel (2024) propongono due approcci per stimare la crescita della produttività aggregata, che portano a due risultati differenti. Il primo è basato sulle precedenti rivoluzioni tecnologiche e porta ad un intervallo che si estende da 0,8 a 1,3 punti percentuali, mentre il secondo fornisce una mediana di 0,68 pp.

Altri studiosi, quali Bergeaud (2024) e soprattutto Acemoglu (2024) adottano stime più conservative. Acemoglu, in particolare, afferma che non si verificherà un aumento della produttività totale dei fattori superiore allo 0,71% in 10 anni. Tale percentuale è frutto di calcoli basati sull'esposizione all'AI e sui miglioramenti della produttività a livello di singole attività. La complessità di queste ultime, tra l'altro, può influenzare la percentuale mediana del 0,71% in quanto compiti più difficili godranno di benefici più limitati.

In sintesi, le stime quantitative riportate si dimostrano molto oscillanti: alcuni studiosi presentano valori quasi trascurabili mentre altri sono molto più fiduciosi, proponendo percentuali molto più significative. L'incertezza è alta ed è frutto del processo di sviluppo e diffusione che sta caratterizzando la vita dell'AI.

Tabella 8: Effetti dell'AI a livello produttivo e di crescita

Paper	Tipologia di dato	Effetti trovati
Brynjolfsson & McAfee (2014)	Analisi teorica e storica (<i>The Second Machine Age</i>)	L'AI è vista come GPT a causa del suo potenziale rivoluzionario, paragonabile alla rivoluzione industriale.
Sanjeev et al. (2017)	Studio prospettico	Forte balzo nella produttività degli USA tra il 2028 ed il 2033.
Gordon (2016)	Analisi storica macroeconomica	I guadagni da nuove tecnologie probabilmente saranno inferiori a quelli ottenuti nel secolo scorso da altre ondate tecnologiche.
Brynjolfsson et al. (2019).	Analisi macroeconomica	La diffusione lenta delle nuove tecnologie ne spiega la stagnazione della produttività.
Acemoglu et al. (2014)	Analisi empirica su settori IT-intensive	Produttività più alta nei produttori che nei distributori a causa degli effetti dovuti alla riallocazione dell'occupazione e non all'aumento dell'output.
Brynjolfsson, Rock, Syverson (2021)	Modello teorico	Individuazione delle 3 fasi della J-Curve: costi iniziali, rallentamento, benefici a lungo termine.
McElheran et al. (2024)	Survey su imprese (USA)	Solo il 6% delle imprese adotta AI (18% se ponderato per

		lavoratori). Questo suggerisce un ampio divario tra grandi imprese e PMI.
Goldfarb, Taska, Teodoridis (2020)	Analisi su dataset Burning Glass	AI viene sfruttata al massimo potenziale quando combinata ad altre tecnologie ad essa correlate.
Bresnahan & Greenstein (1996)	Analisi teorico-storica	Co-invenzione: l'AI richiede adattamenti, combinazioni con altri fattori ed investimenti in complementi.
McKinsey (2023)	Report consulenziale	Produttività globale in crescita di 0,2–3,3 pp annui con automazione (considerando AI pura: 0,1–0,6 pp annui).
Baily, Brynjolfsson, Korinek (2023)	Modello teorico	Combinazione tra produttività generale e dei lavoratori cognitivi fino a raggiungere una crescita del 2,4% annuo.
Goldman Sachs (2023)	Modelli econometrici	Produttività in crescita di 1,5 pp annui negli USA, ipotizzando adozione diffusa in 10 anni.
Aghion & Bunel (2024)	Analisi storica e <i>task based</i>	Approccio storico: +0,8–1,3 pp annui. Approccio <i>task-based</i> : mediana 0,68 pp.
Bergeaud (2024)	Modello <i>task based</i>	+2,9% cumulativo in 10 anni (0,29% annuo) con forti differenze a livello geografico.
Acemoglu (2024)	Modello <i>task based</i>	Scenario conservativo con upper bound di +0,71% cumulativo in 10 anni.

IMF (2024)	Analisi macroeconomica (UK)	Crescita della produttività che può raggiungere i 1,5 pp annui, fortemente dipendente da complementarità e grado di adozione.
Commissione AI francese (2024)	Analisi storica comparata	Crescita della produttività fino a 1,3 pp annui.

3.3 Organizzazione aziendale

3.3.1 Trasformazione delle strutture organizzative ed approccio *data-driven*

Come già discusso nel paragrafo precedente, l'intelligenza artificiale è una tecnologia abilitante il cui impatto non è legato solo all'adozione di strumenti predittivi o di automatizzazione, ma anche alla riorganizzazione aziendale profonda che è necessaria per un'adozione efficace e consapevole.

Una delle conseguenze principali dell'AI è la trasformazione delle organizzazioni a livello strutturale: si passa da un modello gerarchico e rigido ad uno più piatto, agile ed adattivo. Le nuove strutture sono più reattive rispetto alle precedenti e, proprio grazie a questa caratteristica, si contraddistinguono per la spiccata abilità di adattarsi ai cambiamenti repentini delle condizioni esterne ed alle informazioni in input. Tutto questo è possibile grazie alla progressiva disintermediazione dei livelli decisionali. Gli algoritmi di ML, i sistemi predittivi avanzati e gli strumenti di automazione cognitiva permettono di delegare molti compiti intermedi alle macchine, liberando forza lavoro che può essere indirizzata su altre mansioni. In questo modo si riduce il bisogno di figure umane nei livelli intermedi della catena gerarchica, ovvero quelli preposti a attività di supervisione, controllo, coordinamento e trasmissione di informazioni. Inoltre, tramite l'adozione di questa struttura, è possibile svolgere sperimentazione continua e garantire un'interazione fluida tra le varie funzioni aziendali.

La delega a sistemi intelligenti è resa possibile dalla grande quantità di dati associata ai modelli di ML ed alle altre tecnologie sopra citate, che permettono all'AI di prendere decisioni informate e consapevoli. I dati, quindi, rappresentano il cuore pulsante del processo decisionale: si tratta di un'organizzazione *data-driven*. Per garantire decisioni ottimali, sia dal punto di vista strategico che operativo e tattico, bisogna rispettare alcuni requisiti:

1. Disponibilità di dati di qualità facilmente accessibili, ma soprattutto affidabili e aggiornati in tempi brevi.

2. I dati devono poter essere interpretati correttamente da coloro che hanno le competenze e abilità adatte per farlo.
3. La cultura aziendale deve essere fondata su valori precisi: oggettività, evidenza, apprendimento continuo. Bisogna essere disposti a lavorare ed investire continuamente per garantire miglioramento.

In sintesi, la riorganizzazione aziendale costituisce un elemento imprescindibile, che deve procedere di pari passo con l'adozione dell'AI.

3.3.2 Modularità e interdipendenze

La trasformazione della struttura organizzativa da un modello gerarchico ad uno piatto non è l'unica ad essere coinvolta nel processo di adozione dell'AI. Infatti, la nuova ondata tecnologica è destinata a influenzare radicalmente anche la modularità aziendale.

Un'organizzazione è un sistema interdipendente di decisioni, quindi nell'analisi della performance legata all'adozione dell'intelligenza artificiale non andrebbero considerate le singole attività, bensì l'impresa nel suo complesso. Quest'ultima tipicamente può essere divisa in moduli differenti che comunicano ed interagiscono fra loro. L'interdipendenza fra i diversi moduli è strettamente correlata all'AI, in quanto l'utilizzo di quest'ultima potrebbe richiedere una ristrutturazione modulare.

Lo studio di Agrawal, Gans e Goldfarb (2024) si occupa di studiare come varia l'efficacia della nuova tecnologia in base alle interdipendenze riscontrate fra le varie decisioni aziendali. Gli autori elaborano un modello basato su due decisioni, D1 e D2. La prima è quella principale e dipende da fattori esterni, come ad esempio la domanda del mercato. Se presa correttamente genera un beneficio α . La seconda è una decisione di supporto e dipende dall'allineamento con D1. Se c'è corretto allineamento essa genera un beneficio γ . Se entrambe D1 e D2 sono corrette si ha un beneficio totale β che è maggiore della somma di $\alpha + \gamma$. Il premio sinergico è quindi dato da:

$$\delta = \beta - (\alpha + \gamma).$$

Nella pratica, se si prende come esempio un contesto retail, D1 può essere la raccomandazione di un prodotto e D2 la presenza in magazzino del prodotto raccomandato. L'interdipendenza fra le decisioni è modellata con il parametro μ , compreso fra 0 e 1. Un basso valore di μ indica decisioni più modulari ed indipendenti, mentre un alto valore indica una forte interdipendenza e quindi un'alta probabilità che un errore in D1 comprometta anche D2.

In assenza d'AI, il decisore di D1 sceglie sempre l'azione focale, ovvero quella che è corretta con più probabilità. Il responsabile di D2, di conseguenza, riuscirà ad allinearsi senza problemi

e soprattutto senza necessità di comunicazione, in quanto sa a priori quale sarà D1. Si tratta quindi di un equilibrio di Nash, in cui il ritorno per l'organizzazione è dato da

$$R = (1 - \mu)(p\alpha + \gamma) + \mu p\beta,$$

dove p indica la probabilità che D1 sia corretta.

In presenza d'AI la situazione si complica in quanto il decisore sceglie D1 in modo più preciso, affidandosi ad algoritmi predittivi che permettono di avere informazioni corrette e puntuali sullo stato esterno. D1 diventa quindi variabile e l'agente che è responsabile di D2, comportandosi come se D1 fosse scelta come in assenza di AI, commetterà un errore quando D1 si scosterà dalla decisione focale. Il ritorno complessivo diventa:

$$R = (1 - \mu)(\alpha + p\gamma) + \mu p\beta,$$

in cui questa volta p indica la probabilità che D2 sia allineata a D1.

In conseguenza al meccanismo descritto, l'intelligenza artificiale verrà adottata solamente se il contributo α supera la mancanza di allineamento, rappresentata dalla perdita di γ , ovvero se la decisione autonoma è più importante del sincronismo. Quindi, in definitiva, si tratta di valutare il peso relativo della correttezza della decisione core D1 e del coordinamento fra D1 e D2. Inoltre, per migliorare la performance complessiva in caso di adozione di AI, si può valutare l'introduzione di comunicazione fra i due agenti. Il costo associato può essere variabile, se legato a ogni singolo messaggio inviato, o fisso, se si tratta, ad esempio, dell'intera infrastruttura IT. In quest'ultimo caso, modellando il costo fisso come C , il payoff complessivo diventa:

$$R = (1 - \mu)(\alpha + \gamma) + \mu\beta - C.$$

Per quanto riguarda la modularità, essa può essere soggetta a cambiamenti per favorire la buona riuscita dell'adozione della nuova tecnologia. Infatti, un'architettura più modulare implica minori interdipendenze fra le attività e, di conseguenza, una minore necessità di coordinamento e comunicazione. Se, d'altra parte, i benefici sinergici catturati dal parametro β sono molto alti, conviene rimodellare la struttura in modo da aumentare μ . Amazon è un esempio di azienda fortemente modulare, con μ basso e senza bisogno di un forte coordinamento interno: il motore di raccomandazione suggerisce solo prodotti già presenti in magazzino. Queste caratteristiche rendono la piattaforma il prototipo perfetto per l'utilizzo di macchine ed algoritmi intelligenti. In conclusione, l'adozione dell'AI in sistemi modulari è facilitata, mentre per quanto riguarda sistemi più integrati e caratterizzati da forti interdipendenze è necessario investire in comunicazione e coordinamento per poter sfruttare i benefici sinergici al massimo del loro potenziale.

3.3.3 Capitale umano e complementarità

L'adozione dell'AI non è un processo *plug and play*. La tecnologia, se considerata in modo isolato, non è in grado di generare valore duraturo: è necessario che sia accompagnata da adattamenti organizzativi. Per raggiungere l'efficacia del deployment è necessario sviluppare, assorbire e adattare continuamente complementi intangibili. L'AI richiede la formazione di capitale umano qualificato in quanto quest'ultimo diventa responsabile di routine aziendali, governance ed investimenti in infrastrutture digitali. Le capacità manageriali avanzate risultano essenziali per l'efficacia della nuova ondata tecnologica. In questo contesto tornano alla luce anche i concetti di ITIC e co-invenzione (Bresnahan & Greenstein, 1996).

Le principali frizioni che accompagnano la nuova tecnologia sono di tipo organizzativo, tecnico e culturale. Concentrandosi principalmente sulle prime, risulta essenziale: riorganizzare processi interni e prodotti, ridefinire le responsabilità decisionali, aggiornare la configurazione dei flussi informativi e modificare l'interfaccia tra uomo e macchina. Nel quadro complesso che si viene a creare si va quindi ad affermare un processo di co-invenzione, in cui la GPT si lega continuamente in modo dinamico ad altre tecnologie in modo da potere garantire la corretta riuscita delle trasformazioni necessarie.

Per riuscire nell'intento è necessario sviluppare competenze dinamiche e flessibilità. Gli step principali da affrontare per ogni campo soggetto a trasformazione sono: l'analisi dei processi o prodotti, l'identificazione delle criticità, la sperimentazione di nuove soluzioni e, infine, la ristrutturazione. Le aziende che eccellono in queste fasi sono tipicamente quelle che riescono a prosperare ed a guadagnare un vantaggio competitivo sostenibile. Per riuscire a raggiungere tale obiettivo bisogna possedere competenze digitali, manageriali e strategiche, cultura organizzativa, orientamento ai dati e approccio agile. L'ITIC, risultato dell'elenco di caratteristiche appena riportate, consente alle organizzazioni di estrarre il massimo valore strategico.

L'AI, quindi, porta con sé numerosi problemi di installazione e frizioni, che possono essere superate solamente attraverso l'utilizzo di tecnologie complementari, la presenza di dati e di forza lavoro competente in grado di ristrutturare l'organizzazione sia internamente che nel suo complesso, a partire dalle routine operative fino alle logiche di governo strategico.

3.3.4 Cambiamento dei processi decisionali e nuovi ruoli organizzativi

L'emergere dell'intelligenza artificiale e la sua iniziale adozione da parte delle imprese ha messo in luce il profondo cambiamento a cui saranno sottoposti i processi decisionali, il quale a sua volta porterà ad una ristrutturazione dei ruoli organizzativi, nonché ad una riconfigurazione delle funzioni aziendali.

I sistemi intelligenti sono in grado di fare predizioni accurate e precise se alimentati da ampi dataset di informazioni di qualità, affidabili ed aggiornate. I modelli decisionali tradizionali, che prima erano fondati sull'esperienza accumulata, le soft skills e l'intuizione manageriale, vengono quindi affiancati o sostituiti dall'AI. In questo contesto riemerge il dibattito riguardante l'effetto sostituzione o complementarità in termini occupazionali: l'uomo potrà mantenere il controllo decisionale diretto, appoggiandosi al ML e alle altre tecnologie solo come supporto, oppure potrà vedersi totalmente rimpiazzato. In uno stadio iniziale, come quello in cui si trova ora l'intera ondata tecnologica, non si vedrà una completa sostituzione ma piuttosto un affiancamento uomo-macchina, quindi diventa essenziale sviluppare complementarità fra le due parti.

Il management, in conseguenza al cambiamento generale, vede una trasformazione del suo ruolo. Nel quadro complesso creatosi, esso sarà responsabile non tanto del processo decisionale in sé quanto dell'interpretazione dei risultati generati in output dagli algoritmi intelligenti e nella supervisione degli algoritmi stessi. Ciò si rivela necessario in quanto il ML è soggetto a bias (trattati in maniera approfondita nel capitolo 4) che potrebbero far prendere decisioni errate o rinforzare idee socialmente sbagliate. Una volta che il management si è assicurato della bontà dei sistemi intelligenti e dell'output, integrerà le informazioni ricevute nei processi strategici e operativi.

La trasformazione del ruolo del management porta con sé la costruzione di team eterogenei, i quali saranno caratterizzati da una leadership attiva, capace di rimanere al passo con l'innovazione in modo da potersi migliorare continuamente. La costituzione di squadre multidisciplinari permette di affrontare le sfide imposte dall'AI in modo integrato e con competenze che coprono varie discipline. Si osserva, infatti, la costruzione di team formati da personale con indirizzo tecnico, come data scientist ed ingegneri del software, affiancati da esperti legali o professionisti HR. Nei modelli tradizionali queste figure erano spartite in differenti task force, mentre ora collaborano ed interagiscono fra loro in una struttura più dinamica e flessibile, confermando la sfumatura dei confini funzionali.

3.3.5 Sfide organizzative dell'AI nelle imprese

L'introduzione di una nuova tecnologia all'interno delle imprese non è mai un processo facile ed immediato. Le barriere che si riscontrano nella maggior parte dei casi, come conferma anche la precedente ondata legata all'ICT, sono: la mancanza di competenze digitali adeguate e di dati di qualità in grandi dimensioni, un approccio non sufficientemente analitico e, soprattutto, la resistenza culturale alla trasformazione profonda.

In particolare, per quanto riguarda l'AI, numerose realtà stanno iniziando a integrare i sistemi intelligenti nel recruiting e nel settore logistico. Per quanto riguarda il primo, gli algoritmi sono in grado di analizzare e selezionare i CV, condurre screening preliminari e valutare i candidati attraverso lo studio automatico di interviste. Queste attività, però, non possono essere svolte dal ML senza la supervisione dell'uomo in quanto i sistemi intelligenti, come già anticipato, soffrono di bias, legati sia ai dati che agli algoritmi, che potrebbero portare a scartare candidati promettenti o ad assumere persone non idonee o non in linea con gli obiettivi aziendali.

Nel settore logistico, ci si riferisce principalmente alla previsione della domanda, alla gestione dei magazzini e delle scorte ed all'ottimizzazione dei flussi e delle rotte per quanto riguarda i trasporti. Per poter costruire un sistema efficace è necessaria un'adeguata comunicazione fra le varie funzioni aziendali, nonché la presenza di dati in grandi quantità e soprattutto di qualità, in modo da poter fare previsioni affidabili.

In definitiva, l'adozione dell'intelligenza artificiale è strettamente legata a una profonda riconfigurazione della governance, delle pratiche decisionali e dell'organizzazione stessa nella sua interezza. In assenza di tali trasformazioni anche le innovazioni più promettenti rischiano di scontrarsi con barriere così grandi da comprometterne lo sfruttamento, o addirittura di impedirne il deployment sin dalle fasi iniziali.

Tabella 9: Effetti dell'AI a livello organizzativo

Paper	Tipologia di dato	Effetti trovati
Agrawal, Gans & Goldfarb (2024)	Modello formale	L'efficacia dell'AI dipende dal grado di interdipendenza. Nei sistemi integrati servono investimenti in comunicazione per sfruttare al meglio le sinergie.
Bresnahan & Greenstein (1996)	Analisi teorico-storica	Co-invenzione: l'AI richiede adattamenti, combinazioni con altri fattori ed investimenti in complementi.

3.4 Innovazione e modelli di business

3.4.1 Piattaforme digitali

Le piattaforme digitali hanno un ruolo centrale nel paradigma economico che si sta affermando con l'avvento dell'intelligenza artificiale e, proprio a causa della loro importanza, stanno vivendo un periodo di forte cambiamento.

Esse sono nate come sistemi di intermediazione pura, in grado di mettere in contatto più gruppi diversi di persone e, proprio per questa peculiarità, la letteratura economica le definisce *two-sided markets* (Rochet & Tirole, 2003; Armstrong, 2006), ovvero mercati a due (o più) lati. I meccanismi caratteristici delle piattaforme permettono: l'incontro tra domanda e offerta, la riduzione dei costi di ricerca e di transazione, la formazione di una vera e propria rete di utenti. Quest'ultima è uno dei punti chiave, in quanto il beneficio in termini di utilità associato alla partecipazione in una piattaforma è legato alla presenza di altri utenti.

Le piattaforme permeano in settori molto diversi fra loro, basti pensare agli esempi classici di Airbnb, Amazon o Facebook: ognuna di esse ha obiettivi diversi e aiuta al superamento di barriere di differente tipo. Airbnb si occupa di far incontrare la domanda e l'offerta nel settore immobiliare, mettendo in comunicazione proprietari e affittuari, Amazon permette a venditori e consumatori di interagire su scala globale, Facebook è orientata al mondo social e, quindi, alla condivisione di contenuti. In tutti i casi, la piattaforma svolge il ruolo di *market maker*, ovvero facilita lo scambio.

In questa prima fase di sviluppo delle piattaforme, i modelli di business erano fondati su pubblicità o commissioni di intermediazione, mentre la capacità chiave che generava innovazione era l'offerta di un punto di contatto affidabile fra più gruppi di persone, interessati reciprocamente dagli altri, in grado di abbattere barriere informative.

3.4.2 AI come fonte di personalizzazione

Con l'affermarsi delle tecnologie relative all'apprendimento automatico, le piattaforme hanno superato la mera funzione di intermediazione, assumendo un ruolo più attivo nella gestione delle interazioni: esse diventano parte integrante nelle scelte di vendita e di consumo. La grande quantità di dati legata ai modelli di intelligenza artificiale permette agli algoritmi di *machine learning* di analizzare in tempo reale enormi quantità di dati forniti in input dagli utenti. Questo permette di creare raccomandazioni, prezzi dinamici e contenuti personalizzati.

Un esempio emblematico di tale processo trasformativo è Netflix, piattaforma che nasce per il noleggio di videocassette e DVD. Negli anni ha subito un'evoluzione che le ha permesso di diventare uno dei principali colossi globali del servizio di streaming. L'uso sistematico

dell'intelligenza artificiale le ha consentito di sviluppare sofisticati algoritmi di raccomandazione, i quali suggeriscono contenuti diversi ad ogni singolo utente, in base alle sue recensioni o alle serie TV o film già visti in precedenza. La piattaforma, quindi, non si limita più a veicolare contenuti, ma orienta direttamente le decisioni di consumo, diminuendo i costi di ricerca e prolungando la permanenza degli utenti nel proprio ecosistema.

Grazie ai modelli ML si crea quindi un'esperienza altamente personalizzata, e ciò può essere positivo o negativo. Se da un lato l'utente può beneficiare della presenza di contenuti di suo interesse, riducendo i tempi ed i costi di ricerca, dall'altro si possono avere problemi come il *filter bubble* o l'*echo chamber*. Questi sono fenomeni per cui si rafforzano le opinioni pregresse dell'utente grazie all'esposizione continua a contenuti altamente allineati ai suoi ideali ed interessi. In questo modo la persona vede solo ciò in cui crede e interagisce solo con utenti del suo stesso tipo, evitando così di uscire dai suoi schemi per incontrare nuove vie di pensiero o interesse, riducendo la possibilità di esplorare alternative.

Dal punto di vista economico si osserva una transizione nei modelli di ricavo: la pubblicità e le commissioni di intermediazione vengono affiancate a logiche più complesse, basate sulla monetizzazione della personalizzazione. Le piattaforme, così, diventano curatori delle interazioni e non solo più intermediari.

3.4.3 Il futuro dell'AI generativa

Il nuovo orizzonte delle piattaforme è segnato dall'avvento dell'AI generativa e dei *foundation models* (FMs). Essi sono sistemi addestrati su grandi quantità di dati multimodali, capaci di elaborare informazioni, fare previsioni e soprattutto generare nuovi contenuti originali. Quest'ultima è la caratteristica che permette una svolta nei modelli di business delle piattaforme: esse non sono più semplici intermediari o curatori, bensì diventano produttori di valore creativo.

