



**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
ruolo delle partnership strategiche**

**Relatore:
Prof. Carlo Cambini**

**Candidato:
Vincenzo Salvatore**

Anno accademico 2024-2025

1 Introduzione.....	4
2 Struttura mercato dell'AI stack.....	6
2.1 Introduzione al concetto di AI stack.....	7
2.1.1 L'AI stack nell'ecosistema dell'intelligenza artificiale	7
2.1.2 Funzioni concorrenziali e punti nevralgici dello stack	8
2.2 Struttura dell'AI Stack.....	9
2.2.1 Hardware e infrastrutture di calcolo	9
2.2.2 Servizi cloud	12
2.2.3 Dati.....	14
2.2.4 Foundation Models	18
2.2.5 Applicazioni.....	24
2.3 Dinamiche competitive	28
2.3.1 Concentrazione per livello e interdipendenze	28
2.3.2 Meccanismi economici comuni nell'AI stack	29
2.3.3 Integrazione verticale e leve cross-layer	30
3 Effetti economici ed organizzativi dell'AI	33
3.1 Occupazione.....	34
3.1.1. L'AI ed il futuro dell'occupazione	34
3.1.2 L'approccio <i>task-based</i> e le implicazioni organizzative	34
3.1.3 Evidenze empiriche.....	35
3.1.4 Salari e qualifiche	36
3.1.5 La polarizzazione occupazionale	38
3.2 Produttività e crescita.....	39
3.2.1 Il paradosso di Solow e la J-curve.....	39
3.2.2 Evidenze empiriche.....	41
3.2.3 Co-invenzione e disomogeneità tra imprese	43
3.2.4 Implicazioni macroeconomiche	44
3.2.5 Stime e previsioni quantitative sulla produttività.....	44
3.3 Organizzazione aziendale	48
3.3.1 Trasformazione delle strutture organizzative ed approccio <i>data-driven</i>	48
3.3.2 Modularità e interdipendenze	49
3.3.3 Capitale umano e complementarità	51
3.3.4 Cambiamento dei processi decisionali e nuovi ruoli organizzativi	51
3.3.5 Sfide organizzative dell'AI nelle imprese	52
3.4 Innovazione e modelli di business.....	54
3.4.1 Piattaforme digitali.....	54
3.4.2 AI come fonte di personalizzazione	54
3.4.3 Il futuro dell'AI generativa.....	55
3.5 Educazione.....	56
3.5.1 Trasformazione della domanda di competenze e AI literacy	56
3.5.2 Competenze umanistiche e nuove figure professionali	57
3.6 Concorrenza e struttura dei mercati.....	59
3.6.1 Effetti pro-competitivi	59
3.6.2 Effetti anti-competitivi	60
3.6.3 Discriminazione di prezzo e potere di mercato legato ai dati	61
3.6.4 Rischi di collusione nei mercati digitali	64
3.6.5 Fusioni ed acquisizioni: il caso americano.....	65
3.6.6 Venture Capital, imprenditorialità e barriere per i nuovi entranti	65
4 Regolamentazione e policy dell'AI	68

4.1 I <i>driver</i> della regolamentazione	68
4.2 Strumenti regolatori e governance	71
4.2.1 AI Act.....	71
4.2.2 Oltre l'AI Act: le altre norme europee di riferimento	72
4.2.3 Sandboxes.....	73
4.3 Contesto internazionale e cooperazione multilaterale	74
4.4 AI continent action plan EU	76
4.4.1 Obiettivo e scopo	76
4.4.2 Iniziative sulle risorse computazionali e le infrastrutture	77
4.4.3 Verso una governance del calcolo	80
4.4.4 Data Labs ed European data union strategy	80
4.5 Sfide e prospettive future	81
5 Partnership.....	83
5.1 Il ruolo delle partnership.....	84
5.2 Analisi comparativa delle partnership	86
5.3 Evoluzione delle principali partnership verticali e orizzontali	88
5.4 Partnership cross-industry	102
5.4.1 Sanità	102
5.4.2 Automotive