Le piattaforme che si stanno specializzando in questi termini si appoggiano sul modello di distribuzione *AI as a service*: un provider di servizi ospita le applicazioni, rendendole accessibili agli utenti spesso in cambio di un abbonamento, come ad esempio ChatGPT Plus di OpenAI. L'alternativa è l'utilizzo di API a consumo integrate in applicazioni di terzi, come accade per Microsoft Azure. Questi due strumenti permettono alle organizzazioni di interfacciarsi e sfruttare i modelli di AI generativa e gli output da essi prodotti senza dovere sviluppare in house infrastrutture costose e competenze specialistiche, in modo da potere rimanere concentrati sulle attività core del proprio business.

In definitiva, le piattaforme che si basano sull'utilizzo dell'AI generativa stanno diventando portatrici di un processo innovativo distribuito in quanto gli utenti stessi possono diventare co-

creatori: interrogano i modelli fornendo prompt, possono dare feedback per il miglioramento, influenzano gli output. Da funzioni di intermediazione e personalizzazione, ci si evolve verso ecosistemi creativi e iterativi, in cui algoritmi e comunità di utenti cooperano in un processo di sperimentazione continua.

Tabella 10: Effetti dell'AI su innovazione e modelli di business

Paper	Tipologia di dato	Effetti trovati
Rochet & Tirole (2003)	Modello teorico	Definizione dei <i>two-sided markets</i> come intermediari che riducono costi di transazione e creano effetti di rete.
Armstrong (2006)	Modello teorico	Analisi di pricing e strategie dei <i>multi-sided markets</i> , con particolare focus sull'equilibrio creato tra domanda ed offerta su vari fronti.

3.5 Educazione

3.5.1 Trasformazione della domanda di competenze e AI literacy

L'adozione dell'intelligenza artificiale non si limita a produrre effetti sul lavoro, sulla produttività, sull'organizzazione aziendale e sull'innovazione, ma genera profondi cambiamenti anche su campi meno economici, come educazione e formazione. Come conseguenza diretta degli effetti di complementarità e sostituzione, ampiamente discussi nel paragrafo 3.1, nasce l'esigenza di nuove competenze ed abilità, che dovranno essere acquisite tramite corsi, per coloro già inseriti nel mondo lavorativo e che quindi potrebbero vedersi sostituiti, totalmente o solamente in alcune attività, dai sistemi intelligenti, o tramite un percorso formativo vero e proprio, per quanto riguarda gli attuali studenti e le nuove generazioni. L'*AI literacy*, alfabetizzazione all'intelligenza artificiale, è un concetto chiave in quanto comprende l'abilità di comprendere, utilizzare, monitorare e riflettere criticamente sulle applicazioni di AI. Il modello *task based* (Brynjolfsson & Mitchell, 2017) afferma che la sostituzione potenziale dei lavoratori non avviene a livello di occupazione ma di compiti, quindi la chiave sta nello sviluppare competenze ed abilità in tutti i ruoli non-SML, ovvero non adatti ad essere svolti da macchine. L'ambito di intervento della formazione è duplice, in quanto il focus non è da porre solo sulle abilità tecniche ma anche in quelle umanistiche.

La crescente pervasività delle tecnologie di AI richiederà la presenza di lavoratori altamente qualificati specializzati negli ambiti STEM (*science, technology, engineering, mathematics*). La necessità di tali figure aveva già accompagnato l'ondata tecnologica legata all'ICT, in quanto era necessario l'apprendimento e lo sviluppo di *skills* legate alla digitalizzazione ed alla trasformazione digitale. La pervasività del ML e di tutti gli algoritmi ad essa associati aumenterà ancora di più la domanda per tali competenze, tanto che sarà opportuno avviare sistemi di protezione sociale intelligenti e dinamici, in grado di attuare processi di *reskilling* di figure, in modo da poter aggiornare o reinventare il proprio profilo professionale. Per quanto riguarda, invece, la formazione di studenti non ancora inseriti in un contesto lavorativo, sarà necessario rivedere in profondità i metodi di insegnamento e i contenuti didattici affrontati sin dalle scuole medie e superiori, per poi proseguire con corsi di specializzazione nelle università. I sistemi educativi saranno chiamati a integrare corsi su tematiche emergenti come l'intelligenza artificiale, l'apprendimento automatico e la sicurezza dei dati, nonché su discipline come data science, programmazione, statistica e logica computazionale, in modo da garantire un allineamento efficace tra percorsi educativi e bisogni occupazionali.

3.5.2 Competenze umanistiche e nuove figure professionali

Limitarsi allo sviluppo di competenze tecnico-scientifiche, però, risulterebbe miope. Molti studiosi, come già esplicitato, temono il rischio di sostituzione o, addirittura, prevedono la possibilità di una singolarità economica (Lu e Zhou, 2021), in cui l'AI potrebbe giungere a un livello di sviluppo così avanzato da rendere automatizzabile una quota importante della produzione di conoscenza e di ricchezza. In tale scenario ipotetico, la centralità dell'essere umano non deriverebbe più dal contributo operativo diretto, ma dalla sua capacità di interpretare i risultati, orientare i processi innovativi e salvaguardare la dimensione umana all'interno delle scelte decisionali. In questo contesto, quindi, le competenze umanistiche potrebbero guadagnare valore, in quanto saranno il campo in cui le macchine non potranno superare l'uomo, richiedendone così il supporto. Anche in contesti meno estremi, in cui l'AI potrà sostituire l'uomo ma senza portare ad una singolarità economica, le abilità umanistiche saranno comunque soggette a rafforzamento in quanto diventeranno cruciali per affrontare dilemmi legati alla privacy, alla trasparenza e alla giustizia sociale, oltre che per governare l'interazione uomo-macchina ed interpretare le decisioni algoritmiche. Filosofia, psicologia, diritto e sociologia acquisiranno un nuovo peso. La vera sfida consisterà nel superare la dicotomia tra STEM e *humanities*, promuovendo lo sviluppo di un approccio collaborativo e di competenze ibride che integrino pensiero critico, creatività e problem solving, come suggerito dall' Organizzazione per la Cooperazione e lo Sviluppo Economico. Questa necessità è

rappresentata dall’emergere di figure, le cui competenze combinano elementi tecnici, creativi e strategici, come quelle di:

- *Prompt engineer*, specialista nell’ottimizzare i prompt con l’obiettivo da ottenere risposte pertinenti, precise e coerenti con i bisogni dell’utente o dell’organizzazione.
- *AI ethicist*, professionista che si occupa di analizzare i rischi etici connessi all’utilizzo dell’AI in modo da garantire i principi di equità, trasparenza e responsabilità sociale.
- *AI product manager*, responsabile della gestione di cicli di vita di prodotti *AI-based*.
- *Data steward*, figura che lavora a stretto contatto coi dati con lo scopo di garantirne accuratezza, coerenza, sicurezza e disponibilità.
- *AI operations specialist*, tecnico specializzato nella manutenzione, aggiornamento continuo e governance dei sistemi di AI in esercizio.

In definitiva, l’educazione merita una particolare attenzione nelle politiche pubbliche e aziendali in quanto tramite un ripensamento della formazione e delle competenze, oltre che alla nascita di nuove figure, potrebbe permettere di arginare il rischio di disuguaglianza, disoccupazione e polarizzazione.

Tabella 11: Effetti dell’AI su educazione e formazione

Paper	Tipologia di dato	Effetti trovati
Brynjolfsson & Mitchell (2017)	Analisi su database O*NET	Introduzione del concetto di SML. Poche occupazioni sono totalmente automatizzabili ma la maggior parte hanno almeno alcuni compiti che lo sono.
Lu e Zhou (2021)	Analisi teorico-prospettica	Analisi dello scenario di singolarità economica, in cui le <i>humanities</i> diventerebbero le competenze chiave per il successo dell’uomo.
OECD (2021-2023)	Report istituzionali	Necessità di superare la dicotomia tra STEM e <i>humanities</i> in modo da valorizzare competenze ibride

		che possano sostenere nuove figure professionali.
--	--	---

3.6 Concorrenza e struttura dei mercati

3.6.1 Effetti pro-competitivi

Uno degli aspetti più discussi, riguardante l'impatto che l'AI avrà sull'economia, è quello che pone il focus sulla concorrenza e sulla struttura dei mercati. La letteratura non è ancora concorde sulla previsione degli effetti attesi e propone diverse visioni.

Un primo filone (Goldfarb et al., 2020) afferma che, in quanto GPT, l'AI potrà diffondersi in vari settori e diventare portatrice di innovazione, generando un aumento della concorrenza ed un beneficio collettivo in termini di mercato, come avvenuto ad esempio per Internet nel passato: la tecnologia, dopo una prima fase di diffusione, è riuscita ad affermarsi in modo omogeneo, senza determinare squilibri concorrenziali o monopoli. Al contrario, è riuscita a ridurre costi di transazione e di ricerca, abbassando le barriere all'entrata. In un contesto così favorevole molte sono state le imprese che si sono affermate nel mercato, sfruttando le opportunità create da Internet.

Per l'AI si prevede lo stesso trend: la diffusione degli algoritmi predittivi e la riduzione dei costi di informazione, coordinazione, replicazione, trasporto, tracciamento e verifica (Goldfarb & Tucker, 2019) su larga scala permetterà di avere un mercato più contendibile. In questo modo gli entranti potranno sfidare gli incumbent con minori barriere e, come conseguenza, i consumatori beneficeranno di prezzi più competitivi e di una maggiore varietà di prodotto. In questi termini gioca un ruolo chiave l'adozione del *dynamic pricing*, meccanismo che consente un aggiustamento continuo tra domanda e offerta e che, di conseguenza, garantisce l'equilibrio nel mercato evitando eccessi di offerta o domanda insoddisfatta. Esso si basa su una moltitudine di fattori, quali la disponibilità di stock, i vincoli di capacità, i prezzi dei concorrenti e le oscillazioni della domanda. Tuttavia, questo metodo potrebbe generare confusione nei consumatori, che si troverebbero a far fronte a variazioni di prezzo continue.

Per quanto riguarda la predizione come principale utilizzo degli algoritmi di ML, alcuni settori che stanno mostrando guadagni in termini di efficienza sono quello assicurativo e quello finanziario. I sistemi intelligenti possono offrire supporto all'uomo nelle decisioni di portfolio, suggerendo quali titoli tenere, vendere o comprare, in base alle informazioni disponibili in tempo reale. Nel campo assicurativo, invece, possono aiutare nella valutazione della rischiosità dei clienti, nella formulazione di offerte automatiche e nella gestione dei sinistri: *The Economist*

(2017) riporta che un assicurato può ricevere il rimborso tre secondi dopo aver presentato la richiesta sull'app. Il piccolo arco di tempo indicato è sufficiente per la macchina per esaminare la pratica, eseguire 6618 algoritmi antifrode, approvare il rimborso, inviare le istruzioni di pagamento alla banca ed informare il cliente. Inoltre, l'AI può innescare anche efficienza dinamica, promuovendo innovazione continua grazie alla diffusione ampia della tecnologia. Infine, gli algoritmi possono agire a favore del consumatore, riducendo le asimmetrie informative e rendendolo consapevole di altre dimensioni oltre al prezzo, come ad esempio la qualità di prodotto.

3.6.2 Effetti anti-competitivi

Purtroppo però, lo scenario non è semplice e ci sono difficoltà nell'adottare sistemi intelligenti: i modelli di ML e in generale l'AI, necessitano di dati, potenza di calcolo e risorse finanziarie ingenti. Con questa consapevolezza, si afferma un'altra scuola di pensiero, la quale sostiene che l'emergere dei sistemi intelligenti porta con sé il grande rischio dell'incremento del potere di mercato di alcune imprese, le cosiddette *superstar firms*. Come già spiegato nei paragrafi precedenti, è sempre più condivisa l'idea secondo cui le imprese che si muovono prima nel mercato dell'AI e che possiedono complementi in termini di asset intangibili (ITIC) o di competenze pregresse nel mondo digitale, godranno di un vantaggio competitivo se comparate ad altre aziende meno avanzate grazie ai rendimenti crescenti derivanti da economie di scala e *network effects*. Questo non vale solamente per coloro che si occupano di processori e si inseriscono quindi nella catena produttiva dell'AI, dove un ruolo significativo è ricoperto da Nvidia, ma si estende anche ad altri grandi leader, come ad esempio Google, Microsoft, Oracle e Meta, i quali potranno consolidare ancora di più la loro posizione dominante. I colossi citati hanno quotazioni con guadagni superiori all'indice aggregato *Standard & Poor's 500* e stanno avviando sempre più partnership con l'obiettivo di espandersi nel mercato AI.

Inoltre, le *superstar firms* sono facilitate nell'aumentare rapidamente il loro potere grazie agli effetti di rete ed agli *switching costs*: più cresce il numero di utenti che adottano un servizio di AI, più altri utenti vengono attirati. Questo porta ad un utilizzo più intenso da parte di questi ultimi e quindi alla raccolta di più dati individuali (apprendimento *within-users*) ma non solo: tale meccanismo permette alle piattaforme di raccogliere più informazioni da più persone (apprendimento *across-users*). Con la grande quantità di dati disponibili, successivamente, si va ad aumentare la qualità del servizio offerto. Si innesca così un processo di miglioramento continuo dei prodotti e servizi offerti dai big player, che porta ad un vantaggio cumulativo che rischia di eclissare le piccole realtà emergenti portando al consolidamento del potere degli incumbent (Hagiwara & Wright, 2020). Tale posizione dominante è favorita anche dal

funzionamento dei motori di ricerca: l'accesso agli storici degli utenti più lunghi permette loro un miglioramento più rapido dei risultati rispetto a concorrenti altrettanto efficienti (Schafer & Sapi, 2019). In ultima istanza, uno degli elementi più importanti che consente il rafforzamento del potere di mercato è la discriminazione di prezzo, trattata nel dettaglio nel paragrafo seguente.

3.6.3 Discriminazione di prezzo e potere di mercato legato ai dati

Sempre più spesso le decisioni di pricing sono delegate ad algoritmi di ML. Questi ultimi sono abili nel riconoscere schemi in dataset di grandi dimensioni e ciò consente un targeting e una segmentazione più raffinati, aumentando di molto le potenzialità della discriminazione di prezzo, ovvero dell'applicazione di prezzi differenti per diversi tipi di consumatori, in base alle loro preferenze, utilità, disponibilità a pagare e comportamenti passati. In termini di potere di mercato, quanto più un'impresa dispone di dati e di capacità computazionale, tanto più sarà in grado di segmentare i consumatori ed estrarre surplus, consolidando la posizione dominante degli incumbent.

La discriminazione di prezzo può seguire diversi modelli e, di conseguenza, si riconoscono 3 categorie:

- Discriminazione di prezzo di primo grado, in cui il prezzo diventa personalizzato per ciascun individuo. È l'ideale teorico in cui l'impresa conosce perfettamente la *willingness to pay* di ciascun consumatore.
- Discriminazione di prezzo di secondo grado, legata alla quantità acquistata: propone sconti all'aumentare delle unità.
- Discriminazione di prezzo di terzo grado, la quale applica prezzi diversi a consumatori appartenenti a segmenti diversi (ad esempio “clienti fedeli” vs “nuovi clienti”). In questo caso è essenziale riuscire ad analizzare approfonditamente la base clienti, in modo da individuare trend che permettano la segmentazione del mercato.

Prima dell'affermarsi delle tecnologie AI, il raggiungimento della discriminazione di prezzo di primo grado era molto difficile e complesso; il modello predominante era il terzo. Con il progresso tecnologico, lo sviluppo di algoritmi e la grande attenzione volta ai dati sono stati fatti passi in avanti, al punto che i sistemi intelligenti sono ormai in grado non solo di segmentare in modo molto più accurato ma addirittura di studiare ogni singolo consumatore proponendogli un prezzo personalizzato.

In aggiunta alle classiche tre categorie sopra riportate, si afferma la *behaviour-based price discrimination* (BBPD), in cui il prezzo è fissato sulla base dei comportamenti passati di

acquisto dei consumatori (Ezrachi & Stucke, 2016). In questo contesto si utilizzano in modo massiccio dati dinamici e raccolti nel tempo, con l'obiettivo di distinguere tra consumatori nuovi e ricorrenti e per stimare la loro propensione a cambiare fornitore. È uno scenario diverso rispetto alla tradizionale discriminazione di terzo grado, in quanto quest'ultima è basata su caratteristiche osservabili e non sullo studio di ampi dataset riguardanti dati personali per l'individuazione di emozioni, bias o disponibilità massima a pagare. Nel modello esplicativo della BBPD si agisce in due periodi. Nel primo i consumatori effettuano l'acquisto di un prodotto o un bene e l'impresa registra informazioni a riguardo, come ad esempio la frequenza di consumo, gli item acquistati, in che quantità e a che prezzo. Nel secondo stadio l'impresa sfrutta i dati collezionati per proporre prezzi differenti a diverse persone: è possibile che ci sia una riduzione dei prezzi per coloro che nel primo stadio avevano effettuato l'acquisto da un concorrente, in modo tale da attirarli nel proprio business e scatenando la cosiddetta "guerra per i clienti". Per quanto concerne, invece, i propri clienti si possono adottare due comportamenti in base alla loro fedeltà e mobilità: coloro che difficilmente cambieranno fornitore a causa della alta fidelizzazione o dei costi di *switching* potrebbero venire penalizzati, in quanto sono caratterizzati da domanda piuttosto rigida a variazioni di prezzo; al contrario, coloro che sono considerati "a rischio di fuga" verso i concorrenti (perché i loro costi di *switching* sono bassi o perché hanno già mostrato in passato propensione a confrontare alternative) verranno premiati con sconti. Si osserva, quindi, una redistribuzione del surplus dei consumatori, che viene trasferito da quelli più fedeli a quelli più mobili o ai nuovi clienti.

In termini di concorrenza e potere di mercato la letteratura ha evidenziato come la BBPD e, più in generale, la discriminazione di prezzo, possano avere un effetto ambivalente. In scenari di simmetria informativa c'è un'intensificazione della competizione in quanto l'analogo accesso ai dati da parte delle imprese e la possibile adozione delle medesime strategie permette di anticipare la mossa del rivale, spingendole ad abbassare i prezzi. In questo meccanismo si rivede la configurazione del dilemma del prigioniero: le imprese beneficerebbero collettivamente dalla rinuncia alla discriminazione, ma poiché la strategia rappresenta la miglior risposta alla mossa del concorrente, entrambe sono incentivate ad adottarla comunque. Se si prende in considerazione, invece, lo scenario caratterizzato da asimmetria informativa, la situazione cambia. Nel breve termine si può osservare un incremento della competizione, con l'intensificarsi di strategie di rivalità e la spinta delle imprese a ridurre i prezzi. Tuttavia, nel lungo periodo il rischio che le imprese che dispongono di dataset completi o più accurati rispetto ai rivali si impossessino del potere è alto in quanto diventa cruciale il ruolo svolto dai dati, dalle infrastrutture digitali e dalle risorse computazionali. Solo chi ha un vantaggio competitivo in

termini di tali asset riuscirà a guadagnare vantaggio cumulativo difficilmente replicabile da altre imprese. Inoltre, con un numero maggiore di imprese in grado di applicare la pricing dinamica, coloro che dispongono di un dataset storico più ampio, come un motore di ricerca ben consolidato o una piattaforma digitale, possono affinare i loro algoritmi in modo più rapido ed efficiente, aumentando così il loro vantaggio competitivo sugli altri (Schäfer & Sapi, 2020). Di conseguenza la discriminazione di prezzo diventa una leva per ampliare il divario, contribuendo al rischio di *tipping* del mercato verso un leader dominante. Per i consumatori ciò si traduce in un equilibrio fragile in quanto in alcuni casi possono anche beneficiare di prezzi ridotti, ma spesso a scapito della tutela della privacy e con un'esposizione crescente allo sfruttamento dei propri dati personali.

Il terzo ed ultimo scenario individuato dalla lettura riguarda il ruolo degli intermediari dei dati, soggetti terzi che raccolgono, aggregano e rivendono dataset alle imprese. Essi rivendono i dataset in modo strategico, ossia non a tutti e soprattutto non alle stesse condizioni. Di conseguenza, le imprese competono ad armi non pari. Questo tempera la pressione competitiva nel mercato a valle, con conseguenze negative sia per i consumatori che per le imprese: mentre gli intermediari estraggono una quota rilevante del surplus a monte, le condizioni competitive nel mercato finale tendono a peggiorare.

In conclusione, la discriminazione di prezzo nell'era dell'AI assume forme sempre più sofisticate e legate alla disponibilità di dati. Essa può stimolare la competizione e incrementare il surplus dei consumatori. Tuttavia, più frequentemente, favorisce i grandi player dotati di dataset storici e infrastrutture computazionali e, nel caso degli intermediari, sposta il potere a monte, riducendo l'efficienza complessiva del mercato.

3.6.4 Rischi di collusione nei mercati digitali

Un altro aspetto centrale del dibattito sulla concorrenza è la possibilità d'adozione di strategie collusive dagli algoritmi di ML. La collusione è la strategia che porta più imprese a coordinarsi, in modo tacito o esplicito, per fissare prezzi superiori a quelli che si adotterebbero in una situazione perfettamente concorrenziale, così da guadagnare di più a scapito dei consumatori. Se i concorrenti riescono a coordinarsi, preferiscono non violare l'accordo e mantenere prezzi sopra il livello competitivo.

Gli algoritmi di pricing alimentati da modelli di AI possono favorire la collusione tramite due metodi principali. In primis, essi sono in grado di imparare a reagire in modo molto più repentino ai prezzi dei rivali rispetto all'uomo. Questo permette punizioni tempestive per deviazioni e, quindi, aiuta a stabilire strategie collusive. In secondo luogo, l'apprendimento di strategie tramite tentativi ed errori porta a fissare automaticamente prezzi sopra-competitivi,

conducendo a un allineamento involontario fra imprese, che finiscono per adottare strategie simili. La nuova ondata tecnologica porta, inoltre, alla possibilità di collusione senza alcun tipo di comunicazione, fenomeno difficilmente realizzabile quando invece si fa riferimento ad un contesto gestito totalmente da personale umano. Ciò porta problematiche regolatorie e di policy in quanto diventa complicato individuare accordi collusivi.

Tuttavia, la capacità predittiva degli algoritmi può disincentivare la collusione. Infatti, l'abilità di previsione della domanda in momenti specifici permette di individuare i picchi in cui è più promettente abbassare il prezzo per attirare una fetta di mercato maggiore rispetto ai competitor, come ad esempio succede nel settore aereo. In questi termini, le caratteristiche core dei modelli ML rendono più turbolenti e instabili gli accordi collusivi, aumentando la tentazione di deviazione. Tale fenomeno porta ad un potenziale aumento della pressione concorrenziale, portando a prezzi più bassi e a un incremento del surplus del consumatore.

In sintesi, l'impatto finale dell'utilizzo di algoritmi altamente orientati a meccanismi collusivi dipende dal contesto specifico e dalle caratteristiche dei mercati considerati.

3.6.5 Fusioni ed acquisizioni: il caso americano

La struttura di un mercato è fortemente influenzata dalle fusioni e acquisizioni, le quali possono ridurre il numero di concorrenti effettivi e consolidare il potere degli incumbent. Considerando il contesto americano, è importante citare lo studio di McElheran (2024) basato sull'Annual Business Survey. Esso indaga sulle acquisizioni in ambito tech. Un numero sempre più elevato di imprese conclude acquisizioni in tale ambito, con l'obiettivo di accrescere le proprie competenze ed abilità appoggiandosi a realtà già formate, accelerando tempi e minimizzando l'utilizzo di risorse. Tramite tali operazioni è possibile anche espandersi in settori non appartenenti al proprio core business, diversificandosi e risentendo meno della concorrenza. Le aziende target solitamente sono start-up o PMI fortemente innovative che, d'altro lato, mirano a espandersi ed affermarsi su scala più ampia. Le acquisizioni, quindi, possono portare vantaggi sia alle acquirenti che alle acquisite, con l'eccezione delle cosiddette acquisizioni killer, in cui grandi imprese mirano ad inglobare realtà più piccole e più innovative che però si posizionano nello stesso segmento di mercato, potendo di fatto diventare competitor. Tali fenomeni portano all'eliminazione delle target in modo da sopprimere la concorrenza, impedendo loro di sviluppare pienamente il proprio potenziale innovativo.

Guardando all'adozione vera e propria dell'AI, inoltre, si nota prevalenza di utilizzo nelle grandi imprese, a scapito di quelle di minori dimensioni, come anticipato nei paragrafi precedenti. In aggiunta, si evidenzia una concentrazione a livello geografico: i poli tecnologici di maggiore rilievo, come la Silicon Valley, il Research Triangle in North Carolina, e alcune

aree emergenti nel Sud e nell’Ovest degli Stati Uniti (come Nashville, San Antonio, Tampa, Las Vegas) hanno tassi d’adozione molto alti rispetto ad altre zone. Ne deriva il rischio del cosiddetto *AI divide*, che pone preoccupazioni in quanto le regioni meno dinamiche e lontane dai centri metropolitani potrebbero venire escluse da questa ondata tecnologica, ampliando i divari già esistenti.

3.6.6 Venture Capital, imprenditorialità e barriere per i nuovi entranti

Oltre alle barriere tecniche ed economiche, un altro elemento cruciale riguarda il comportamento degli investitori. Il capitale di rischio ricopre un ruolo essenziale nello sviluppo dell’AI, ma i suoi meccanismi intrinseci tendono a favorire la concentrazione di mercato. I *venture capitalist* infatti tendono a supportare le iniziative che hanno maggiori probabilità di affermarsi su scala globale, trascurando e non facendo leva sulle realtà minori che invece potrebbero avere un gran potenziale che non viene espresso. Questa dinamica rafforza gli squilibri competitivi, mettendo ancora più in luce i grandi leader già affermati sul mercato. Il fenomeno verificatosi è altamente negativo in quanto l’evidenza empirica (McElheran et al., 2024) mette in luce il fatto che i giovani imprenditori, spesso a capo di piccole startup, hanno più probabilità di avere successo nell’implementazione di sistemi volti all’utilizzo dell’intelligenza artificiale. Le realtà ampiamente orientate alla crescita, che quindi si discostano dal modello tradizionale di azienda, mostrano maggiore apertura verso l’innovazione, una mentalità sperimentale e una più elevata familiarità con le tecnologie emergenti. L’insieme di tali fattori aumenta le probabilità di successo nella trasformazione digitale, che però purtroppo non vengono sfruttate al massimo per mancanza di capitale, capacità computazionali ed infrastrutturali, senza contare i vincoli normativi complessi presenti. In questo contesto si inseriscono strumenti come le AI sandboxes, ambienti controllati in cui gli innovatori possono testare rapidamente nuove tecnologie sotto lo sguardo attento dei regolatori. Il funzionamento e le implicazioni di tali strumenti verranno approfonditi nel capitolo successivo.

Tabella 12: Effetti dell’AI a livello concorrenziale

Paper	Tipologia di dato	Effetti trovati
Goldfarb, Taska & Teodoris (2020)	Analisi teorico-concettuale	L’AI può diffondersi trasversalmente, riducendo costi di ricerca e di transazione e abbassando le barriere all’entrata,

		favorendo così la concorrenza.
Goldfarb & Tucker (2019)	Analisi di letteratura con framework teorico sui costi digitali	L'AI porta alla riduzione di svariati costi che, di conseguenza, aiutano l'affermarsi di meccanismi pro-concorrenziali.
The Economist (2017)	Articolo giornalistico con caso applicativo	Esempio di applicazione dell'AI nel caso della gestione di sinistri.
Hagiu & Wright (2020)	Analisi teorico-manageriale	Gli effetti di rete e la presenza di <i>switching costs</i> generano rendimenti crescenti, facilitando in tal modo l'affermarsi di <i>superstar firms</i> ed il rafforzamento del potere degli incumbent.
Schafer & Sapi (2020)	Modello teorico con evidenze empiriche	Storici utenti più lunghi permettono di avere un apprendimento migliore e, di conseguenza, un vantaggio cumulativo.
Ezrachi & Stucke (2016)	Analisi teorico-legale	La BBPD prevede che i prezzi siano fissati sul comportamento dei consumatori. Essa porta con sé problematiche legate alla privacy, all'estrazione di surplus e alla diminuzione della concorrenza.

McElheran et al. (2024)	Analisi empirica su Annual Business Survey (USA)	Aumento delle acquisizioni in ambito tech, rischio di acquisizioni killer. Start-up e giovani imprenditori sono più propensi a successo nel mondo dell'AI. Rimangono vincoli tecnici che possono però frenarne l'entrata.
-------------------------	--	---

4 Regolamentazione e policy dell'AI

La rapida diffusione dell'intelligenza artificiale sta sconvolgendo l'economia globale e, con essa, la società. Gli impatti che gli studiosi stanno riscontrando o si aspettano di osservare, ampiamente discussi nel capitolo precedente, pongono preoccupazioni e fanno emergere la necessità di un quadro regolatorio adeguato. La capacità di incidere sui diritti fondamentali, di modificare gli equilibri concorrenziali e di influenzare processi decisionali in ambiti critici rende inevitabile l'intervento pubblico. Con l'obiettivo di equilibrare i valori democratici e la tutela del mercato interno si afferma una duplice esigenza: se da un lato si vuole favorire l'innovazione e la competitività, dall'altro però risulta essenziale tutelare la sicurezza, i diritti fondamentali e la stabilità economico-sociale.

Nel capitolo si analizzerà la situazione attuale, facendo un confronto fra diversi Paesi, quali Europa, Stati Uniti, Giappone e Regno Unito, e si discuteranno le principali sfide che ad oggi rimangono aperte. Nell'ottica della discussione è opportuno differenziare tra regolamentazione e policy. La prima riguarda l'insieme di norme giuridiche che impongono obblighi e divieti a soggetti ed organizzazioni con lo scopo di prevenire o correggere fallimenti di mercato, tutelare interessi pubblici e, più in generale, contenere o arginare rischi di vario tipo. In questi termini risultano di significativa importanza: l'AI Act (2024), il Digital Markets Act (DMA, 2022) ed il General Data Protection Regulation (GDPR, 2016). Quando, invece, si introduce il concetto di policy ci si riferisce a strategie, programmi e strumenti di intervento che si pongono l'obiettivo di promuovere l'innovazione, la competitività e lo sviluppo economico e tecnologico. Qui diventa rilevante l'AI Continent Action Plan (2025), il quale contiene strategie volte a fare dell'UE un AI continent leader.

4.1 I *driver* della regolamentazione

I razionali alla base della regolamentazione sono molteplici e comprendono ambiti differenti che riflettono la portata trasversale dell'AI. In questa sezione verranno approfonditi quelli di maggiore rilevanza, quali: bias, privacy e data governance, proprietà intellettuale e copyright, sicurezza nazionale e finanziaria. Sulla questione inerente alla concorrenza non verrà fornita un'ulteriore trattazione in quanto l'argomento è già ampiamente trattato nei capitoli precedenti. In primo luogo, se si guarda ai diritti fondamentali dell'uomo, le problematiche più significative riguardano i bias e le discriminazioni ad essi correlate. I dati con cui vengono addestrati e testati gli algoritmi di ML spesso riflettono pregiudizi storici e squilibri sociali che, di conseguenza, vengono riportati nell'output fornito durante la fase di deployment effettiva. Sono vari gli esempi di situazioni in cui gli algoritmi hanno incentivato a trattamenti discriminatori; basti

pensare a come l'etnia di un individuo possa diventare un fattore determinante nella concessione di mutui o credito. A causa di tali meccanismi discriminatori innescati dai dati e dagli algoritmi, molte persone di colore subiscono rifiuti da banche nonostante la loro situazione possa non essere così disastrosa. Nella pratica, una persona bianca con lo stesso curriculum redditizio potrebbe ottenere il credito richiesto. Questo non è un caso isolato, il medesimo contesto si crea anche in materia di assunzioni. Il discorso inerente ai bias è ampio in quanto essi possono essere innescati non solo a partire dalla presenza di un dataset di bassa qualità, ma anche da un'etichettatura non adeguata o da un utilizzo improprio degli algoritmi. Inoltre, l'opacità intrinseca dei modelli di ML non permette di guardare a fondo nei meccanismi, minando la scoperta dei punti critici che andrebbero monitorati e, più in generale, la comprensione ed interpretazione delle decisioni di un modello di AI (*explainability*). Nella pratica, quindi, si ha a che fare con una *black box* e, proprio a causa di ciò, la regolamentazione insiste sulla necessità di spiegazioni comprensibili, di documentazione accurata e di meccanismi di controllo indipendente. Questi strumenti risultano utili per rendere gli individui informati, in modo che possano comprendere e contestare decisioni che incidono sui loro diritti. Considerando proprio questi ultimi, un altro aspetto negativo che richiede attenzione è la generazione di contenuti *deepfake*, ovvero manipolazioni di immagini o video con scopo di ricreare verosimilmente persone reali.

Un altro tasto dolente dell'adozione dell'AI riguarda l'utilizzo dei dati e le rispettive privacy e governance. La nuova tecnologia, per definizione, ha bisogno di un grande quantitativo di informazioni per poter addestrare i modelli di ML in modo accurato e puntuale; questa necessità spesso porta all'attuazione di processi di *web scraping* o di consultazione di archivi non progettati per un uso secondario, con l'obiettivo di collezionare dati sensibili e personali. Si sollevano così interrogativi riguardanti al rispetto del GDPR, insieme armonizzato di norme applicabili a tutti i trattamenti di dati personali da parte di organizzazioni situate nello Spazio economico europeo (SEE) o che si rivolgono a persone nell'UE. Oltre al processo di raccolta dati, emergono problematiche relative al loro possibile attacco ed a rischi di *leakage* e *re-identification*. Il primo indica la fuoriuscita indesiderata dei dati, sia a livello di dataset, quando informazioni personali finiscono nei set di addestramento senza adeguata protezione, sia a livello di output, quando un modello restituisce informazioni private se sollecitato dalle giuste richieste. La *re-identification*, invece, si riferisce al fenomeno per cui un individuo può venire riconosciuto nonostante l'anonimato a causa della combinazione di più dati, la quale fornisce indizi per ricostruire l'identità di un soggetto. A ciò si aggiunge il nodo del controllo sui dataset in quanto i dataset più significativi per l'addestramento sono concentrati nelle mani di poche

piattaforme globali, in particolare fornitori di Cloud e Big Tech, quali Google, Microsoft, Meta e Amazon. La governance dei dati quindi si espande: da mera tutela della privacy individuale diventa inclusiva anche di questioni riguardanti l'equità nell'accesso e nel controllo.

Un ulteriore fonte di complessità, richiedente azione regolatoria, è la proprietà intellettuale e, con essa, il diritto d'autore. La questione è duplice, in quanto riguarda sia i contenuti utilizzati per l'addestramento che quelli generati in output dai sistemi intelligenti. Per effettuare l'addestramento dei modelli sono necessarie ingenti quantità di dati, e nell'insieme di informazioni utilizzate ci sono anche immagini, video, testi, musica o altre opere protette da copyright. Ci si pone quindi la domanda se l'utilizzo di tali contenuti sia lecito e, soprattutto, se gli autori possano opporsi rivendicando un diritto di esclusione. In Europa esistono regole che permettono il *text and data mining*, cioè l'estrazione e l'analisi automatica di grandi quantità di dati per facilitare lo sviluppo tecnologico, ma ci si chiede come muti la situazione quando si considerano contenuti soggetti a copyright. Per quanto concerne i risultati in output dagli algoritmi di AI non è ancora chiaro se essi possano essere tutelati da copyright. Inoltre, uno degli aspetti maggiormente dibattuti riguarda le azioni da intraprendere nel caso in cui venga generato un contenuto che imita l'opera di un artista riconoscibile. In quest'ottica si stanno sviluppando strumenti in grado di distinguere fra opere generate da AI o dall'uomo, come ad esempio il *watermarking*, una firma digitale invisibile inserita all'interno di contenuti generati dall'AI, ed i metadati di provenienza, che sono informazioni aggiuntive presenti nel file stesso e che indicano chi l'ha creato, quando e con quali strumenti.

Infine, l'ultimo razionale che viene analizzato riguarda la dimensione della sicurezza nazionale e della stabilità finanziaria. Infatti, l'intelligenza artificiale può essere sfruttata per guadagnare benefici in innovazione, progresso tecnologico e, auspicabilmente, in economia, ma in modo totalmente opposto si può far leva su essa per diffondere disinformazione, fare attacchi informatici o addirittura usarla in sistemi militari autonomi. Nello scenario complesso che viene a formarsi la sicurezza nazionale non è garantita, e proprio per questo sorge la necessità di regolamentazione. Inoltre, la concentrazione delle potenze di calcolo e di cloud nelle mani di pochi leader mondiali rende l'Europa in una posizione vulnerabile e strettamente dipendente da altri Paesi, nello specifico gli Stati Uniti. Se i servizi forniti esternamente si bloccassero o venissero limitati, molti settori ne risentirebbero, quali finanza, energia o trasporti, provocando effetti disastrosi su tutti i fronti.

Come conseguenza a tutti i punti toccati nasce il bisogno di un quadro di governance che garantisca equilibrio tra progresso tecnologico e tutela dei diritti umani e della sicurezza globale. Esso deve contemplare standard di sicurezza, porre norme vincolanti e progettare piani

di resilienza, in modo da ridurre al minimo gli effetti negativi che stanno sorgendo e potrebbero sorgere con la nuova ondata tecnologica.

4.2 Strumenti regolatori e governance

4.2.1 AI Act

Gli strumenti regolatori ideati e adottati dall'Unione Europea per arginare i rischi legati all'intelligenza artificiale si inseriscono in un contesto più ampio di governance digitale e sono volti a garantire, nel loro complesso, un equilibrio tra promozione dell'innovazione e tutela dei diritti fondamentali e della concorrenza. Sono di centrale importanza tre regolamenti che sono stati istituiti precedentemente alla nuova ondata tecnologica, ovvero il Digital Markets Act (DMA, 2022), il Digital Services Act (DSA, 2022) ed il Regolamento generale sulla protezione dei dati (GDPR, 2016). Prima di passare alla rassegna di questi ultimi, si analizza a fondo l'AI Act (2024), primo regolamento orizzontale al mondo sull'AI, ovvero che non si concentra su un unico settore ma copre tutti quelli esistenti.

L'AI Act adotta un approccio basato sul rischio (*risk-based*), in cui i sistemi di intelligenza artificiale sono classificati in base ai pericoli ad essi associati. Ad ogni categoria corrisponde poi una trattazione particolare in termini di obblighi e trasparenza. Nello specifico, il regolamento suddivide in quattro livelli di rischio: inaccettabile, elevato, limitato e minimo o nullo.

Procedendo in ordine dal più al meno grave, si prendono in esame i sistemi AI considerati inaccettabili, ovvero quelli relativi alla violazione dei diritti umani o dei diritti fondamentali dell'Unione Europea, come ad esempio la sorveglianza di massa o il *social scoring*. Prendendo in considerazione quest'ultimo risulta evidente come un sistema basato su tale principio sia dannoso socialmente. Esso si basa sul guadagno o perdita di punti da parte di cittadini in base al loro comportamento. Se inizialmente può sembrare ottimale per incentivare alla buona condotta, riflettendo meglio sui suoi meccanismi, ci si rende conto che può ledere diritti fondamentali: multe non pagate o comportamenti non idonei online possono limitare l'accesso a trasporti pubblici o ad occasioni lavorative. Tali divieti possono portare le persone ad adottare comportamenti ancora peggiori a causa della privazione di tali diritti. Come conseguenza a tutto ciò, i sistemi appartenenti a questa categoria sono considerati vietati.

Proseguendo con il livello di rischio elevato, si passa a sistemi che possono essere adottati a patto che vengano rispettati requisiti specifici, quali: l'obbligo di adozione di misure di sicurezza, la gestione del rischio, la trasparenza sulla modalità di funzionamento del sistema e la supervisione umana. Inoltre, devono essere forniti valutazioni e test di conformità

obbligatoria. In questa categoria rientrano i sistemi che in caso di malfunzionamento potrebbero portare a danni gravi, come ad esempio: il *credit scoring*, i sistemi a guida autonoma, i sistemi di recruiting, le applicazioni in ambito medico ed i sistemi di identificazione biometrica o riconoscimento facciale. Prendendo in considerazione quest'ultimo esempio, Nijeer Parks, un afroamericano residente in New Jersey, è stato condannato ingiustamente a causa di un errore dovuto ad un sistema di AI. L'uomo in realtà era innocente ma siccome gli algoritmi, anche se avanzati, possono confondersi nel riconoscimento facciale man mano che la pelle diventa più scura, è stato confuso con il vero malvivente, anch'egli di colore.

La terza categoria si riferisce ai sistemi con rischio limitato, ovvero a quelli che presentano rischi minori, come chatbot generici o software destinati alla rielaborazione di immagini. Entrambi gli esempi riportati, infatti, non hanno accesso ai dati sensibili o personali degli utenti, limitando così le implicazioni etiche e i rischi di privacy. In questo caso si richiede semplicemente trasparenza e robustezza.

Infine, l'ultimo livello comprende i sistemi a rischio minimo o nullo. Anch'essi non hanno accesso ad informazioni sensibili e, oltretutto, non si appoggiano ad altre infrastrutture o sistemi critici. Questo permette di escludere gran parte dei pericoli che invece possono caratterizzare le categorie sopra citate, motivo per cui questi sistemi non sono soggetti a obblighi particolari o restrizioni. Esempi significativi di questo livello di rischio sono le calcolatrici o i videogiochi basati su AI per smartphone.

Il punto comune a tutte le categorie è che la violazione delle norme previste è oggetto di sanzione, ennesima conferma del fatto che l'AI Act pone leggi vincolanti e non linee guida etiche.

4.2.2 Oltre l'AI Act: le altre norme europee di riferimento

Ad affiancare il regolamento europeo destinato all'AI, ce ne sono altri, che seppur non strettamente legati all'intelligenza artificiale, ne completano il quadro normativo.

In primis si analizzano il Digital Markets Act ed il Digital Services Act, denominati nel loro insieme come Digital Services Package, che rappresentano il core della strategia europea per riequilibrare il potere delle Big Tech e garantire un ecosistema più sicuro ed equo. Il DMA, a orientazione economica, contiene una serie di obblighi e divieti per le piattaforme considerate *gatekeeper*, ovvero occupanti una posizione dominante nell'intermediazione tra imprese e consumatori. Esso mira ad assicurare condizioni di interoperabilità e contestabilità nei mercati digitali e, per garantire ciò vieta: l'auto-preferenza, come ad esempio il favorire propri servizi nei motori di ricerca da parte di Google, la combinazione di dati raccolti da servizi diversi senza consenso esplicito e l'impedimento ai venditori di offrire prezzi migliori su altre piattaforme.

D'altra parte, invece, il DSA è più indirizzato alla sicurezza online e, proprio con tale obiettivo, disciplina la responsabilità degli intermediari attraverso sistemi di segnalazione e rimozione di contenuti illegali, valutazioni annuali del rischio e richiesta di trasparenza. Quest'ultima, in particolare, è strettamente legata all'AI generativa.

A completare il contesto normativo di riferimento si colloca il GDPR, che seppur non essendo progettato per l'AI risulta essenziale in quanto fortemente legato ai dati, risorsa chiave dei nuovi algoritmi. Con particolare attenzione alla privacy, nasce qui il concetto di *accountability*: le organizzazioni non devono limitarsi a rispettare le norme inerenti alla privacy ma devono anche fornire dimostrazioni concrete di ciò. Le problematiche elencate nel paragrafo precedente, quali il *web scraping*, la *re-identification* e la collezione di dati sensibili sono da risolvere anche alla luce del GDPR.

Un ulteriore strumento regolatorio è rappresentato dalla politica antitrust, operante ex-post, che si pone lo scopo di controllare la concentrazione del mercato. In questi termini è di fondamentale importanza porre attenzione alle acquisizioni delle Big Tech, in quanto il rischio che esse diventino acquisizioni killer a sfavore di startup e PMI desta sempre più preoccupazione.

4.2.3 Sandboxes

Accanto alle normative vincolanti, si stanno affermando strumenti flessibili e promotori di innovazione, ovvero le regulatory sandboxes. L'articolo 57 del Capitolo VI dell'AI Act definisce le sandboxes come ambienti controllati che favoriscono l'innovazione e facilitano lo sviluppo, l'addestramento, il testing e la validazione di sistemi di AI innovativi per un tempo limitato prima della loro immissione sul mercato. Le organizzazioni possono quindi sperimentare soluzioni sotto il controllo delle autorità preposte al controllo e nei rispetti dell'AI Act.

I vantaggi principali associati a questi strumenti riguardano la possibilità di sperimentare soluzioni innovative trattando anche dati personali senza il consenso specifico degli utenti grazie a deroghe al GDPR. Si ha, quindi, più libertà rispetto al contesto normativo standard. In questi termini è facilitata anche la comunicazione fra organizzazioni e regolatori. Inoltre, piccole realtà che internamente non dispongono delle risorse necessarie, possono beneficiare del supporto di figure specialistiche in termini di compliance, auditing e validazione dei modelli. Infine, la partecipazione alle sandboxes aiuta ad acquisire fiducia da parte di clienti ed investitori, elemento da non sottovalutare in quanto necessario per la raccolta di capitale da poter reinvestire in nuovi progetti all'avanguardia.

Ai vantaggi appena elencati, però, si affiancano anche delle criticità e dei limiti. In primis emerge il rischio che tali strumenti si trasformino in meccanismi di *compliance learning*, ovvero luoghi dove le imprese imparano ad adeguarsi alla regolamentazione vigente, senza un effettivo importo in termini di innovazione. In secondo luogo, l'efficacia delle sandboxes potrebbe essere ostacolata dagli iter burocratici che potrebbero essere potenzialmente lenti e poco reattivi: le imprese ed i centri di ricerca potrebbero non sfruttare appieno lo strumento in quanto percepito come poco snello e rapido in termini di sperimentazione, in contrapposizione con l'obiettivo di minimizzare il *time to market* di soluzioni innovative.

A livello geografico, le sandboxes sono state introdotte dall'Unione Europea tramite l'AI Act, ma in realtà il concetto non è nuovo: il Regno Unito nel 2016 aveva lanciato le prime sandboxes, che nel tempo hanno fatto sbocciare molte imprese, tra cui Zilch, una rete di pagamento sovvenzionata dalla pubblicità che ha raggiunto lo status di doppio unicorno nel 2020. Dopo tale successo, è stata inaugurata la *Supercharged Sandbox* a giugno 2025, in collaborazione con Nvidia, con l'obiettivo di focalizzarsi sull'AI. Anche altri Paesi hanno supportato l'iniziativa. La Spagna è stata la prima a lanciare una sandbox dopo l'AI Act e, inoltre, offre accesso al *Barcelona Supercomputing Center*, uno dei supercomputer più potenti d'Europa. La Norvegia ha lanciato *Datatilsynet*, anch'essa orientata all'AI, mentre la Francia ha avviato progetti settoriali specifici per la sanità ed i servizi pubblici.

4.3 Contesto internazionale e cooperazione multilaterale

L'AI è una tecnologia globale ed interconnessa: i modelli vengono sviluppati, addestrati e successivamente distribuiti attraverso catene del valore internazionali, mentre i rischi e le opportunità generati hanno effetti su scala mondiale. Di conseguenza, diventa necessario creare una cooperazione multilaterale, in modo da rendere il quadro normativo meno frammentato e più in armonia tra i diversi Stati. Attualmente i vari Paesi hanno ancora approcci molto diversi tra loro. L'Europa, come spiegato nel paragrafo precedente, ha un approccio regolativo molto forte volto a tutelare i diritti fondamentali. Gli altri Paesi che verranno analizzati in questa sezione, quali Stati Uniti, Cina, Regno Unito e Giappone si discostano da tale modello.

Considerando in prima battuta gli Stati Uniti, si rileva subito l'importanza centrale dell'Executive Order on AI (EO) firmato da Biden nel 2023 con l'obiettivo di promuovere lo sviluppo e l'uso sicuro, protetto e affidabile delle tecnologie di AI. Esso enuncia otto principi e finalità da seguire:

1. Garantire la sicurezza dell'AI
2. Promuovere l'innovazione e la concorrenza

3. Supportare i lavoratori attraverso la formazione e comprendere l'impatto dell'AI sul lavoro
4. Promuovere l'eguaglianza e i diritti civili
5. Proteggere consumatori
6. Proteggere la privacy
7. Promuovere l'uso di AI nel governo federale
8. Rafforzare la leadership americana all'estero

La struttura dell'EO rende la regolamentazione statunitense molto più decentralizzata rispetto a quella europea, in quanto si delega ad agenzie e dipartimenti federali il compito di elaborare ed implementare linee guida sui singoli principi e priorità. Si rivolge particolare attenzione ai *dual-use foundation models*, modelli di grandi dimensioni, caratterizzati da miliardi di parametri e capacità molto ampie, per i quali vengono introdotti obblighi di trasparenza e test di sicurezza. Inoltre, si mira ad aumentare la sicurezza per quanto riguarda la prevenzione di cyberattacchi e la gestione delle infrastrutture. Infine, si affronta anche la questione del copyright e dell'impatto economico che l'AI può suscitare, in particolare modo nei confronti dei lavoratori. L'EO è affiancato da accordi volontari stipulati con le Big Tech, che servono come strumento rapido per fissare standard minimi comuni e dimostrare impegno nella gestione del rischio, e da un dibattito bipartisan al Senato volto a costruire un consenso legislativo sul futuro della disciplina dell'IA coinvolgendo in primis le stesse imprese.

Passando al contesto cinese, si fa riferimento ad una regolamentazione algoritmica centralizzata, dove il governo ha un controllo centrale e diretto sugli algoritmi, sui dati e sulle piattaforme, stabilisce standard obbligatori per i fornitori e interviene rapidamente in caso di rischi. Inoltre il focus è sulla sovranità tecnologica: ci si pone l'obiettivo di ridurre al massimo la dipendenza da altri Stati, sviluppando internamente le infrastrutture digitali necessarie, anche allo scopo di garantire la sicurezza nazionale. La Cina, in particolare, pone particolare attenzione al rispetto ed alla tutela dei valori fondamentali del socialismo, fortemente caratterizzanti la sua cultura. I contenuti generati da AI devono essere conformi a tali valori e non devono minare la sicurezza pubblica o il sistema politico nazionale, pena sanzioni pesanti. A questo scopo, il governo adotta controlli stringenti sugli algoritmi, sui dati di addestramento e sugli stessi modelli.

Un approccio ancora diverso è adottato da Regno Unito e Giappone, i quali mostrano un atteggiamento pro-adozione e flessibile attraverso la definizione di principi generali applicabili a vari contesti, come sicurezza, equità e trasparenza. In questi termini il White Paper del marzo 2023 risulta fondamentale per quanto riguarda il contesto britannico. Tramite tale approccio,

quindi, non va ad affermarsi una legislazione rigida e stringente, come invece risulta essere quella europea, ad esempio. Tokyo sostiene l'adozione di regole flessibili che incentivino all'innovazione tecnologica, con particolare attenzione e apertura a standard condivisi e cooperazione internazionale. Il Regno Unito prosegue alla stessa maniera, muovendosi in modo ancora più attivo con l'inaugurazione dell'AI Safety Institute e l'AI Safety Summit, istituzioni volte rispettivamente alla definizione di standard tecnici di sicurezza ed alla promozione di cooperazione globale.

Analizzando infine i paesi emergenti, si nota un forte disallineamento rispetto alle altre potenze globali a causa delle risorse scarse e delle regolamentazioni adeguate che li caratterizzano. Gli altri Paesi si sono adoperati per agevolare realtà più arretrate in modo da favorire il *capacity building* tramite la formazione di personale qualificato e funzionari pubblici, lo sviluppo di infrastrutture digitali, il supporto finanziario e la creazione di quadri normativi di base. Nel contesto in via di definizione e sviluppo che viene a crearsi il Brasile ha deciso di adottare un sistema basato sul rischio, con stampo simile all'AI Act, mentre l'India ha seguito il modello britannico e giapponese improntato all'innovazione ed alla flessibilità.

In sintesi, escludendo i Paesi più arretrati, il quadro regolatorio è frammentato: USA e Cina, accompagnati da Regno Unito e Giappone, sono più orientati al mercato, mentre l'Europa adotta un approccio più rigido improntato alla tutela dei cittadini. Nell'ampio contesto globale nasce la necessità di convergenza. Organizzazioni come G7, l'OCSE e l'IMF stanno stilando linee guida comuni e raccomandazioni nel tentativo di unificare la regolamentazione mondiale e di arginare il rischio di una competizione regolatoria, il quale potrebbe portare ad un aumento dei costi di compliance e soprattutto rischierebbe di accentuare le differenze esistenti.

4.4 AI continent action plan EU

4.4.1 Obiettivo e scopo

Le iniziative messe in atto dall'Unione Europea partono dalla creazione di un piano d'azione per il continente sull'AI (AI continent action plan). L'obiettivo è far diventare l'EU un leader nel campo dell'intelligenza artificiale, promuovendo lo sviluppo e la diffusione di soluzioni AI a beneficio della società e dell'economia nei campi di applicazione come l'assistenza sanitaria, il settore dell'*automotive* e in ambito scientifico. L'intento è valorizzare i talenti e le forti industrie tradizionali europee trasformandole in acceleratori per l'AI.

Il piano prevede una serie di iniziative strategiche che costituiscono un investimento complessivo, insieme ai fondi dell'InvestAI, di 200 miliardi per l'intelligenza artificiale. Uno dei punti cardine è la costruzione di tredici *AI Factories* e di cinque *AI Gigafactories* che

riuniranno la potenza di calcolo di oltre 100.000 processori di AI avanzati, investendo una cifra intorno ai 20 miliardi di euro.

I principali scopi dell'AI continent action plan sono quattro. Il primo è l'aumento dell'accesso a dati di alta qualità, che è essenziale per sviluppare e addestrare modelli AI avanzati. L'idea è creare una *data union strategy*, che promuova un mercato unico dei dati, e dei *data labs* all'interno delle *AI Factories* per raccogliere e organizzare dati di alta qualità provenienti da diverse fonti, fornendo così ai ricercatori e agli sviluppatori gli strumenti necessari per favorire l'innovazione.

Il secondo scopo è la promozione dell'AI nei settori strategici, con l'intento di colmare il divario rispetto agli altri paesi e stimolare l'adozione della tecnologia in determinati campi, poiché attualmente l'AI è utilizzata solo dal 13,5% delle imprese nell'Unione Europea. Gli obiettivi principali della strategia comprendono: incentivare l'uso dell'intelligenza artificiale nei settori industriali chiave, favorire l'integrazione dell'AI in ambiti strategici come il settore pubblico e l'assistenza sanitaria e, per attuare concretamente la strategia, utilizzare le infrastrutture esistenti, quali le *AI Factories* e i *Digital Innovation Hubs* europei. Questi ultimi rappresentano degli sportelli unici di supporto alle PMI e alle organizzazioni del settore pubblico nella loro trasformazione digitale, offrendo accesso a tecnologie avanzate, competenze, servizi di innovazione e reti di contatti, al fine di aumentare la loro competitività e crescita, come parte del programma europeo Digital Europe.

Il terzo punta sul rafforzamento delle competenze e dei talenti nel settore dell'AI. La crescente domanda di competenze specialistiche in questo campo richiede interventi mirati per sviluppare e trattenere talenti qualificati all'interno dell'Europa. La Commissione Europea prevede, pertanto, di formare e educare la prossima generazione di esperti, di incentivare i professionisti europei del settore a rimanere o a rientrare nel continente e di attrarre talenti qualificati da paesi terzi, compresi ricercatori e specialisti, al fine di potenziare il capitale umano necessario per sostenere l'innovazione e la competitività europea.

L'ultimo scopo è più di natura pratica e prevede la semplificazione dell'attuazione del regolamento sull'intelligenza artificiale, al fine di facilitare l'adozione e l'applicazione del Regolamento europeo.

4.4.2 Iniziative sulle risorse computazionali e le infrastrutture

La capacità di calcolo, come già analizzato esaurientemente nel capitolo 2, rappresenta oggi uno dei colli di bottiglia più rilevanti nello sviluppo dell'intelligenza artificiale sia dal punto di vista del cloud che dei chip. In assenza di infrastrutture alternative a quelle degli *hyperscaler* e ai

chip del colosso Nvidia, le startup e le PMI europee si trovano costrette a dipendere da fornitori esteri, con effetti negativi sia sulla competitività sia sulla resilienza del sistema economico.

Per ridurre la dipendenza da questi player globali e preservare la competitività del proprio ecosistema sono state messe in atto più iniziative all'interno dell'AI continent action plan. Tra queste il progetto principale è EuroStack, che si propone di promuovere la sovranità digitale europea creando un'alternativa alle soluzioni tecnologiche provenienti da Stati Uniti e Cina. All'interno di questa iniziativa l'obiettivo è creare uno *stack* tecnologico, che comprenda la gestione di cloud, AI, IoT (*Internet of things*), cybersecurity e altri strumenti essenziali per un'Europa tecnologicamente autosufficiente. Questo risulta possibile promuovendo infrastrutture che siano interoperabili e che possano essere utilizzate da vari settori privati e pubblici, oltre che dai cittadini in modo trasparente e accessibile.

A fianco di questa idea si pone un'altra proposta: il Chips Act europeo che mira a rafforzare la capacità produttiva di semiconduttori sul continente, con l'obiettivo di raggiungere il 20% della produzione mondiale entro il 2030. Il piano prevede investimenti per oltre 43 miliardi di euro, finalizzati sia a sostenere la ricerca e lo sviluppo sia ad attrarre nuovi impianti di produzione in Europa. Ad oggi la Commissione ha approvato sette decisioni in materia di aiuti di Stato relative ad impianti di semiconduttori, che rappresentano un investimento pubblico e privato complessivo di oltre 31,5 miliardi di euro. Attualmente l'UE copre il 10% della produzione mondiale di chip. Il piano prevede la creazione di una vera e propria filiera a livello locale e allo stesso tempo interconnessa tra i vari impianti produttivi che man mano saranno costruiti con un co-finanziamento da parte dei Paesi membri. Inoltre, le PMI possono beneficiare di finanziamenti pubblici tramite partnership con grandi aziende tecnologiche o enti di ricerca, potendo così accedere a tecnologie avanzate riducendo i costi di ingresso in questo settore.

Si tratta di un intervento cruciale, poiché le catene di approvvigionamento di chip rimangono concentrate in pochi paesi extra-UE, in particolare Taiwan, Corea del Sud e Cina, esponendo l'Europa a vulnerabilità geopolitiche e industriali.

Parallelamente, l'UE ha promosso l'iniziativa EuroHPC Joint Undertaking, che coordina la realizzazione di una rete di supercomputer europei ad alte prestazioni. L'obiettivo è potenziare la capacità computazionale dell'Europa, ridurre il gap con altre potenze tecnologiche, come USA e Cina, e supportare la ricerca scientifica avanzata e l'innovazione industriale attraverso la costruzione dell'*AI Factories*. Questa potenza fornita dovrebbe consentire l'elaborazione di grandi volumi di dati per l'AI, simulazioni scientifiche, modelli predittivi e altre applicazioni avanzate, dando la possibilità di sfruttarla per il training degli FM su larga scala. In questa maniera si ha la possibilità di contrastare le principali aziende nel settore del cloud computing,

grazie alla creazione di un'infrastruttura connessa, autonoma e competitiva. La caratteristica dell'EuroHPC è quella di essere una piattaforma di calcolo accessibile a costo contenuto sia per grandi enti, come università e governi, sia per PMI e startup. Questo è particolarmente importante per queste ultime due categorie che potrebbero non avere i mezzi per poter accedere a risorse computazionali così avanzate.

Alcuni esempi di infrastrutture per ora realizzate in quest'ottica sono: LUMI, realizzata in Finlandia e considerata una dei più potenti siti costruiti; MareNostrum5 in Spagna, progettato per la simulazione avanzata e la ricerca scientifica; Leonardo in Italia, focalizzato su simulazioni e innovazione; Meluxina in Lussemburgo; Discoverer e Vega rispettivamente in Bulgaria e Slovenia. Nella *figura 9* vengono mostrate tutte le *AI Factories* europee.

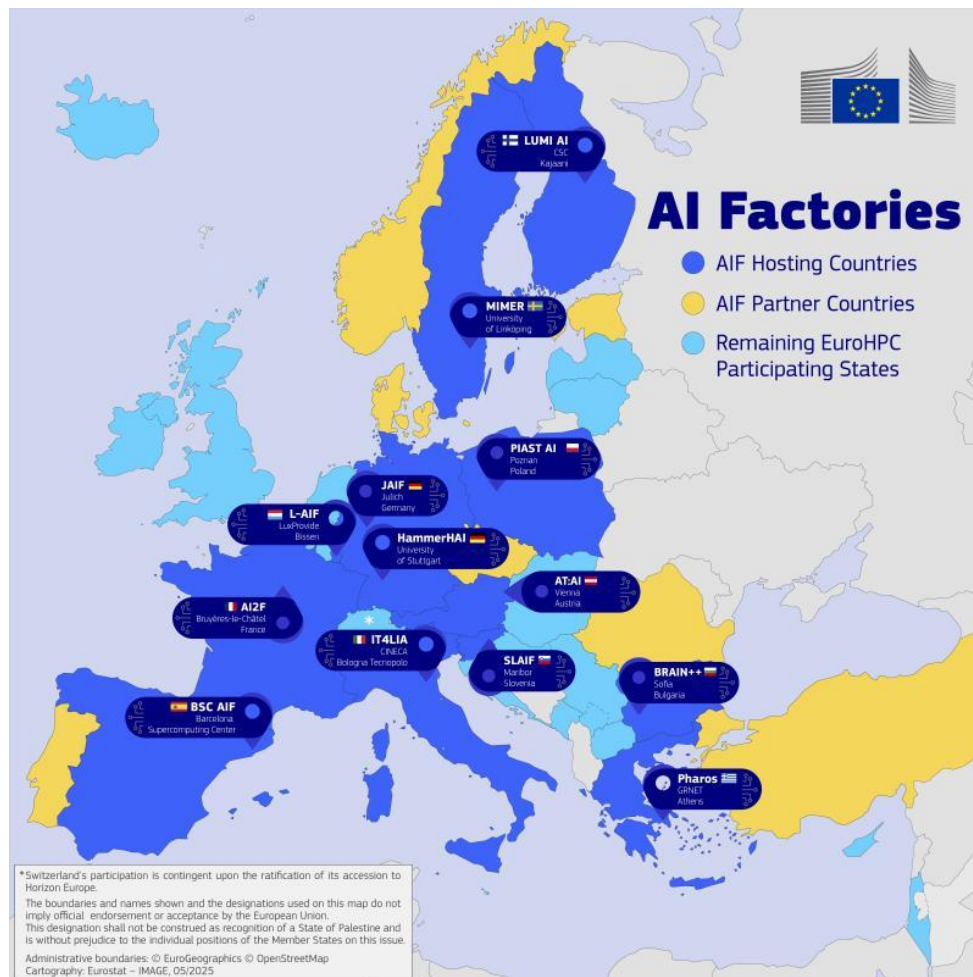


Figura 9: elenco AI Factories dell'UE

Accanto a questi programmi all'interno della Commissione Europea si sta sempre più ragionando in ottica di *capacity pooling*, ovvero la creazione di riserve comuni di calcolo destinate a soggetti pubblici e privati. Tale approccio, ispirato a logiche di mutualizzazione,

consentirebbe di garantire un accesso più equo e diffuso alle risorse, contrastando le dinamiche di esclusione dovute agli alti costi di mercato.

4.4.3 Verso una governance del calcolo

L'emergere del calcolo come risorsa scarsa e strategica ha stimolato un dibattito crescente sulla necessità di una vera e propria compute governance. Alcuni studiosi e policy maker sostengono che, al pari dell'energia o delle telecomunicazioni, il calcolo debba essere considerato una utility strategica, il cui accesso equo è essenziale per la concorrenza e l'innovazione. Nelle sedi comunitarie l'idea è di garantire un accesso che non vada ad introdurre regole di discriminazione nell'utilizzo delle GPU e dei cloud, e di istituire un'autorità europea per la supervisione del calcolo.

La questione non è solo economica ma anche geopolitica: la concentrazione della produzione di chip in Asia orientale e il dominio delle piattaforme cloud statunitensi rendono l'Europa vulnerabile a shock esterni e tensioni internazionali, come dimostrato dalla “*chip war*” tra Stati Uniti e Cina. In questo contesto diventa necessario ed essenziale garantire la resilienza delle catene di approvvigionamento globali, costruendo delle alternative continentali, in modo da avere come obiettivo quello di preservare l'autonomia strategica europea. In questo senso, la governance del calcolo riguarda anche la sicurezza, la sostenibilità e la stabilità dell'intero ecosistema digitale.

4.4.4 Data Labs ed European data union strategy

Per favorire l'innovazione e garantire la sovranità digitale europea, l'Unione Europea sta promuovendo la creazione di *Data Labs* e la strategia di unione dei dati, che hanno l'obiettivo di arricchire l'European Data Spaces. Queste infrastrutture permettono la condivisione dei dati e ne garantiscono l'utilizzo sicuro tra imprese, centri di ricerca e università. Inoltre, stimolano un ecosistema collaborativo in cui i dati possono essere utilizzati per l'addestramento di modelli di intelligenza artificiale, compreso l'addestramento e le implementazioni di applicazioni avanzate dei *foundation models*, senza dover dipendere da imprese esterne come gli *hyperscaler*.

Le iniziative riportate hanno lo scopo di garantire la sovranità dei dati, cioè il pieno controllo da parte dei cittadini e delle istituzioni europee, e l'interoperabilità, attraverso l'adozione di standard comuni che consentono di integrare dataset provenienti da diversi settori e Paesi membri. Questo approccio permette una riduzione del rischio di *lock-in* tecnologico, stimolando allo stesso tempo la collaborazione tra attori pubblici e privati, che determina e favorisce la diffusione dell'innovazione in tutto il continente.

Gli spazi sono accessibili anche a realtà che non dispongono di risorse proprie per l'acquisizione e la gestione di dataset su larga scala, come le PMI e alle startup. Le piattaforme nominate danno la possibilità alle imprese di sperimentare nuove applicazioni AI, di ottimizzare i processi industriali e di sviluppare modelli predittivi, contribuendo alla creazione di un ecosistema digitale resiliente, sicuro e competitivo a livello globale.

La strategia europea dei dati, oltre a portare a un'innovazione diffusa, garantisce un accesso equo ai dati, la protezione dei diritti dei cittadini e uno sviluppo sostenibile di tecnologie all'avanguardia.

4.5 Sfide e prospettive future

L'adozione e la regolamentazione dell'intelligenza artificiale in Europa presentano una serie di sfide strutturali e prospettive di lungo termine che dovranno essere affrontate per garantire uno sviluppo equilibrato e sostenibile del settore. In primo luogo, si evidenzia una contrapposizione tra l'innovazione rapida e la lentezza istituzionale, infatti le tecnologie AI evolvono a ritmi vertiginosi, mentre la regolamentazione e le procedure amministrative spesso faticano a tenere il passo. Questo meccanismo crea un divario tra capacità tecnologiche disponibili e strumenti normativi in grado di governarle efficacemente, con il conseguente rischio di ritardi nell'adozione delle soluzioni più avanzate, soprattutto da parte di PMI e startup, che potrebbero essere penalizzate dalla complessità burocratica.

Un secondo nodo critico riguarda il rischio di *over-regulation*, ossia la possibilità che regole troppo stringenti o complesse creino oneri e vincoli eccessivi, rallentando l'innovazione e riducendo la competitività delle aziende europee. In particolare, le startup e le PMI si potrebbero trovare in difficoltà avendo poche risorse a disposizione rispetto agli oneri necessari per conformarsi alle norme, limitando la loro capacità di partecipare pienamente all'ecosistema digitale continentale. Al contrario, una regolamentazione bilanciata e flessibile può promuovere la fiducia nel mercato, favorendo gli investimenti e la collaborazione.

Un ulteriore punto strategico riguarda la scelta tra modelli *open source* e *closed source* nei *foundation models*. I modelli open source garantiscono una maggiore trasparenza, accessibilità e un'opportunità di collaborazione tra attori pubblici e privati, ma possono comportare rischi di sicurezza e di utilizzo improprio. I modelli *closed source*, invece, offrono un migliore controllo e protezione della proprietà intellettuale, ma rischiano di accentuare la concentrazione tecnologica e di creare barriere all'ingresso per i nuovi attori. L'obiettivo delle istituzioni europee è trovare un equilibrio tra apertura, sicurezza e competitività.

Infine, emerge la questione della capacità di *enforcement* e della trasparenza dei modelli AI. La crescente complessità dei sistemi intelligenti rende difficile verificarne il funzionamento, valutare la conformità alle normative e garantire l'assenza di bias o discriminazioni. Per questo motivo, la governance futura dovrà includere degli strumenti di *auditing* indipendente, degli standard di documentazione dei modelli, delle procedure di validazione continue e dei meccanismi per il monitoraggio della sicurezza e della responsabilità delle soluzioni AI. La trasparenza dei modelli è un elemento fondamentale sia per costruire fiducia tra i cittadini, le imprese e le istituzioni, sia un requisito etico.

5 AI Agents

L'evoluzione dell'intelligenza artificiale ha recentemente portato allo sviluppo di un nuovo fenomeno, ovvero l'AI agentica. Essa va oltre ai semplici sistemi di raccomandazione, chatbot o assistenti virtuali, i quali si limitano a generare risposte sulla base di prompt dell'utente. Come verrà spiegato nel capitolo, gli AI agents hanno caratteristiche peculiari, come la capacità di agire in modo autonomo, di interagire con l'ambiente esterno e digitale e di coordinarsi tra loro per svolgere attività complesse. Queste proprietà li distinguono dagli esistenti sistemi basati sull'intelligenza artificiale, rendendoli un potenziale punto di svolta per l'economia ma, soprattutto, per la ridefinizione dell'interfaccia uomo-macchina e dell'approccio che l'umanità avrà con la tecnologia.

Il settore dell'AI è uno dei più promettenti: si stima, infatti, un tasso di crescita medio annuo del 28% tra il 2024 ed il 2028². Se si prende in considerazione la fetta specifica di mercato riservata agli AI agents, si rilevano previsioni strabilianti, con un tasso annuo medio composto di crescita del 44,8% da ora al 2030 ed un mercato che supererà i 47 miliardi di dollari³.

Il loro potenziale rivoluzionario, però, sta facendo sorgere preoccupazioni dal punto di vista regolatorio in quanto potrebbero emergere varie problematiche, di cui alcune riconducibili ai precedenti progressi tecnologici, mentre altre totalmente nuove. La combinazione tra le due sta sempre più attirando l'attenzione delle autorità di controllo dei mercati, in quanto l'obiettivo è preservare la concorrenza con lo scopo di non rafforzare ulteriormente il potere di mercato delle Big Tech, cruciali nel contesto che si sta affermando considerando che gli AI agents si baseranno su componenti facenti parte dell'AI stack, in particolare modo sui LLM.

Nel capitolo verranno affrontate le questioni sopra citate, partendo con una definizione di AI agent, in modo da poter tracciare la differenza, seppur ancora sfumata, con i precedenti sistemi basati su AI, proseguendo poi con la questione regolatoria e, infine, esponendo il contesto industriale attuale, analizzando sia i primi tentativi falliti che i modelli in fase di lancio.

5.1 Caratteristiche degli AI agents

L'evoluzione dell'intelligenza artificiale ha permesso lo sviluppo di nuovi paradigmi, tra cui gli AI agents. Essi si pongono ad un livello successivo rispetto ai sistemi che già sono largamente diffusi, come ad esempio gli assistenti virtuali quali ChatGPT di OpenAI, in quanto mentre questi ultimi si limitano a rispondere all'utente in base alle richieste in input, i nuovi agenti saranno in grado di agire in autonomia. Tutto ciò è reso possibile tramite alcune

² Anitec-Assinform, Agenti di IA. Conoscere l'IA – Paper #3, luglio 2025

³ <https://www.digital4.biz/marketing/ai-agent-cosa-sono-vantaggi-agenti-autonomi/>

caratteristiche peculiari che contraddistinguono la nuova AI agentica: l'autonomia, la persistenza, la proattività, la capacità di ragionamento e quella multimodale.

Le proprietà appena elencate si intrecciano tra loro, dando vita ciò che è un AI agent.

Esso si contraddistingue in primis per la sua autonomia, in quanto è in grado di pianificare ed eseguire attività complesse senza la necessità di un umano che lo informi con prompt continui.

Una volta che l'obiettivo principale è stato fissato, l'agente è capace di portare a termine l'azione necessaria e, in caso di alta complessità, riesce a scomporla in più compiti separati, in modo da ottimizzare il risultato finale e la velocità di risposta. Questo meccanismo garantisce efficienza dal momento che permette, tra le altre cose, di utilizzare i primi task come base per ragionare su quelli successivi: in questo modo l'agente è in grado di adattarsi alle circostanze e, in caso qualcosa negli step precedenti possa non essere andato come previsto inizialmente, di modificare i suoi piani in modo da raggiungere l'obiettivo finale. L'autonomia, quindi, fa sì che gli AI agents si discostino in modo sorprendente dal funzionamento classico dei sistemi AI, i quali si limitano a dare risposta a comandi puntuali, senza possibilità di adattarsi in corso d'opera e di orientarsi su obiettivi di lungo periodo (*long term-planning*), caratteristica che invece contraddistingue il nuovo paradigma e la sua *agency*.

Strettamente legata all'autonomia è la capacità di ragionamento (*reasoning*). Gli agenti, infatti, non si limitano a seguire regole predefinite ma sono in grado di controllare il proprio flusso operativo e prendere decisioni autonome: si va dal semplice instradamento per scegliere l'alternativa migliore tra un set predefinito, alla pianificazione di obiettivi complessi, come ad esempio la suddivisione in sotto-compiti, come si anticipava sopra. Questo implica sapersi adattare alle circostanze mutevoli e, di conseguenza, conferisce loro grande flessibilità, rendendoli strumenti versatili nel reagire a contesti inattesi.

Un altro pilastro fondamentale è la memoria. I chatbot tradizionali sono caratterizzati da memoria a breve termine e, infatti, ad ogni nuova iterazione parte un nuovo ciclo da zero. Gli agenti superano questa problematica dotandosi di due tipi di memoria, ovvero quella breve e quella a lungo termine. La prima garantisce il mantenimento del contesto di una conversazione o di un compito in corso, in modo da non perdere informazioni tra un passaggio e l'altro; la seconda permette di archiviare conoscenze e interazioni pregresse e di recuperarle quando necessario. Questo rende l'agente capace di personalizzare il suo comportamento sulla base dell'evoluzione delle richieste e delle preferenze dell'utente, diventando così un vero e proprio collaboratore digitale che utilizza le analisi sulla storia pregressa per fondare le scelte future. In questi termini diventa di notevole importanza il concetto di *reinforcement learning*, un paradigma di apprendimento che si basa su un meccanismo di ricompensa ed errore grazie al

quale gli agenti non si limitano ad agire sulla base del passato, ma piuttosto riescono a migliorare nel tempo, sperimentando, valutando e correggendo il proprio comportamento, in modo tale da proporre un'esperienza via via migliore all'utente.

Gli agents, inoltre, si contraddistinguono per la multimodalità. Essi sono in grado di processare, sia in input che in output, dati di diversa tipologia e provenienza, come ad esempio: linguaggio scritto, voce, foto, video e informazioni provenienti da sensori. Quest'ultima categoria risulta nuova rispetto ai tradizionali sistemi AI in quanto permette agli agenti di entrare in comunicazione e captare segnali dal mondo esterno in modo diretto, senza bisogno dell'intermediazione da parte dell'utente che inserisce richieste in input. La grande differenziazione dei dati trattati permette di ampliare enormemente le interazioni possibili, avvicinandosi sempre più al concetto stesso di agente e di interfaccia "umana", riducendo in tal modo la barriera tra utente e macchina.

Infine, l'ultima proprietà che contraddistingue gli AI agents è la proattività: essi monitorano costantemente l'ambiente esterno e, di conseguenza, captano segnali, in modo tale da anticipare bisogni futuri. Questo è possibile tramite un ciclo, che viene ripetuto ad ogni azione svolta da parte dell'agente. In primis si crea una percezione a partire dai segnali rilevati dal mondo esterno. Successivamente si valutano le condizioni e le regole predefinite (*reasoning*), si prende la decisione e, infine, si svolge l'azione necessaria per raggiungere l'obiettivo prefissato. Il meccanismo illustrato differisce sostanzialmente da quello dei sistemi precedenti, in quanto essi erano reattivi ma non proattivi: reagivano solamente se sollecitati da prompt ma non svolgevano azioni in autonomia.

In sintesi, le caratteristiche analizzate rendono gli AI agents dei collaboratori digitali attivi ed intelligenti, capaci di aiutare le imprese ed i consumatori nel soddisfacimento dei loro bisogni, addirittura anticipandoli. Si sta assistendo ad un cambio di paradigma: dall'interazione reattiva e frammentata si sta progredendo verso un'interazione continua, adattiva e orientata agli obiettivi di lungo periodo. Inoltre, le proprietà non sono isolate ma interagiscono tra loro, rafforzandosi a vicenda: ad esempio, la proattività diventa realmente possibile grazie alla persistenza della memoria, che permette l'interpretazione di segnali e l'allineamento con le esperienze pregresse. Il complesso di tali qualità rende gli agenti un punto di svolta potenziale per ridefinire il rapporto uomo-macchina ed aprire nuove prospettive in termini di collaborazione intelligente ed aumento della produttività ed efficienza.

5.2 Architettura e struttura organizzativa dell'agency

Il passaggio da AI generativa ad AI agentic rappresenta una possibile discontinuità nello sviluppo dell'intelligenza artificiale. I modelli linguistici di grandi dimensioni, cuore pulsante dei sistemi di AI, hanno dimostrato di avere importanza centrale nella generazione di contenuti di vario tipo ma, nonostante ciò, rimangono strumenti principalmente reattivi, bisognosi di prompt e indicazioni da parte degli utenti per fornire risposte in output. Gli AI agents, come già discusso nel paragrafo 5.1, sono caratterizzati da un funzionamento differente e, di conseguenza, si appoggiano a vari strumenti nuovi che, però, vanno ad integrare i classici modelli LLM, in modo tale da poter garantire le caratteristiche sopra esplicitate, quali l'autonomia o la proattività, oltre che la focalizzazione su obiettivi a lungo termine (*goal-seeking*).

I LLM, quindi, rimangono il nucleo cognitivo dell'agente. Come già ampiamente discusso nel capitolo 2, si distinguono per le loro avanzate capacità di: comprensione del linguaggio, generazione contenuti, ragionamento flessibile al contesto e interazione naturale con utenti e sistemi. Da soli, però, non bastano ad elevare i tradizionali sistemi di AI, per cui sono affiancati da due elementi innovativi, ovvero la memoria a lungo termine ed i *tools*. Questi ultimi sono fondamentali per la differenziazione degli agenti rispetto ai precedenti chatbot o assistenti. Essi sono un insieme di strumenti, quali API, connettori e interfacce, che l'agente attiva per interagire con piattaforme, sistemi digitali o fisici ed applicazioni esterne, ampliando le proprie capacità operative ed integrandosi ad ambienti tecnologici preesistenti (si veda la *figura 10*).

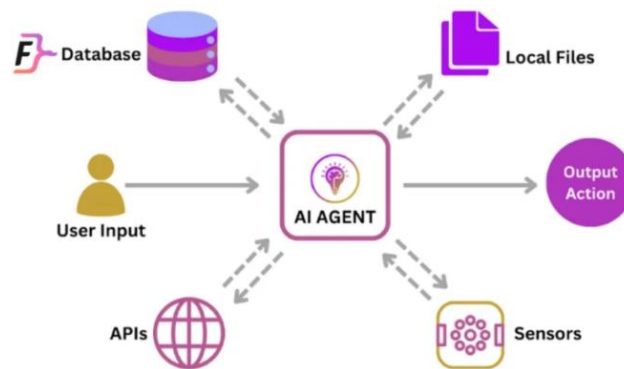


Figura 10: Architettura dell'AI agent con focus sui tools

Inoltre, passi in avanti sono stati compiuti anche grazie alla presenza dei *retriever-enhanced agents*, che sfruttano la *Retrieval-Augmented Generation* (RAG). Quest'ultima permette di ottenere un'affidabilità maggiore in quanto gli agenti non si limitano a consultare le informazioni apprese in fase di addestramento ma interrogano fonti esterne, riuscendo così a estrarre dati più aggiornati e, di conseguenza, prendere decisioni più informate. Questo è un

punto di svolta in quanto permette di essere più accurati nel processo decisionale in contesti mutevoli nel tempo.

Un altro elemento fondamentale nell'architettura propria degli agenti intelligenti è quello dell'orchestrazione. Nel momento in cui l'obiettivo di lungo termine presuppone lo svolgimento di azioni complesse, queste vengono frammentate in task più piccoli e di più facile esecuzione. In conseguenza a ciò diventa essenziale sviluppare una corretta organizzazione e coordinazione, in modo da portare a termine le attività programmate nel migliore modo possibile e rispettando gli standard di efficienza. In questi termini risulta utile combinare più agenti differenti, così da smistare i compiti tra più entità. Per coordinare i vari agenti viene istituito un agent manager, il quale fa da orchestratore e si occupa della pianificazione strategica, dell'assegnazione delle priorità, della suddivisione dei compiti e del monitoraggio dell'avanzamento. In questo modo si garantisce l'adattamento a situazioni complesse e di difficile gestione se mantenute integrate nelle responsabilità di un agente unico. Talvolta, se anche i task sono articolati, vengono costruite *crew*, in cui più agenti lavorano assieme sfruttando le sinergie con l'obiettivo di massimizzare l'efficienza e la produttività. I vari agenti, quindi, si trovano ad interagire in un ecosistema multi-agent, scambiandosi informazioni e dati. Molti agenti, inoltre, adottano un modello *plugin-based*, il quale permette di appoggiarsi a componenti aggiuntivi, sviluppati internamente o da terze parti, per aumentare le capacità ed estendere il nucleo stesso dell'agente.

Per rendere possibile la cooperazione, adottare protocolli standardizzati diventa fondamentale e, in questi termini, nasce il Model Context Protocol (MCP), che regola lo scambio di informazioni tra agenti ed ambienti esterni, in modo tale da garantire un'interazione uniforme con diverse fonti di dati, sistemi e strumenti. Esso definisce in modo chiaro e preciso il contesto operativo, i compiti e le risorse disponibili. L'intelligenza artificiale, così, diventa un'infrastruttura autonoma e distribuita, in cui i vari nodi, costituiti dagli AI agents, rimangono interconnessi. Tale protocollo, in particolare, risulta fondamentale per le *open agentic platforms*, piattaforme che incentivano all'interoperabilità tra agenti multipli e che, quindi, richiedono lo scambio di dati e la cooperazione in ecosistemi differenti.

L'insieme delle componenti appena descritte dà vita al framework agentico, architettura complessa e modulare che descrive i componenti tecnologici ed organizzativi che caratterizzano gli agenti intelligenti. Questi ultimi, grazie alle varie integrazioni che permettono di distinguerli da sistemi tradizionali, diventano entità adattive, in grado di reagire in maniera rapida in scenari dinamici e complessi.

5.3 I vari tipi di AI agents

5.3.1. Categorie di alto livello

Gli AI agents non sono una categoria omogenea, ma si articolano in diverse tipologie, che rispecchiano la loro versatilità d'uso in diversi contesti. Procedendo con ordine, si analizzeranno prima le suddivisioni più generali, per poi passare a fattori caratterizzanti di tipo specifico. Le categorie analizzate sono sintetizzate nel grafico ad albero riportato in *figura 11*.

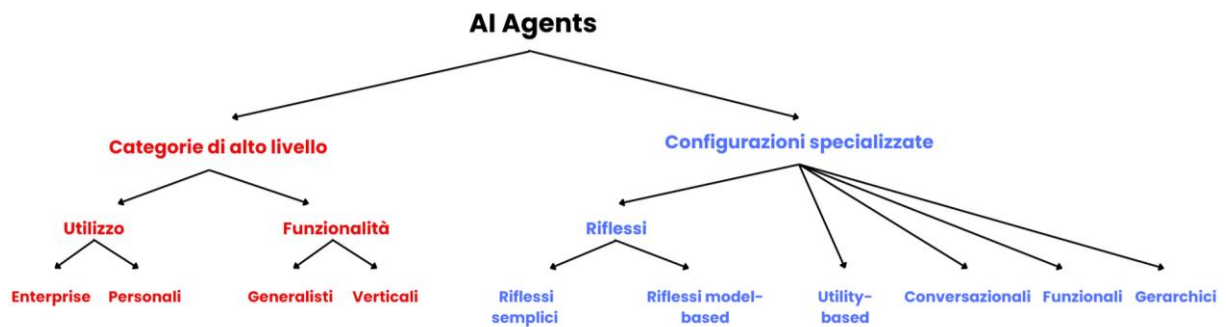


Figura 11: Tipologie di AI agents

Innanzitutto bisogna distinguere l'utilizzo da parte di imprese o di utenti singoli. A questo scopo nascono gli agenti enterprise e quelli personali. I primi sono pensati per i contesti aziendali ed i loro obiettivi principali sono quelli di supportare i processi organizzativi e decisionali, ottimizzare i flussi di lavoro e automatizzare la gestione dei dati, a partire dalla pulizia ed etichettatura, fino alle questioni relative alla privacy. In definitiva sono lo strumento chiave per permettere la trasformazione digitale del mondo enterprise. I limiti associati sono inerenti alla forte integrazione richiesta con l'organizzazione stessa, che quindi può essere soggetta a necessità di cambiamento e di riconfigurazione. Gli agenti personali, invece, sono progettati per affiancare l'utente singolo nella gestione delle attività quotidiane, come ad esempio la schedulazione di attività sul calendario. In questo senso risulta fondamentale riuscire a sviluppare agenti con un alto grado di personalizzazione, in modo da poter anticipare e prevedere i bisogni dell'uomo.

La seconda distinzione riguarda lo spettro di funzioni che l'agente è in grado di performare e, a tal proposito, nascono gli agenti generalisti e quelli verticali. I primi sono "collaboratori universali" in quanto abili nello svolgimento di task eterogenei tra loro, con assenza di un dominio specifico di riferimento. Questa loro caratteristica li rende estremamente flessibili, sebbene non altamente specializzati, e capaci di interagire con un numero alto di strumenti, sistemi e piattaforme esterne. I secondi, al contrario, sono molto più specialistici e si concentrano su settori distinti, come ad esempio la finanza, la sanità o la logistica. Il loro modus

operandi li rende molto più accurati e precisi, con possibilità di accesso a dati strutturati e regole standard, sacrificando però la versatilità tipica degli agenti generalisti.

5.3.2 Configurazioni specializzate

Accanto alle due categorizzazioni di alto livello si delinea uno scenario più articolato, formato da molteplici tipologie di agenti identificate da analisti e tech vendor.

Innanzitutto si considerano i riflessi, distinguendo così tra agenti a riflessi semplici e agenti a riflessi *model-based*. I primi operano seguendo standard specificati al momento dello sviluppo e reagiscono agli input senza prestare attenzione al contesto più ampio in cui ci si inserisce. Proprio per queste caratteristiche, si collocano vicino ad assistenti e chatbot assistenti: non sfruttano appieno il potenziale intrinseco degli agenti intelligenti. Al contrario, quelli con riflessi *model-based* mappano l'ambiente rilevante per l'obiettivo da raggiungere grazie ad un modello integrato e, di conseguenza, prima di decidere e, successivamente, di agire, valutano più accuratamente gli effetti attesi delle varie alternative disponibili.

Procedendo per grado di sviluppo, emergono gli agenti basati su obiettivi. Essi seguono lo schema degli agenti *model-based* ma se ne discostano per il fatto che nascono per specializzarsi nel conseguimento di obiettivi di lungo periodo.

Gli agenti *utility-based* o *task-based*, invece, sono molto più specializzati e vengono richiamati da altri agenti per svolgere compiti precisi, come ad esempio inviare mail, recuperare documenti o eseguire calcoli. Fanno parte di questa categoria gli agenti RAG, essenziali per il reperimento di dati aggiornati da dare in input agli LLM, quelli specializzati in coding e quelli addetti alla ricerca.

In ambito ERP o IoT si stanno sviluppando gli agenti conversazionali, i quali interagiscono col mondo esterno ma, per ora, rimangono limitati all'interazione umana, senza escludere, però, che in futuro potranno avanzare verso un'integrazione con altri software o apparecchiature digitali. Questo permetterebbe un salto di qualità, soprattutto per quanto riguarda il contesto dell'Internet of Things.

Considerando le organizzazioni in senso stretto diventano rilevanti gli agenti funzionali, che sono pensati per essere associati ad una funzione specifica (o ad un ruolo organizzativo). Essi svolgono attività inerenti ad un campo ristretto: ad esempio, si possono concentrare sul recruiting, sul credito, sull'assistenza clienti o sul marketing. Differenziandosi per l'ambito di applicazione, utilizzano strumenti specifici e, inoltre, comunicano tra loro per portare a termine compiti che coinvolgono più reparti aziendali.

Infine, gli agenti gerarchici si contraddistinguono per la loro struttura multilivello: quelli negli strati superiori suddividono le attività in task elementari, che vengono assegnati

successivamente agli agenti di livello più basso. Per l'architettura stessa di questa categoria è richiesto un alto grado di coordinamento che è svolto dagli agenti che smistano i vari compiti, i quali agiscono in maniera analoga ad un agent manager.

In sintesi, la pluralità di agenti esistenti o in via di sviluppo dimostra come l'agency sia un concetto flessibile, che propone soluzioni adatte a esigenze diverse. La varietà di tipologie che si stanno affermando è indice della natura emergente e sperimentale dell'AI agentica, ma al contempo suggerisce che ci sono alte probabilità che essa sarà il punto di svolta per il futuro dell'intelligenza artificiale.

5.4 Settori d'impatto ed applicazioni

La natura collaborativa degli AI agents li rende adatti per una vasta gamma di settori ed applicazioni. Considerando l'utilizzo a livello di singolo consumatore, essi possono essere d'aiuto nella gestione della quotidianità attraverso l'automatizzazione di azioni routinarie, come l'invio di e-mail e documenti oppure la programmazione o la cancellazione di eventi a calendario. Inoltre, si rivelano utili nel settore del turismo, dello shopping e dell'assistenza sanitaria. Per quanto riguarda il primo, gli agenti saranno in grado di gestire la prenotazione di una vacanza nel suo complesso, a partire dai trasporti e dall'alloggio fino alla pianificazione delle attività da svolgere durante il soggiorno. Nell'ambito dello shopping saranno utili per confrontare alternative in termini di prezzo e qualità, facendo così risparmiare tempo all'utente, che non dovrà più scorrere pagine web per confrontare diversi prodotti. Per quanto riguarda l'*healthcare* si tratterà di fornire assistenza ai pazienti ed agli operatori sanitari per il monitoraggio di parametri vitali o strettamente correlati alla problematica specifica riscontrata. Inoltre si avrà la possibilità di rendere automatizzato l'invio di notifiche specifiche per facilitare l'aderenza alla terapia, nonché di messaggi di promemoria per quanto riguarda visite, la cui prenotazione potrà anch'essa essere automatizzata.

Sul versante B2B, gli impatti principali riguarderanno il customer care, la logistica ed il supporto alle vendite. In termini di assistenza clienti ci saranno progressi significativi, come testimonia Klarna, fintech unicorno, che ha sviluppato con OpenAI un agente in grado di gestire due terzi delle richieste di assistenza, affrontando problematiche relative a: rimborsi, resi, pagamenti, cancellazioni, contestazioni ed errori di fatturazione. In ambito logistico si avrà supporto per il coordinamento dei flussi della merce, per il monitoraggio real-time dell'andamento della supply chain e per l'ottimizzazione delle rotte. Considerando le vendite, si avranno vantaggi legati alla previsione della domanda e, di conseguenza all'anticipazione di bisogni, la quale permetterà di esercitare un alto grado di personalizzazione.

In linea di massima, in ogni settore che si prenda in esame, c'è una forte tendenza all'automazione completa di alcuni processi. Se l'andamento previsto dovesse affermarsi, si passerebbe dall'avere agenti *human-in-the-loop*, in cui l'uomo ricopre ancora un ruolo fondamentale, ad esempio fornendo dati mancanti o supervisionando l'attività dell'AI agent, ad una *agent-to-agent economy*, in cui si ci si affida ad altri agenti intelligenti invece che richiedere il supporto a figure umane. Questo permetterebbe di massimizzare la crescita della produttività, superando addirittura i guadagni già previsti grazie all'adozione dei tradizionali sistemi di intelligenza artificiale, ampiamente discussi nel capitolo riguardante gli effetti. Inoltre, la possibilità di affidare numerose mansioni agli agenti permetterebbe di eliminare attuali colli di bottiglia, ad esempio in processi manifatturieri, grazie alla scalabilità propria dei sistemi intelligenti. Questo, tra le altre cose, porterebbe all'uomo benefici in termini di tempo risparmiato e, soprattutto, consentirebbe di dedicare più attenzione ad attività a maggiore valore aggiunto.

5.5 Posizionamento degli AI agents nello stack

Come illustrato nel capitolo 2, l'AI stack si articola su più strati, ovvero: hardware, servizi cloud, dati, *foundation models* ed applicazioni. I recenti AI agents si collocano in un nuovo livello intermedio, introducendo una componente del tutto nuova rispetto a quelle preesistenti. Gli agenti intelligenti si basano sui *foundation models*, che costituiscono il loro nucleo cognitivo, ma accanto a essi costruiscono un'architettura dotata di numerose funzionalità aggiuntive. D'altra parte, non coincidono nemmeno con le applicazioni finali, punto di contatto cruciale con il consumatore, in quanto non si limitano ad offrire servizi di raccomandazione, assistenza o chatbot, ma si spingono oltre, collegando i modelli sottostanti con l'ambiente esterno, inteso come l'insieme di applicazioni, piattaforme o servizi messi a disposizione da terze parti. Nella pratica, il nuovo paradigma si instaura in uno strato trasversale, in quanto mette in comunicazione varie componenti del classico AI stack.

La posizione intermedia degli agenti permette loro di poter esercitare un duplice ruolo. In primis, essi possono facilitare l'accesso ai *foundation models* per le imprese e gli utenti in quanto questi ultimi non dovrebbero più interfacciarsi con le barriere tecniche tipiche degli algoritmi sottostanti; il punto di contatto diventerebbe l'agente. In tal modo si abbasserebbero le competenze richieste per governare i modelli, con una conseguente amplificazione della platea di utilizzatori.

D'altra parte, però, c'è il rischio che essi possano divenire un punto di controllo strategico, in quanto chi sviluppa e controlla l'agente diventa il responsabile dell'interazione uomo-macchina

e, quindi, dell’orientazione di consumi e scelte da parte degli utenti o di processi decisionali da parte delle imprese. In questi termini, gli AI agents iniziano a essere percepiti come i potenziali successori dei sistemi operativi o dei browser, come si approfondirà più avanti con riferimento al DMA, anche grazie alla loro adozione sempre più diffusa. Il recente studio di PwC, *AI Agent Survey* (2025), mette in luce la situazione attuale (*figura 12*), evidenziando che l’80% delle imprese statunitensi analizzate utilizza tecnologie legate agli AI agents, nonostante il 60% si collochi ad un’adozione parziale (*broad o limited adoption*). Ciò riflette la presenza di ostacoli tecnologici e la necessità di un’integrazione graduale. Inoltre, un numero più limitato di aziende (circa il 15%) è in fase di esplorazione, seppur non abbia ancora adottato in via definitiva i collaboratori nei propri processi. Questo suggerisce che la direzione intrapresa è ormai tracciata: l’AI agentica è un’evoluzione in atto ed è destinata a trasformare l’equilibrio competitivo e l’interazione tra uomo e tecnologia. Questo è confermato anche dall’aumento dei budget indirizzati a tali progetti, come viene riportato in *figura 13*.

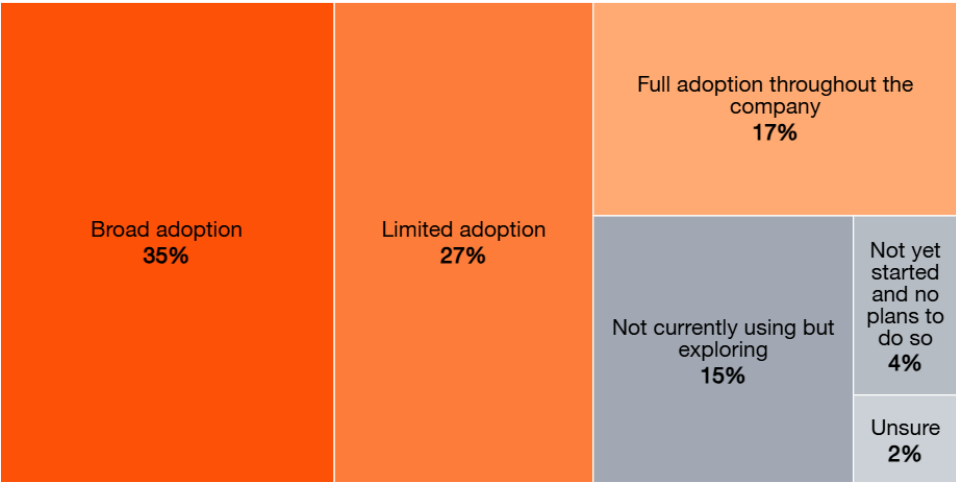


Figura 12: Utilizzo di AI agents negli Stati Uniti (fonte PwC’s AI Agent Survey, 2025)

AI budget increases due to agentic AI

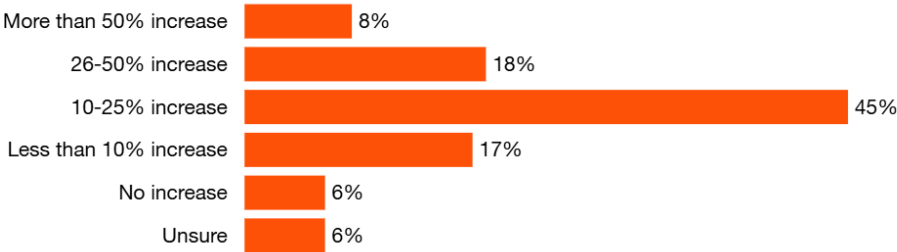


Figura 13: Aumento del budget destinato all’AI agentica (fonte:PwC’s AI Agent Survey, 2025)

5.6 Impatto sulle piattaforme e sui modelli di business

Il posizionamento degli AI agents in un livello intermedio della catena del valore dell’AI fa vacillare l’equilibrio competitivo nei mercati digitali, con rischi elevati per le piattaforme

esistenti, le quali potrebbero subire un impatto non indifferente per quanto riguarda i propri modelli di business. Gli agenti, interfacciandosi in modo diretto con l'utente, potrebbero essere i successori di motori di ricerca e web browser. I precedenti chatbot, assistenti vocali o digitali facilitavano l'accesso a servizi esistenti, mentre gli agenti AI diventeranno lo strumento chiave per la ricerca informazioni da parte dell'utente. Quest'ultimo, di conseguenza, non dovrà più comparare alternative o concludere transazioni in quanto sarà il collaboratore digitale ad occuparsi di tali attività.

Tale trasformazione mette in discussione la stabilità delle piattaforme esistenti in quanto c'è rischio di una potenziale disintermediazione da parte degli AI agents. Questi ultimi agiscono mettendosi in comunicazione direttamente con i fornitori di servizi e prodotti, eliminando di conseguenza la necessità di intermediari. Le piattaforme, nate e sviluppate proprio con tale scopo, potrebbero quindi incontrare difficoltà nella sopravvivenza in un mondo digitale dominato da agenti. Ad esempio, considerando il settore turistico, Booking e piattaforme simili potrebbero risultare superflue in quanto il loro obiettivo di far incontrare domanda e offerta diventerebbe irrilevante nel momento in cui un collaboratore digitale sarebbe capace di prenotare da capo a fondo una vacanza interfacciandosi direttamente con i fornitori di alloggi, hotel, trasporti o servizi di intrattenimento. Lo stesso potrebbe succedere nel campo del food delivery, in cui l'agente, su richiesta dell'utente, potrebbe ordinare direttamente da ristoranti, *by-passando* app come Deliveroo o Glovo.

D'altro canto, però, non è detto che le piattaforme debbano scomparire dal mercato. Esse potrebbero essere soggette ad un processo trasformativo che le porti, ad esempio, ad essere un luogo di interazione tra agenti. Come già evidenziato, essi avranno bisogno di appoggiarsi gli uni agli altri in caso di compiti complessi e, in questi termini, le piattaforme potrebbero reinventarsi per facilitare la loro connessione. Inoltre esse potrebbero diventare un marketplace di API e strumenti esterni, in modo tale da essere consultate dagli agenti per ampliare le loro capacità intrinseche.

Considerando i modelli di business, le imprese potrebbero dover rivedere le loro strategie di marketing e distribuzione in quanto in un mondo governato dagli agenti e non più dall'uomo, cambiano anche le entità con cui interfacciarsi. Si passa, quindi, da modelli business-to-consumer (B2C) a modelli business-to-agent (B2A), in cui l'obiettivo delle aziende diventa essere attrattive per gli agenti. Essi hanno criteri diversi di preferenza rispetto all'uomo: mentre quest'ultimo è più influenzabile dalle emozioni, come dimostra l'attaccamento al brand, i primi si basano su informazioni più oggettive, quali il prezzo, la qualità, l'affidabilità e la coerenza con le preferenze dell'utente (a meno di integrazioni verticali, come verrà approfondito in

seguito). Come conseguenza a tali criteri, le imprese diventeranno più propense a fornire dati strutturati ed aggiornati, accessibili tramite API, a rispettare requisiti di trasparenza e ad adottare protocolli che rendano più fluida l'interazione con gli agenti. Per le organizzazioni entrare nelle preferenze dei collaboratori digitali sarà cruciale in quanto permetterà loro di ottenere quote di mercato che, nel mondo allo stato attuale, avrebbero richiesto ingenti investimenti in marketing e pubblicità.

5.7 Impatto sulla concorrenza

5.7.1 Premesse generali

L'avvento degli AI agents ed il loro posizionamento in un livello intermedio dell'AI stack desta preoccupazioni in quanto potrebbe ridefinire profondamente gli equilibri concorrenziali. Il loro impatto appare duplice e contraddittorio. Se da un lato essi agiscono in modo pro-concorrenziale, abbassando le barriere all'ingresso per favorire l'entrata di nuovi player, come ad esempio sviluppatori di modelli o di soluzioni open source, dall'altro rappresentano un nodo di controllo in quanto diventano l'interfaccia principale tra uomo e macchina, sostituendosi al primo per la scelta dei servizi o prodotti. In questi termini, diventano i potenziali successori dei motori di ricerca o dei browser web e, di conseguenza, emerge il rischio che essi si affermino come nuovi gatekeeper.

La natura emergente dei collaboratori digitali rende il quadro competitivo incerto e attualmente le previsioni appaiono contraddittorie. Un primo filone sostiene che essi rafforzeranno il potere dei leader. La capacità di apprendimento legata alla qualità e quantità dei dati potrebbe portare a vantaggi competitivi per gli operatori attualmente dominanti nei LLM, nelle infrastrutture cloud e negli ecosistemi operativi. In tal modo le Big Tech potrebbero sfruttare i *feedback loop* sui dati e le economie di scala per imporre i propri agenti come interfaccia dominante: l'agente così diventerebbe capace di indirizzare consumi, preferenze e persino scelte strategiche delle imprese. Questo implicherebbe il rafforzamento di meccanismi anti-concorrenziali e farebbe emergere preoccupazioni ulteriori in termini di regolamentazione. In questi termini è emblematico l'esempio di Meta, che controlla Whatsapp, Instagram, Facebook e Messenger. Dal 2023 ha iniziato a distribuire Meta AI, assistente conversazionale improntato a trasformarsi in agente, in maniera integrata ai servizi già offerti. La posizione di Meta all'interno del mercato digitale e la quantità di dati a cui ha accesso tramite la moltitudine di piattaforme controllate rende il caso delicato e lo pone sotto l'attenzione delle autorità di controllo.

Il secondo scenario ipotetico, invece, pone le sue fondamenta sull'interoperabilità e la modularità dell'AI agentica. Tali proprietà potrebbero favorire la contestabilità dei mercati in

quanto stimolerebbero la collaborazione tra agenti e, in parallelo, incentiverebbero l'adozione di protocolli adeguati, come ad esempio l'attuale MCP, e lo sviluppo di piattaforme aperte per l'interazione, come già anticipato nel paragrafo precedente.

Considerando i nuovi entranti e le startup lo scenario atteso è ambivalente: se da un lato la presenza di un nuovo *layer* apre l'opportunità di offrire soluzioni innovative e differenziate rispetto a quelle offerte dalle Big Tech, dall'altro lato il rischio di barriere all'ingresso rimane elevato, soprattutto se i leader dovessero integrare i propri agenti con sistemi operativi o ecosistemi già consolidati.

Alla luce di queste dinamiche, risulta evidente che gli AI agents saranno i promotori di un cambiamento a livello concorrenziale. Nonostante le loro caratteristiche intrinseche possano portare ad un atteggiamento positivo, molti sono gli indizi che invece suggeriscono che potrebbero innescarsi dinamiche volte a rendere il mercato meno contestabile. Nello specifico, i rischi individuati sono due: in primo luogo gli agenti potrebbero venire esclusi a causa dell'iniquità di risorse chiave, in aggiunta gli agenti stessi potrebbero agire per escludere altri operatori grazie alla loro posizione privilegiata della catena del valore, che li rende un nodo strategico nel controllo di essa. I paragrafi successivi si concentreranno sull'analisi dei due scenari appena citati.

5.7.2 Esclusione degli AI agents

5.7.2.1 Esclusione a monte

La prima modalità attraverso cui gli agenti potrebbero essere soggetto di esclusione è quella legata agli input necessari al loro funzionamento. Come spiegato nel paragrafo relativo alla loro architettura, si appoggiano su *foundation models*. Essi sono liberi di svilupparli internamente o di rivolgersi a terzi. La prima opzione è molto costosa e, quindi, raramente utilizzata. In genere si fa riferimento a modelli di altri, che possono essere open source o chiusi, quindi accessibili tramite API. Il problema sorge nel momento in cui i fornitori potrebbero decidere di interrompere o limitare la disponibilità di tali asset, ad esempio adottando strategie del tipo “*open early, closed late*”. Un'altra criticità è legata ai chip AI integrati nei dispositivi: essi a volte supportano solamente alcuni tipi modelli, limitando enormemente la scelta di questi ultimi.

Il secondo vincolo riguarda i dati che alimentano i modelli. I dataset proprietari, ricchi e costantemente aggiornati, sono risorse limitate, a differenza della potenza di calcolo (chip e servizi cloud) che, pur rimanendo concentrata, risulta comunque accessibile. Seppure esistano siti web aperti che forniscono informazioni, il loro utilizzo non consente di avere un vantaggio

competitivo in quanto di libero accesso. Alcuni leader, come Google, stanno iniziando a sviluppare agenti intelligenti appoggiandosi sui dati di loro proprietà, nel caso citato provenienti da YouTube, ad esempio. Il controllo di informazioni in grandi quantità e di una buona qualità permette loro di sviluppare agenti più avanzati, spingendoli quindi a tenere per sé i dataset per non condividere il beneficio guadagnato in termini di competitività. In alternativa, ci si potrebbe appoggiare a licenze, che però presentano criticità dovute al fatto che si tende a dare l'esclusività ad un solo sviluppatore di AI agents, permettendo a quest'ultimo di imporsi sul mercato a scapito di altri.

In aggiunta, spesso è difficile garantire equo accesso a talenti qualificati, che attualmente sono rari e costosi.

Un altro aspetto fondamentale da considerare è che gli agenti per funzionare in modo efficace devono recuperare informazioni dall'ambiente, dall'utente, da altri servizi e, inoltre, necessitano una profonda integrazione con l'hardware ed il software del dispositivo "contenitore" che li ospita, tipicamente uno smartphone o un computer. In questi termini, giocano un ruolo cruciale i produttori di dispositivi ed i proprietari dei sistemi operativi (OS), dal momento che essi si posizionano in una via intermedia tra input e canale distributivo. Nel caso in cui i produttori o proprietari di OS siano anche sviluppatori di AI agents diventa chiaro che tenderanno a limitare l'accesso ai propri device escludendo altre tipologie di agenti. Lo stesso discorso vale in caso di partnership e accordi esclusivi, che infatti sono sotto l'attenzione delle autorità di regolamentazione.

Inoltre, un altro collo di bottiglia potrebbe essere rappresentato da applicazioni esterne e servizi di terzi a cui gli agenti si appoggiano per portare a termine le proprie attività ed alleggerire quelle a capo dell'utente, in quanto potrebbero siglare accordi esclusivi solamente con alcuni collaboratori digitali, escludendone la restante parte. Basti pensare alla richiesta di acquisto di un certo prodotto: è necessario poter accedere ad un'applicazione di shopping o ad un browser per navigare in Internet, ad un sistema di pagamento e, infine, al calendario, in modo da poter programmare la consegna. Se anche solo uno di questi servizi venisse meno l'intera performance dell'agente sarebbe compromessa e non potrebbe raggiungere l'obiettivo fissato inizialmente.

In definitiva, il controllo sugli input essenziali allo sviluppo degli AI agents potrebbe influenzare la concorrenza per il mercato, che rischia di rimanere concentrata nelle mani di pochi attori. Questo solleva interrogativi sulla necessità di standard regolatori che garantiscano l'accesso equo a tali risorse, stimolando così la contestabilità del mercato.

5.7.2.2 Esclusione a valle

L'altra tipologia di esclusione di AI agents si verifica a valle ed è legata al contenitore che ogni agente necessita per poter operare in modo efficace. Come anticipato, si può trattare di un computer, di un dispositivo mobile, di un sistema operativo o di un motore di ricerca. I problemi sorgono nel momento in cui il fornitore del contenitore sviluppa internamente il proprio agente intelligente o, in alternativa, ne privilegia uno a causa di partnership o accordi di *revenue sharing*. Questo porta a problematiche di tipo concorrenziale in quanto altri AI agents riscontrano di conseguenza difficoltà a raggiungere una larga fetta di utenti, i quali sono legati al collaboratore digitale offerto dal provider del contenitore tramite preinstallazione o maggiore visibilità.

Esempi emblematici della dinamica che si viene a creare sono rappresentati da Android ed iOS, che favoriscono rispettivamente l'utilizzo di Google Search e Siri. Gli utenti sono indotti a utilizzare tali servizi, senza essere incentivati a ricercare alternative, che quindi rimangono inesplorate e con meno possibilità di crescita. Questo porta a concorrenza sleale e soprattutto al guadagno di un vantaggio competitivo ingannevole, in quanto i consumatori non sono nemmeno consapevoli degli altri servizi disponibili. La scelta rimane quindi limitata e soprattutto correlata a problematiche come il *lock-in* ed i *switching costs*. Nel momento in cui si vuole passare ad un altro servizio, si incontrano costi economici, legati ad esempio all'installazione di nuove applicazioni, e costi cognitivi: l'utente, infatti, deve abituarsi a nuove opzioni, interfacce e modi di interagire coi sistemi digitali.

La pratica di *self-preferencing* appena descritta porta ad una diminuzione della contestabilità del mercato in quanto aumenta le barriere all'ingresso per i nuovi entranti, rafforzando il potere degli incumbent.

In definitiva, l'esclusione a valle mette in evidenza come la proprietà ed il controllo del contenitore costituiscano una leva competitiva cruciale per condizionare in modo significativo la scelta degli utenti e ridurre le opportunità di successo dei concorrenti. Ciò rischia di rafforzare la posizione degli incumbent ma soprattutto di erodere la contestabilità complessiva del mercato.

5.7.3 Esclusione da parte degli AI agents

Lo scenario che si contrappone a quello appena descritto riguarda l'esclusione da parte degli agenti intelligenti. Essi, seppure trovandosi attualmente in uno stadio embrionale, sono in via di sviluppo e destinati ad assumere un ruolo centrale nei processi decisionali, che per ora sono in capo all'uomo. Il potere che essi acquisiranno permetterà loro di influenzare la domanda in

quanto si troveranno ad agire per conto dei consumatori, concludendo acquisti, transazioni e, più in generale, azioni, non sempre in modo totalmente trasparente.

Anche in questo scenario si presenta la problematica relativa alla *self-preferencing*: il potere decisionale e la capacità di raccomandazione degli agenti possono essere sfruttati per promuovere i servizi del loro sviluppatore. In questi termini, emergono i concetti di *leveraging* e *tying*. Il primo permette di affermarsi in un mercato sfruttando la dominanza in un altro, come avvenuto nel caso Google Search/Google Shopping; il secondo lega un servizio al prodotto principale, come ha fatto Apple tramite la preinstallazione di Apple Store o Meta con l'integrazione di Meta AI nelle piattaforme controllate. Se nel caso dei motori di ricerca il comportamento sleale è abbastanza evidente, per quanto riguarda i nuovi collaboratori la situazione si complica in quanto vi è un deficit di trasparenza. Quando un agente analizza l'ambiente, prende una decisione ed agisce, si muove "dietro le quinte". L'utente viene informato quando tutto l'iter è concluso, senza avere informazioni su ciò che è accaduto nel mentre. Egli potrebbe richiedere specificità legate al processo, ma è raro che i consumatori si pongano domande in questo senso. L'agente nasce per semplificare processi, doversi informare su tutto il meccanismo sottostante all'AI agent rallenterebbe le procedure e snaturerebbe la logica stessa del nuovo paradigma. La monopolizzazione del ranking, quindi visibile solo all'agente, può portare ad una riduzione del benessere in quanto spinge ad adottare scelte che non massimizzano l'utilità dell'utente ma che sono influenzate dalla proprietà dei servizi correlati. Si potrebbe, quindi, essere vittime di acquisti avvenuti a condizioni di prezzo svantaggiose o ritrovarsi con la ricezione di ordini contenenti prodotti di qualità inferiore allo standard che avrebbero potuto garantire fornitori non affiliati.

Inoltre, uno dei principali rischi che le autorità regolatorie stanno affrontando riguarda il fatto che potrebbe affermarsi una dinamica del tipo "*winner-takes-all*", in cui un singolo agente prende il pieno controllo. Per arrivare a tale contesto, si devono aggiungere effetti di rete indiretti, legati ad esempio al fatto che applicazioni e servizi di terzi potrebbero dover essere ripensati e ristrutturati per garantire un fit preciso con i collaboratori digitali. Questo ovviamente porta a costi e, di conseguenza, i fornitori di tali app e servizi potrebbero affidarsi ad un solo agente. Nel far ciò potrebbero scegliere colui che è più popolare, in quanto dotato di più collegamenti con altre piattaforme e, quindi, più efficiente. Questo rafforza ancora di più la concentrazione del mercato, innalzandone ulteriormente le barriere all'ingresso. In secondo luogo, man mano che gli utenti si affidano ad un unico AI agent, si materializza sempre di più il rischio di *lock-in* indotto dai dati. Più un agente ha a disposizione dati, più sarà in grado di apprendere e di fornire un'esperienza personalizzata, quindi gli utenti saranno frenati nel

cambiare agente se non esiste la possibilità di portare con sé i dati e la cronologia delle interazioni dal momento dovrebbero ripartire da zero nel processo di “addestramento”.

In aggiunta, un AI agent potrebbe avere incentivo nel rifiuto a collaborare. Nel caso in cui l'azienda sviluppatrice possieda anche un sistema operativo, potrebbe voler escludere altri OS dall'utilizzo del proprio agente, in modo da favorire il proprio ecosistema ed espandere la propria quota di mercato, a discapito di altri operatori, che soffrirebbero della perdita di competitività. Meccanismi di questo tipo potrebbero portare ad una dinamica monopolistica. Come conseguenza, entra in gioco il diritto della concorrenza, il quale afferma che per poter forzare un'azienda a fornire un servizio (o un bene) devono essere soddisfatte le condizioni di indispensabilità ed eliminazione della concorrenza. Per quanto riguarda la prima, ci si riferisce alla possibilità di affidarsi ad opzioni alternative, come altri agenti, magari meno performanti o affermati nel mercato, ma che di base possano offrire le stesse funzionalità. La seconda invece, afferma che la forzatura all'accesso è possibile solamente se il rifiuto è in grado di eliminare tutta la concorrenza da parte dell'impresa che richiede il servizio. Il fornitore dell'AI agent, quindi, deve avere un potere di mercato sufficiente per influenzare in modo radicale la domanda del mercato. Attualmente, la situazione non soddisfa nessuna delle due condizioni in quanto essendo un mercato emergente non c'è ancora un agente predominante sugli altri, esistono possibilità di alternativa e la concorrenza non è eliminata a prescindere dal rifiuto di un agente, data l'ampia scelta.

5.8 Applicabilità del DMA

5.8.1 Criteri per la designazione

Come accennato del capitolo dedicato alla regolamentazione, il Digital Markets Act è lo strumento principale per disciplinare l'operato e limitare il potere delle grandi piattaforme digitali che vengono etichettate come gatekeeper, in modo tale da preservare la contestabilità dei mercati. Esso si basa su interventi ex ante, imponendo obblighi e divieti alle piattaforme prese in considerazione. La designazione di queste ultime come gatekeeper si basa su criteri quantitativi e qualitativi. Nel momento in cui vengono rispettate le soglie la piattaforma deve seguire una serie di obblighi relativi ai *Core Platform Services* (CPS) che offre. Essi sono le categorie di servizi digitali che il DMA considera dannose e suscettibili di incentivare dinamiche anticoncorrenziali. Se si presta attenzione ai criteri propri del DMA, riportati nella *tabella 13*, emergono problemi interpretativi, non tanto per quanto concerne il criterio inerente ai ricavi e la capitalizzazione di mercato, quanto invece al numero di utenti finali e commerciali che interagiscono con l'agente. Il DMA definisce utente commerciale come “qualsiasi persona

fisica o giuridica che agisca in una veste commerciale o professionale utilizzando servizi di piattaforma di base al fine di, o nel corso di, fornire beni o servizi agli utenti finali”. Nel contesto analizzato non è chiaro se sia l’utente a scegliere attivamente l’agente o se sia quest’ultimo ad appoggiarsi a servizi esterni senza la loro cooperazione attiva. I confini risultano quindi sfumati e poco chiari. Inoltre, il terzo criterio incorpora un ritardo temporale in quanto implica che se anche un agente venisse designato come gatekeeper, questo avverrebbe al più presto tre anni dopo la sua adozione iniziale. In definitiva, risulta evidente che attualmente nessun AI agent rispetti tali soglie in quanto il mercato è in via di sviluppo e ancora non si può parlare di posizione dominante. Come già espresso nei paragrafi precedenti, però, ci sono preoccupazioni relative al fatto che le Big Tech potrebbero sfruttare il proprio potere per espandersi in questo nuovo mercato in via di evoluzione. A tal proposito, si considera la possibilità di etichettare i collaboratori come “gatekeeper emergenti” in modo tale da poter applicare le disposizioni proposte dal DMA. Per fare ciò ci si può affidare ad indagini di mercato che confermino che almeno i criteri qualitativi siano soddisfatti e che, quindi, il fornitore gode di una posizione consolidata e durevole rispetto al proprio AI agent. Tale operazione, tuttavia, risulta ancora complessa data la fase iniziale dell’adozione degli agenti, che quindi non risultano in linea col profilo definito. Il contesto, però, sta mutando velocemente, di conseguenza la situazione potrebbe dare luce ai primi collaboratori digitali dominanti, i quali potrebbero diventare soggetti agli obblighi e divieti imposti dal DMA.

Tabella 13: Criteri di designazione dei gatekeeper nel DMA

Categoria	Criterio qualitativo	Criterio quantitativo
Impatto sul mercato interno	L’impresa deve esercitare un’influenza rilevante all’interno del mercato europeo.	Ricavi all’interno dell’UE $\geq$ 7,5 miliardi di euro in ciascuno degli ultimi 3 esercizi finanziari oppure capitalizzazione media di mercato/valore equo $\geq$ 75 miliardi di euro nell’ultimo esercizio e fornitura dello stesso CPS in $\geq$ 3 Stati membri.
Gateway per collegamento tra imprese ed utenti finali	Il servizio deve essere un punto di accesso	CPS con $\geq$ 45 milioni di utenti finali attivi mensili e $\geq$

	fondamentale, che permetta alle imprese di raggiungere gli utenti finali.	10.000 utenti commerciali attivi annuali nell'UE (ultimo esercizio finanziario).
Posizione consolidata e stabile	L'impresa deve essere in grado di mantenere la sua rilevanza in modo durevole.	Il criterio relativo al gateway deve essere soddisfatto per almeno 3 anni consecutivi.

5.8.2 Categorizzazione dei CPS

Per poter stabilire se il DMA sia sufficiente per la regolamentazione dei recenti collaboratori digitali, risulta essenziale individuare la categoria di CPS in cui classificarli. Ne esistono dieci in totale, ovvero: servizi di social networking online, servizi di piattaforma di condivisione video, servizi di comunicazione interpersonale indipendenti dal numero, servizi di pubblicità online, servizi di intermediazione online, motori di ricerca online, sistemi operativi (OS), browser web, assistenti virtuali, servizi di cloud computing.

Escludendo i primi quattro a priori, rimangono gli altri sei. Tra essi alcuni attualmente risultano ancora poco accurati in quanto gli agenti sono ad uno stadio iniziale e non ancora completamente sviluppati. Ad esempio, essi si appoggiano a servizi cloud, seppur non coincidano con essi. Inoltre, presentano analogie con i motori di ricerca, in quanto effettuano query, raccolgono informazioni e restituiscono risultati. La differenza, però, sta nel fatto che gli agenti sono anche in grado di prendere decisioni ed agire, oltre che a fornire un elenco di pagine web. In aggiunta, essi utilizzano i browser web per poter essere informati e prendere decisioni, si appoggiano a sistemi operativi per l'interfaccia utente, ma ad oggi non si identificano precisamente nelle loro funzionalità, sebbene ci sia margine per poter arrivare a tale punto nel futuro prossimo. I servizi di intermediazione, invece, sono esclusi in quanto gli agenti non si occupano di costituire un luogo virtuale per l'incontro di domanda ed offerta ma piuttosto prendono decisioni autonome, formano loro stessi la curva della domanda da parte dei consumatori.

La categoria che risulta più appropriata, di conseguenza, è quella degli assistenti virtuali, nonostante l'aderenza non sia comunque perfetta. La differenza principale è che mentre device come Alexa si limitano a eseguire azioni sotto richiesta dell'utente o a fornire a quest'ultimo risposte, gli agenti agiscono in autonomia secondo una logica orientata ad obiettivi di lungo periodo, oltre a potersi adattare anche come interfaccia principale tra uomo e macchina. Di conseguenza, le strade percorribili sono due: o si ridefiniscono gli obblighi e divieti relativi agli

assistenti virtuali, ampliandone lo spettro, oppure si può istituire una nuova categoria di CPS apposita per gli AI agents.

5.8.3 Obblighi e divieti previsti per i CPS

In termini di restrizioni ed obblighi, come conseguenza al fatto che i collaboratori digitali potrebbero adattarsi a più CPS differenti, le possibilità sono ampie. Ciascuno di essi presenta un proprio profilo e specifiche regole da rispettare, che a volte possono essere condivise con più categorie. Esistono, però, degli obblighi universali, comuni a tutti i CPS, quali: divieto di aggregare e combinare dati personali collezionati da servizi diversi senza il consenso esplicito dell'utente, obbligo di portabilità dei dati su richiesta dell'utente in modo da poterne effettuare il recupero ed il trasferimento, restrizioni all'uso dei dati degli utenti commerciali per competere direttamente con loro. Questi tre vincoli sono necessari nel contesto formato dagli AI agents in quanto essi lavorano a stretto contatto con una mole ingente di dati, che conferisce loro potere e può generare fenomeni di *lock-in*; con tali restrizioni si incentiva al mantenimento della contendibilità del mercato.

Nel caso in cui gli AI agents venissero classificati come sistemi operativi, si applicherebbero disposizioni volte a consentire la disinstallazione di applicazioni e servizi preinstallati che non siano essenziali per il funzionamento del dispositivo; in tal modo l'utente avrebbe la possibilità di esplorare tra più alternative. In aggiunta, si presenta l'obbligo di garantire la piena interoperabilità ad agenti concorrenti: essi, quindi, devono essere posti nella condizione per cui abbiano accesso equo ad API e funzionalità di sistema proprie del gatekeeper.

Se, invece, l'agente viene classificato come browser o assistente virtuale vengono introdotte nuove restrizioni ed obblighi da rispettare. In primis è necessario garantire la possibilità di scelta effettiva e la modificabilità delle impostazioni di default, obbligo che può essere adempiuto tramite l'utilizzo dei cosiddetti *choice screens*. In secondo luogo è necessario assicurare l'interoperabilità verticale, ovvero la possibilità a fornitori terzi di hardware e software di poter funzionare correttamente con l'agente sviluppato dal gatekeeper.

Per la classificazione all'interno della categoria dei motori di ricerca emergono divieti legati all'auto-preferenza ed obblighi per la condivisione di dati di ranking, query, click e view nel rispetto dei termini FRAND (*fair, reasonable and non-discriminatory*). Quest'ultima disposizione è di complicata applicazione in quanto gli AI agents non sono tipici fornire risultati sotto forma di liste cliccabili o ranking, producono direttamente l'output. In aggiunta, bisogna garantire condizioni di accesso trasparenti, non discriminatorie ed eque per gli utenti commerciali.

Come detto in precedenza, gli assistenti virtuali sono la categoria di CPS che più si adatta ad accogliere gli AI agents. Gli obblighi relativi sono un'integrazione di quelli provenienti dalle altre classificazioni, di conseguenza i collaboratori digitali ricadono su un perimetro di restrizioni più ampio e completo. In particolare si tratta di: possibilità di scelta tramite *choice screens*, possibilità di modificare le impostazioni di default, obbligo di garantire interoperabilità verticale e di offrire condizioni di accesso eque, divieto di auto-preferenziazione. Oltre a questi, si aggiungono gli obblighi condivisi a tutte le categorie citati sopra.

Tuttavia, siccome il DMA non è stato progettato pensando agli AI agent, permangono delle zone di ambiguità. Infatti, non è ancora disciplinato il rischio che l'agente possa agire come livello intermedio obbligato tra l'utente e l'ecosistema digitale. Inoltre, le Big Tech potrebbero sfruttare l'ambiguità attuale per ridefinire i collaboratori in modo tale da non rientrare in nessun CPS e, conseguentemente, sfuggire agli obblighi. In tal caso si verificherebbero fenomeni di *under-enforcement*, con l'applicazione parziale o in ritardo di regolamenti e disposizioni.

In definitiva, il DMA può essere il punto di partenza per la regolamentazione degli AI agents. La presenza di zone di ambiguità, però, pone interrogativi e, soprattutto, porta alla necessità di revisione o di integrazione del regolamento, con la possibile introduzione di una categoria dedicata ai collaboratori digitali ed i rispettivi obblighi adatti alle loro caratteristiche specifiche. Inoltre, rimane applicabile la designazione degli agenti come "gatekeeper emergenti", in modo tale da poterli includere senza attendere una revisione legislativa formale, la quale richiederebbe tempo, non sempre disponibile per innovazioni dirompenti come potrebbe rivelarsi questa.

5.9 Politiche europee di sostegno e regolazione

Come ampiamente discusso nel paragrafo precedente, il DMA è un buon punto di partenza per la regolamentazione degli AI agents ma da solo non basta, a tal proposito si affiancano le politiche europee di sostegno e di regolazione. L'insieme di tali regole, disposizioni ed incentivi è volto a garantire lo sviluppo di un ecosistema competitivo ma soprattutto la contestabilità del mercato in via di formazione. In questi termini emergono le iniziative ed i documenti analizzati nel capitolo inerente alla regolamentazione, seppur essi non menzionino ancora in maniera esplicita delle linee guida specifiche da seguire per quanto riguarda gli agenti. Il tassello centrale rimane rappresentato dall'AI Continent Action Plan, in quanto nato appositamente per tutte le tecnologie correlate all'intelligenza artificiale. Le sandboxes regolatorie giocano un ruolo chiave in quanto permettono la sperimentazione controllata dei collaboratori, ancora in uno stadio embrionale. Grazie ad esse si può migliorare progressivamente consentendo di proporre successivamente prodotti più maturi e performanti, obiettivo ad oggi non ancora

raggiunto, come si vedrà nel prossimo paragrafo. A ciò si affiancano le AI Factories dell'EuroHPC Joint Undertaking ed il Chips Act europeo, che rimane essenziale per quanto riguarda la produzione di semiconduttori, necessari per le componenti sottostanti agli agenti.

Attualmente gli sviluppatori di agenti riscontrano ancora ostacoli tecnici e strategici, dovuti alla mancanza di standard comuni e API condivise, oltre che alla chiusura degli ecosistemi dominati dalle Big Tech. Le politiche sopra menzionate sono volte a mitigare le problematiche che si stanno verificando ma non sono l'unico strumento possibile: stanno emergendo, infatti iniziative a supporto per il raggiungimento di un mercato contestabile. In primis si promuove la portabilità degli agenti, ovvero la possibilità per l'utente di migrare da una soluzione ad un'altra, senza perdere la memoria, le preferenze, i permessi e, più in generale, le interazioni passate. In aggiunta, si sta pensando alla creazione di un *agent registry* europeo: si tratta di un registro pubblico dove i fornitori dichiarano gli agenti messi sul mercato, specificando e certificando le caratteristiche chiave. Inoltre, si stanno sviluppando politiche industriali volte a sostenere uno stack open source per agenti, con l'aspettativa di incentivare l'innovazione anche per le PMI, ridurre la possibilità di *lock-in* e garantire maggiore trasparenza e facoltà di verificare e controllare il funzionamento intrinseco dei sistemi digitali.

In parallelo si stanno sviluppando agenti open source, come Auto-GPT, LangChain Agents e Open Agents, i quali aiutano ancora una volta a rimuovere le barriere per le PMI, nonostante sollevino preoccupazioni in termini di sicurezza e governance. Infine, si discute sul lancio del progetto Gaia-X for Agents, che richiama il modello pensato per il cloud, ovvero uno spazio interoperabile di agenti intelligenti, in cui i dati possono essere scambiati in sicurezza e si possa costruire fiducia ed identità.

In sintesi, il mosaico di strumenti esistenti dimostra come l'Europa stia agendo per promuovere lo sviluppo dei sistemi di intelligenza artificiale e, con essi, degli AI agents. Questi ultimi ereditano la regolamentazione presente, alla quale si presume verranno affiancate regole specifiche ed iniziative capaci di favorire la concorrenza e la contestabilità del mercato.

5.10 Situazione attuale: fallimenti, hype ed aspettative per il futuro

Il mercato degli AI agents si trova attualmente in una fase sperimentale, governata da un forte hype ma, al contempo, dall'assenza di innovazioni efficienti. Per ora, i device che si sono avvicinati maggiormente al concetto di collaboratore intelligente sono stati due dispositivi agentici stand-alone: Rabbit R1, dall'omonima startup fondata da Jesse Lyu, e Humane AI Pin, ideato da due ex-dipendenti Apple, ovvero Bethany Bongiorno e Imran Chaudhri.

Considerando il primo device, riportato in *figura 14*, le aspettative del mercato erano alte. Esso era stato presentato al CES⁴ 2024 come un potenziale sostituto dello smartphone. Si trattava, quindi, di uno strumento che avrebbe dovuto avere le stesse funzionalità dei cellulari, ma con l'integrazione delle recenti scoperte nell'ambito AI. Infatti, la tecnologia innovativa che lo alimentava era la LAM (*Large Action Model*). Quest'ultima si distingue dai LLM in quanto non si limita a generare testo, immagini o altri contenuti, ma arriva ad anticipare i bisogni dell'utente e ad agire in autonomia. Tutto questo, però, solo dal lato teorico. Le recensioni dei consumatori purtroppo sono tutte concordi e negative. Lo strumento che era stato definito come rivoluzionario e capace di semplificare la vita, riducendo il tempo che l'uomo passa davanti allo schermo, in realtà si è rivelato un fallimento: di tutti i compiti che prometteva di poter svolgere, non ne portava a termine nemmeno uno in modo affidabile. Ad esempio, il livello di accuratezza nel riconoscimento di oggetti, come cibo, non era accurato. Oppure riproduceva la canzone sbagliata rispetto a quella richiesta dall'utente. Inoltre sono state rilevate problematiche relative allo stand-by ed alla durata della batteria. In sintesi, il device, pur presentando formalmente le caratteristiche degli AI agents, si comportava come uno smartphone obsoleto e con vari bug da risolvere.



Figura 14: Rabbit R1

Humane AI Pin (riportato in *figura 15*) è l'invenzione che affianca Rabbit R1. L'obiettivo era, anche per esso, di facilitare l'utente e ridurre il tempo di utilizzo di smartphone e computer, evitando di esporre l'uomo agli schermi digitali. Proprio per questo motivo nasce privo di display e dotato unicamente di un proiettore. Le prime criticità emergono proprio da quest'ultimo in quanto, seppur simbolo di innovazione, risulta poco pratico: la proiezione risulta efficiente solo in alcuni ambienti e con determinate condizioni di luminosità. Oltre a tali problemi di utilizzo, le promesse sulle funzionalità, ancora una volta, non sono state rispettate.

⁴ *Consumer Electronics Show*, tenutosi a Las Vegas dall'9 al 12 gennaio 2024

Anche qui, infatti, le recensioni lamentano problemi di latenza, di batteria e di accuratezza nello svolgimento delle azioni. L'esperienza utente si è rivelata insoddisfacente e non commisurata al prezzo proposto. Il device, in definitiva, è stato ritirato dal mercato nel febbraio 2025, ma recentemente è stato lanciato un progetto open source chiamato OpenPin, che ha l'obiettivo di dare una nuova possibilità al device.

I due esempi riportati sono emblematici nel rappresentare il divario esistente tra teoria e pratica ma, soprattutto, sono indicativi del fatto che il futuro non risiede in device stand-alone bensì in soluzioni *software-based* integrate all'interno degli ecosistemi digitali già consolidati, in modo tale da arricchirne i servizi e, soprattutto, migliorare l'esperienza dell'utente. Molte piattaforme si stanno già muovendo in questa direzione. OpenAI ha recentemente introdotto la possibilità di creare GPTs personalizzati, che consentono all'utente di costruire degli agenti su misura, adattando LLM preesistenti a specifici contesti. Microsoft ha lanciato Copilot, definito come “assistente conversazionale basato su intelligenza artificiale che consente di aumentare la produttività e semplificare i flussi di lavoro offrendo assistenza contestuale, automatizzando le attività di routine e analizzando i dati”. Esso è l'emblema dell'integrazione dei collaboratori intelligenti in ecosistemi già esistenti. Infatti, Copilot può affiancare l'uomo nella creazione di documenti, nell'analisi di dati, nella gestione di progetti o nella comunicazione. In sintesi, offre supporto nell'utilizzo del pacchetto Microsoft Office. Considerando, poi, l'ecosistema iOS, Apple prevede di lanciare nel 2026 Siri 2.0, che introdurrà una serie di caratteristiche innovative, le quali permetteranno di superare il solo comando vocale attualmente disponibile. L'assistente Apple permetterà di lavorare con terze parti, offrirà supporto imparando dalle abitudini e preferenze dell'utente e, di conseguenza, creerà un'esperienza ed un'interfaccia altamente personalizzata per ciascun consumatore. Infine, anche Rubrik, azienda di sicurezza informatica, si sta impegnando per supportare iniziative agentiche tramite Rewind.



Figura 15: Humane AI Pin

Quest'ultimo è un caso peculiare in quanto si tratta di un agente orientato alla memoria, improntato al recupero di informazioni e all'estensione cognitiva.

Gli esempi riportati dimostrano che l'AI agentica sta facendo progressi ma, nonostante ciò, alcune tecnologie abilitanti sono ancora immature. Un recente studio di Andrea Viliotti (2025), consulente in ambito AI, mette in luce le criticità attuali analizzando contesti aziendali. In primis, l'assenza di standard condivisi rende complicata l'attività di benchmarking. Questo porta alla valutazione degli agenti tramite simulazioni o test astratti. Inoltre, riemergono i limiti propri dei LLM, quali la generalizzazione e la fatica nel mantenere coerenza nelle sequenze di dati complessi. Tutto ciò, se già destava preoccupazioni nei precedenti sistemi di AI, amplifica i rischi in un contesto dominato da agenti, in quanto la direzione di sviluppo è quella della piena autonomia. In quest'ottica risulta essenziale migliorare le capacità di ragionamento e di *goal decomposition*.

In conclusione, il quadro riportato dimostra come gli AI agent attualmente siano ancora una sfida. I fallimenti di Rabbit R1 e Human AI Pin sono indice delle mancanze e dei limiti che solo il tempo e la ricerca potranno colmare. Numerosi leader (riportati nella *tabella 14*), però, sembrano propensi al supporto dell'AI agentica, come mostrato dalle varie iniziative che si stanno affermando sul mercato. Questo pone ottime basi, che se combinate con i progressi continui che stanno caratterizzando l'ambito AI, hanno potenziale per rivoluzionare l'economia digitale e il rapporto tra l'uomo e le macchine.

Tabella 14: Agenti più popolari nel mercato, sviluppatori e stato attuale di utilizzo

Agente	Sviluppatore	Stato
GPTs personalizzati	OpenAI	In uso
Copilot	Microsoft	In uso
Meta AI	Meta	In beta
Siri 2.0	Apple	In sviluppo
Gemini 2, Project Mariner	Google	In sviluppo
Bedrock Agents & Alexa +	Amazon	In uso/In sviluppo
Rewind	Rubrik	In uso
Auto-GPT & Open Agents vari	Community open source	In uso/In sviluppo
Claude Agent SDK	Anthropic	In beta

6 Conclusioni

L'analisi approfondita svolta attraverso i vari capitoli mostra come l'intelligenza artificiale rappresenti una discontinuità tecnologica al pari delle più grandi rivoluzioni del passato, come quella industriale o quella legata all'informazione. Sebbene gli effetti siano ancora sfumati a causa della mancata adozione su larga scala, emergono già segnali e impatti tangibili che suggeriscono che nel futuro si assisterà ad una ridefinizione degli equilibri in vari settori, quali l'economia, la politica, la tecnologia e la società nel suo complesso.

Considerando i principali aspetti macroeconomici, ci si trova attualmente in una fase di transizione. La crescita di produttività rimane limitata, con però grande potenziale di miglioramento per gli anni a venire, legato principalmente agli investimenti in complementi intangibili. Gli effetti in termini di disoccupazione rimangono in via definizione in quanto, sebbene risulti evidente che molti lavoratori dovranno riconfigurarsi per adattarsi al cambiamento, c'è ampio margine per la costruzione di nuove figure altamente specializzate per il settore. Il rischio principale è quello di vedere inasprite le disuguaglianze presenti: alcuni Paesi, specialmente quelli meno sviluppati, tendono a rimanere indietro, sia per la mancanza di asset fondamentali che per l'assenza di un quadro regolatorio adeguato e che incentivi allo sviluppo.

Sul piano concorrenziale la situazione non è stabile e, anzi, si prevedono importanti trasformazioni. L'attuale struttura dell'AI stack concentra il potere nelle mani di pochi attori globali attraverso il controllo dei dati, della potenza di calcolo e dei canali di distribuzione. L'equilibrio nei prossimi anni sarà determinato dalle scelte infrastrutturali e dal possibile ingresso di nuovi entranti, sebbene esso sia complicato a causa del vantaggio competitivo accumulato dagli attuali leader. In questo contesto il quadro regolatorio gioca un ruolo cruciale. Esso si basa fortemente sul precedente Digital Markets Act, in quanto sono coinvolti i mercati digitali. Quest'ultimo, tuttavia, non è sufficiente, vista la peculiarità dell'ondata tecnologica che si sta affermando. Di conseguenza, un ruolo centrale è e sarà ricoperto dall'AI Act, specifico per l'intelligenza artificiale, e dalle iniziative per la sovranità tecnologica.

Un nuovo paradigma che sta andando affermandosi, inoltre, ha il potenziale per stravolgere ulteriormente gli equilibri competitivi: si tratta degli AI agents. Essi sono in grado di processare informazioni complesse, prendere decisioni, agire e comunicare con l'ambiente esterno, altri sistemi digitali e piattaforme. Tali proprietà li rendono avanzati rispetto ai precedenti sistemi di intelligenza artificiale e, proprio per ciò, fanno emergere nuove preoccupazioni. Si potrà assistere, infatti, ad una esclusione *degli* agenti o *da parte* degli agenti. Il primo scenario è

correlato alle risorse in input ed ai canali di distribuzione, mentre il secondo ai meccanismi sottostanti ai collaboratori digitali, che potrebbero servirsi di pratiche anticoncorrenziali, come ad esempio la *self-preferencing*. In quest'ottica sarà necessario un ripensamento o adattamento delle categorie di CPS che attualmente il DMA propone, in modo tale da poter governare e controllare gli agenti nel modo più opportuno.

In conclusione, la sfida che attualmente caratterizza l'intelligenza artificiale è quella della costruzione di un ecosistema equo, sicuro ed aperto. Bisogna trovare un equilibrio tra innovazione e preservazione della concorrenza, tenendo conto degli aspetti su cui l'AI è in grado di impattare, quali l'economia, la società e la politica. A tale scopo, l'AI Act rappresenta lo strumento di riferimento, con i rispettivi obblighi e restrizioni in termini di utilizzo ma, soprattutto, di trasparenza, elemento cruciale nel quadro che si sta andando a formare, anche per quanto riguarda i recenti AI agents. Le scelte che vengono compiute nel presente plasmeranno il contesto futuro. Da esse dipenderanno la futura società, il rapporto tra l'uomo e la tecnologia ma, soprattutto, l'ambito di applicazione di quest'ultima: si potrà assistere all'affermazione di una nuova tecnologia abilitante oppure allo scenario più temuto della polarizzazione del controllo. L'intelligenza artificiale, in definitiva, è e sarà il punto di incontro tra conoscenza, potere e progresso.

Bibliografia

Abrardi, Laura, et al. *Artificial Intelligence, Firms and Consumer Behavior: A Survey*. Journal of Economic Surveys - Wiley, 2021.

Acemoglu, Daron, et al. “Ai and Jobs: Evidence from Online Vacancies.” *National Bureau of Economic Research*, 2020, <https://doi.org/10.2139/ssrn.3765910>.

Acemoglu, Daron, et al. *Return of the Solow Paradox? IT, Productivity, and Employment in US Manufacturing*. American Economic Review, 2014.

Acemoglu, Daron. *The Simple Macroeconomics of AI **. Massachusetts Institute of Technology, 2024.

AGCM Staff. *Competition in the Artificial Intelligence Tech Stack: Recent Developments and Emerging Issues*. Autorità Garante della Concorrenza e del Mercato, 2024.

Aghion, Philippe, and Simon Bunel. *AI and Growth: Where Do We Stand? **. Banque de France, 2024.

Agrawal, Ajay, et al. *Artificial Intelligence Adoption and System-Wide Change*. Journal of Economics Management and Strategy - Wiley, 2024.

Aguiar, Luis, et al. *Platforms and the Transformation of the Content Industries*. Journal of Economics Management and Strategy - Wiley, 2023.

Anitec-Assinform. *Agenti Di IA (Conoscere l’IA – Paper #3)*. Anitec-Assinform, 2025.

Armstrong, Mark. *Competition in Two-Sided Markets*. RAND Journal of Economics, 2006.

Bergeaud, Antonin. *Monetary Policy in an Era of Transformation the Past, Present and Future of European Productivity*. European Central Bank, 2024.

Bommasani, Rishi, et al. *On the Opportunities and Risks of Foundation Models*. Stanford University / Stanford Institute for Human-Centered Artificial Intelligence, 2021.

Borgogno, Andrea. *Large Language Models and International Coordination*. University of Turin – Department of Law, 2025.

Borreau, Marc, and Zach Meyers. *A Competition Policy for Cloud and AI*. Centre on Regulation in Europe, 2025.

Bostoen, Friso, and Jan Krämer. *AI Agents and Ecosystems Contestability*. Centre on Regulation in Europe, 2024.

Bostoen, Friso, and Jan Kramer. *Is the DMA Ready for Agentic AI?* CERRE, 2025.

Bresnahan, Timothy, and Shane Greenstein. “Technical Progress and Co-Invention in Computing and in the Uses of Computers.” *Brookings Papers on Economic Activity. Microeconomics*, vol. 1996, 1996, p. 1, <https://doi.org/10.2307/2534746>.

Brown, Zach Y., and Alexander MacKay. *Competition in Pricing Algorithms*. American Economic Journal: Microeconomics, 2023.

Brynjolfsson, Erik, and Andrew McAfee. *The Second Machine Age: Work, Progress, and Prosperity in a Time of Brilliant Technologies*. W. W. Norton & Company, 2014.

Brynjolfsson, Erik, and Tom Mitchell. *What Can Machine Learning Do? Workforce Implications*. Science, 2017.

Brynjolfsson, Erik, et al. *Machine Learning, Labor Demand, and the Reorganization of Work*. MIT, 2019.

Brynjolfsson, Erik, et al. *Economic Consequences of Artificial Intelligence and Robotics*. AEA Papers and Proceedings, 2018.

Brynjolfsson, Erik, et al. *How Will Machine Learning Transform the Labor Market?* Hoover Institution, 2019.

Brynjolfsson, Erik, et al. *The Productivity J-Curve: How Intangibles Complement General Purpose Technologies*. American Economic Journal: Macroeconomics, 2021.

Burke, Seán, et al. *Governing Artificial Intelligence in China and the European Union: Comparing Aims and Promoting Ethical Outcomes*. European Parliamentary Research Service, 2023.

Calvano, Emilio, et al. *Artificial Intelligence, Algorithmic Pricing, and Collusion*. *American Economic Review*. American Economic Review, 2020.

Cambini, Carlo. *AI and the Markets*. Politecnico di Torino and SIEPI, 2025.

Carugati, Christophe. *Competition and Generative Artificial Intelligence*. Centre on Regulation in Europe, 2023.

Carugati, Christophe. *Competition Law and Economics of Big Data: A New Competition Rulebook*. Université Paris II Panthéon-Assas, 2020.

Cheng, Hsing Kenneth, et al. *The Rise of Empirical Online Platform Research in the New Millennium*. *Journal of Economics Management and Strategy* - Wiley, 2023.

Chui, Michael, et al. *The Economic Potential of Generative AI - the next Productivity Frontier*. McKinsey & Company, 2023.

CMA Staff. *AI Foundation Models: Initial Report (Full Report)*. UK Competition and Markets Authority, 2023.

CMA Staff. *AI Foundation Models: Short Version*. UK Competition and Markets Authority, 2023.

CMA Staff. *AI Foundation Models: Summary*. UK Competition and Markets Authority, 2023.

CMA Staff. *AI Foundation Models: Technical Update Report*. UK Competition and Markets Authority, 2024.

Comunale, Mariarosaria, and Andrea Manera. *The Economic Impacts and the Regulation of AI: A Review of the Academic Literature and Policy Actions*. International Monetary Fund, 2024.

European Commission. *AI Continent Action Plan*. Brussels, European Commission, 2025.

Ezrachi, Ariel, and Maurice E. Stucke. *Virtual Competition*. *Journal of European Competition Law & Practice*, 2016.

Fabbri, Flavio. *AI, La Marcia Degli Small Language Models (SLM). Un'alternativa Più Accessibile E Sostenibile per Le Imprese*. Key4biz.it., 2025.

Fiordalisi, Mila. “Un Cloud Made in Europe Contro Il Monopolio Big Tech.” *Domani*, 10 May 2025.

Girod, Melanie. *Europe’s AI Sandboxes: Racing in the Sand?* Center for European Policy Analysis, 2025.

Goldfarb, Avi, and Catherine Tucker. *Digital Economics*. Journal of Economic Literature, 2019.

Goldfarb, Avi, et al. *Could Machine Learning Be a General Purpose Technology? A Comparison of Emerging Technologies Using Data from Online Job Postings*. National Bureau of Economic Research, 2020.

Gordon, Robert J. *The Rise and Fall of American Growth: The U.S. Standard of Living since the Civil War*. Princeton; Oxford Princeton University Press, 2016.

GovAI. *Market Concentration and the Implications of Foundation Models: The Invisible Hand of ChatGPT*. Centre for the Governance of AI, 2023.

Goyet, Manuel Mateo Goyet. *The Economic Foundations of AI Infrastructure: What Policymakers Need to Understand*. European Commission, 2025.

Goyet, Manuel Mateo. *The Economic Foundations of AI Infrastructure: What Do Policymakers Need to Understand?* 2025.

Hagi, Andrei, and Julian Wright. *Artificial Intelligence and Competition Policy*. Elsevier B.V., 2025.

Hagi, Andrei, and Julian Wright. *Artificial Intelligence and Competition Policy*. Elsevier B.V., 2024.

Hagi, Andrei, and Julian Wright. *When Data Creates Competitive Advantage*. Harvard Business Review, 2020.

Halaburda, Hanna, et al. *the Business Revolution: Economy-Wide Impacts of Artificial Intelligence and Digital Platforms*. Journal of Economics Management and Strategy - Wiley, 2024.

Hunt, Stefan, et al. *Will 2025 Be the Year of the Agent? A Primer for Competition Practitioners on the next Wave of AI Innovation*. Competition Law & Policy Debate, 2025.

Jin, Ginger Zhe, et al. *M&a and Technological Expansion*. Journal of Economics Management and Strategy - Wiley, 2023.

Kowalski, Klaus, et al. *Competition Policy Brief: Competition in Generative AI and Virtual Worlds*. European Commission, Directorate-General for Competition, 2024.

Kucharavy, Andrei, et al. *Large Language Models in Cybersecurity*. Springer Nature Switzerland AG, 2024.

Lu, Yingying, and Yixiao Zhou. *A Review on the Economics of Artificial Intelligence*. Journal of Economic Surveys - Wiley, 2021.

McElheran, Kristina, et al. *AI Adoption in America: Who, What, and Where*. Journal of Economics & Management Strategy - Wiley, 2024.

Meyers, Zach, and Marc Bourreau. *A Competition Policy for Cloud and AI*. CERRE, 2025.

Misch, Florian, et al. *Artificial Intelligence and Productivity in Europe*. International Monetary Fund, 2025.

Parlamento Europeo e Consiglio dell'Unione Europea. *Regolamento (UE) Del Parlamento Europeo E Del Consiglio Che Stabilisce Norme Armonizzate Sull'intelligenza Artificiale (AI Act)*. Gazzetta ufficiale dell'Unione Europea, 2024.

Parlamento Europeo e Consiglio dell'Unione Europea. *Regolamento (UE) 2023/1781 Del Parlamento Europeo E Del Consiglio Del 13 Settembre 2023 Che Istituisce Un Quadro Di Misure per Rafforzare L'ecosistema Europeo Dei Semiconduttori*. Gazzetta Ufficiale dell'Unione Europea, 2023.

Perrault, Raymond, et al. *AI Index Report 2024*. Stanford University, Human-Centered AI Institute, 2024.

Pino, Flavio. *The Microeconomics of Data – a Survey*. Journal of Industrial and Business Economics, 2022.

Redazione economica. “L’AI Generativa Entra Nelle Imprese: La Rivoluzione Secondo IBM E Granite.” *Il Sole 24 Ore*, 2025.

Ricci, Lara. “L’intelligenza Artificiale Cambia Le Regole Del Gioco.” *La Repubblica*, 14 Nov. 2024.

Rochet, Jean-Charles, and Jean Tirole. *Platform Competition in Two-Sided Markets*. Journal of the European Economic Association, 2003.

Sanjeev, Iraj, et al. *Will Productivity Growth Return in the New Digital Era? An Analysis of the Potential Impact on Productivity of the Fourth Industrial Revolution*. Bell Labs Technical Journal, 2017.

Sastry, Girish, et al. *A Computing Power and the Governance of Artificial Intelligence*. ArXiv – Cornell University, 2024.

Schaefer, Maximilian, and Geza Sapi. *Data Network Effects: The Example of Internet Search*. Deutsches Institut für Wirtschaftsforschung, 2019.

Schmid, Jon, et al. *Titolo: Evaluating Natural Monopoly Conditions in the AI Foundation Model Market*. RAND Corporation, 2024.

Schrepel, Thibault, and Alex Sandy Pentland. *Competition between AI Foundation Models: Dynamics and Policy Recommendations*. MIT Connection Science Group, 2023.

Schrepel, Thibault, and Jason Potts. *Measuring the Openness of AI Foundation Models: Competition and Policy Implications*. Tanford CodeX: the Stanford Center for Legal Informatics, 2024.

Stiftung, Bertelsmann. *EuroStack: A European Alternative for Digital Sovereignty*. Bertelsmann Stiftung – ReframeTech Project, 2025.

The Economist. “A New York Startup Shakes up the Insurance Business.” *The Economist*, 2017.

Uuk, Risto, et al. *A Taxonomy of Systemic Risks from General-Purpose AI*. Centre for the Governance of AI, 2024.

Xu, Fasheng, et al. *The Economics of AI Foundation Models: Openness, Competition and Governance*. SSRN / ResearchGate Working Paper, 2024.

Zhao, Wayne Xin, et al. *A Survey of Large Language Models*. ACM Computing Surveys, 2023.

Sitografia

Agenda Digitale. “L’IA Come Nuova Corsa All’oro: Chi Farà Fortuna E Come.” *Agenda Digitale*, 22 June 2023, www.agendadigitale.eu/mercati-digitali/lia-come-nuova-corsa-alloro-chi-fara-fortuna-e-come/.

Amazon Web Services. “Agenti Bedrock.” *Amazon Web Services, Inc.*, 30 Sept. 2025, <https://aws.amazon.com/it/bedrock/agents/>.

Anthropic. “Building Agents with the Claude Agent SDK.” *Anthropic.com*, 2025, www.anthropic.com/engineering/building-agents-with-the-claude-agent-sdk.

Baily, Martin Neil, et al. “Machines of Mind: The Case for an AI-Powered Productivity Boom.” *Brookings*, 10 May 2023, www.brookings.edu/articles/machines-of-mind-the-case-for-an-ai-powered-productivity-boom/.

Barbera, Diego. “Humane AI Pin Addio, Smetterà Di Funzionare Il 28 Febbraio.” *Wired Italia*, Wired Italia, 19 Feb. 2025, www.wired.it/article/humane-ai-pin-fine-progetto/.

Barbera, Diego. “Meta AI Su Instagram E Facebook, Come Usarla al Meglio.” *Wired Italia*, Wired Italia, 8 Apr. 2025, www.wired.it/article/meta-ai-instagram-facebook-messenger-come-fare/.

Barbera, Diego. “Rabbit R1 E Humane Pin AI, Le Prime Recensioni Sono Un Disastro.” *Wired Italia*, Wired Italia, 3 May 2024, www.wired.it/article/rabbit-r1-humane-pin-ai-recensioni-negative/.

Bastian Nominacher, et al. “Here’s How to Pick the Right AI Agent for Your Organization.” *World Economic Forum*, 23 May 2025, www.weforum.org/stories/2025/05/ai-agents-select-the-right-agent/.

Belcic, Ivan. “AutoGPT.” *Ibm.com*, 15 Oct. 2024, www.ibm.com/think/topics/autogpt.

Bertelsmann Stiftung. “Report - EuroStack – a European Alternative for Digital Sovereignty.” *Bertelsmann-Stiftung.de*, 14 Feb. 2025, www.bertelsmann-stiftung.de/en/our-projects/reframetech-algorithmen-fuers-gemeinwohl/project-news/eurostack-a-european-alternative-for-digital-sovereignty.

Casali, Annalisa. “AI Agent Autonomi E Collaborativi: Cosa Sono, Casi d’Uso E Vantaggi.” *Digital4*, 8 Jan. 2025, www.digital4.biz/marketing/ai-agent-cosa-sono-vantaggi-agenti-autonomi/.

Chokkattu, Julian. “Welcome to Zscaler Directory Authentication.” *Wired.it*, 2025, www.wired.it/article/rabbit-r1-assistente-ai-ces-2024/.

Commissione Europea. “Spazi Comuni Europei Di Dati.” *Plasmare Il Futuro Digitale Dell’Europa*, 2020, <https://digital-strategy.ec.europa.eu/it/policies/data-spaces>.

Commissione Europea. “Continente Dell’IA.” *Commissione Europea*, 2024, https://commission.europa.eu/topics/eu-competitiveness/ai-continent_it.

Commissione Europea. “Regolamento Sui Chip.” *Commissione Europea*, 2019, https://commission.europa.eu/strategy-and-policy/priorities-2019-2024/europe-fit-digital-age/european-chips-act_it.

Concorrenza, della. “AGCM - G7 Competition: Authorities against AI Risks.” *Agcm.it*, 2024, <https://en.agcm.it/en/media/press-releases/2024/10/G7-Competition-Authorities-against-AI-risks>.

Eadicco, Lisa. “Rabbit R1 First Impressions: How I’ve Been Using the Handheld AI Assistant so Far.” *CNET*, 2024, www.cnet.com/tech/mobile/one-day-with-the-rabbit-r1-how-ive-been-using-it-so-far/.

Eisenberg, Seth. “Humane AI Pin: A Disappointing Reality despite the Hype → Fatherhood Channel.” *Fatherhood Channel → the Ultimate Resource for Dads*, 27 Sept. 2024, <https://fatherhoodchannel.com/2024/09/27/humane-ai-pin-a-disappointing-reality-despite-the-hype/>.

EuroHPC. “AI Factories Systems.” *The European High Performance Computing Joint Undertaking (EuroHPC JU)*, 2023, www.eurohpc-ju.europa.eu/ai-factories/ai-factories-systems_en.

European Commission. “Dati Aperti.” *Plasmare Il Futuro Digitale Dell’Europa*, 2019, <https://digital-strategy.ec.europa.eu/it/policies/open-data>.

European Commission. “AI Factories.” *Shaping Europe’s Digital Future*, 2024, <https://digital-strategy.ec.europa.eu/en/policies/ai-factories>.

European Commission. “European Chips Act | Shaping Europe’s Digital Future.” *Digital-Strategy.ec.europa.eu*, 2023, <https://digital-strategy.ec.europa.eu/en/policies/european-chips-act>.

European Commission. “L’UE Lancia l’Iniziativa InvestAI per Mobilitare 200 Miliardi Di € Di Investimenti Nell’intelligenza Artificiale.” *Plasmare Il Futuro Digitale Dell’Europa*, 2025, <https://digital-strategy.ec.europa.eu/it/news/eu-launches-investai-initiative-mobilise-eu200-billion-investment-artificial-intelligence>.

European Commission. “Strategy for Data | Shaping Europe’s Digital Future.” *Digital-Strategy.ec.europa.eu*, 2025, <https://digital-strategy.ec.europa.eu/en/policies/strategy-data>.

European Data Protection Board. “Cos’è Il GDPR? | European Data Protection Board.” *Europa.eu*, 2016, www.edpb.europa.eu/sme-data-protection-guide/faq-frequently-asked-questions/answer/what-gdpr_it.

Goldman Sachs. “AI May Start to Boost US GDP in 2027.” *Goldmansachs.com*, 2023, www.goldmansachs.com/insights/articles/ai-may-start-to-boost-us-gdp-in-2027.

Kargwal, Aryan. “Agenti Verticali Di Intelligenza Artificiale: Comprendere l’IA Orientata Allo Scopo.” *Botpress.com*, 2025, <https://botpress.com/it/blog/vertical-ai-agents>.

Knight, Will, and Adamà Faye. “Google Svela Gemini 2 E Si Butta Sugli Agenti AI.” *Wired Italia*, Wired Italia, 12 Dec. 2024, www.wired.it/article/google-gemini-2-agenti-ai-astramariner-intelligenza-artificiale/.

Lana, Alessio. “L’intelligenza Artificiale Sbaglia: Un Innocente Finisce in Carcere Negli Usa.” *Corriere Della Sera*, 30 Apr. 2021, www.corriere.it/tecnologia/scuola-dad-didattica-in-presenza-futuro/notizie/intelligenza-artificiale-sbaglia-innocente-finisce-carcere-usa-f4ca2cb0-a9a4-11eb-8b01-83c2a483d7f5.shtml.

McKinsey & Company. “Economic Potential of Generative AI | McKinsey.” *Www.mckinsey.com*, McKinsey, 14 June 2023, www.mckinsey.com/capabilities/mckinsey-digital/our-insights/the-economic-potential-of-generative-ai-the-next-productivity-frontier.

Meta. “Introducing the Meta AI App: A New Way to Access Your AI Assistant.” *Meta | Social Technology Company*, 29 Apr. 2025, <https://about.fb.com/news/2025/04/introducing-meta-ai-app-new-way-access-ai-assistant/>.

Microsoft. “Che Cos’è Un Copilot E Come Funziona? | Microsoft Copilot.” *Microsoft.com*, 2025, www.microsoft.com/it-it/microsoft-copilot/copilot-101/what-is-copilot.

Murphy, Hannah. “Retraining Workers for the AI World.” *@FinancialTimes*, Financial Times, 3 June 2024, www.ft.com/content/a6cb4832-c5be-480d-911c-a9ba92c929d7.

OpenContent Scarl. “Stati Uniti d’America – Executive Order on the Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence.” *Biodiritto*, 29 Oct. 2023, www.biodiritto.org/AI-Legal-Atlas/AI-Normativa/Stati-Uniti-d-America-Executive-Order-on-the-Safe-Secure-and-Trustworthy-Development-and-Use-of-Artificial-Intelligence.

Panay, Panos. “Introducing Alexa+, the next Generation of Alexa.” *Aboutamazon.com*, US About Amazon, 26 Feb. 2025, www.aboutamazon.com/news/devices/new-alexa-generative-artificial-intelligence.

Pascucci, Spartaco. “Siri 2.0 Di Apple: Data Di Rilascio, Funzionalità E Compatibilità | Massa Carrara News.” *Massacarraranews.com*, 6 Sept. 2025, www.massacarraranews.com/tecnologia/39103/siri-2-0-di-apple-data-di-rilascio-funzionalità-e-compatibilità/.

Pierce, David. “Humane AI Pin Review: Not Even Close.” *The Verge*, 11 Apr. 2024, www.theverge.com/24126502/humane-ai-pin-review.

Pierce, David. “Rabbit R1 Review: Nothing to See Here.” *The Verge*, 2 May 2024, www.theverge.com/2024/5/2/24147159/rabbit-r1-review-ai-gadget.

Prakash, Lakshmi Devi. “Agentic AI Frameworks: Building Autonomous AI Agents with LangChain, CrewAI, AutoGen, and More.” *Medium*, 4 June 2025, <https://medium.com/%40datascientist.lakshmi/agentic-ai-frameworks-building-autonomous-ai-agents-with-langchain-crewai-autogen-and-more-8a697bee8bf8>.

PricewaterhouseCoopers. “AI Agent Survey: PwC.” *PwC*, 2025, www.pwc.com/us/en/tech-effect/ai-analytics/ai-agent-survey.html.

Redazione di Innovation Post. “AI Agents: Cosa Sono E Come Migliorano L’industria.” *Innovation Post*, 30 Apr. 2025, www.innovationpost.it/tecnologie/ai-agents-cosa-sono-come-e-perche-migliorano-lindustria/.

Rogers, Reece, and Adamà Faye. “ChatGPT, Come Crearne Uno Personalizzato.” *Wired Italia*, Wired Italia, 28 Dec. 2023, www.wired.it/article/chatgpt-gpt-personalizzati-come-crearli-guida/.

Roi Lipman. “AI Agents: Memory Systems and Graph Database Integration.” *FalkorDB Knowledge Graph Database*, 6 Nov. 2024, www.falkordb.com/blog/ai-agents-memory-systems/.

Rossetti, Andrea. “Classificazione E Categorie Di Rischio Nell’AI Act.” *Bnews*, 2024, <https://bnews.unimib.it/blog/classificazione-e-categorie-di-rischio-nellai-act/>.

Rubrik. “Rubrik Agent Rewind: AI Agent Resilience.” *Rubrik*, 2025, www.rubrik.com/products/agent-rewind.

Seivwright, Scott. “The J-Curve Concept Is Commonly Used in Economics, Finance, and Change Management. In the Context of Change Management, Particularly in Organizational and Project Settings, the J-Curve Represents the Initial Dip in Performance or Benefits Following the Implementation of a Change, before a Gradual Re.” *Linkedin.com*, 24 Nov. 2023, www.linkedin.com/pulse/j-curve-lean-agile-scott-seivwright--wuue/.

UK Parliament. “Written Questions and Answers - Written Questions, Answers and Statements - UK Parliament.” *Parliament.uk*, 2025, <https://questions-statements.parliament.uk/written-questions/detail/2025-03-27/hl6273>.

Viliotti, Andrea. “Agenti AI in Azienda: Performance Reali, Limiti E Strategie (Report 2025).” *Andrea Viliotti*, 16 July 2025, www.andreaviliotti.it/post/agenti-ai-in-azienda-performance-reali-limiti-e-strategie-report-2025#viewer-ga3od1033.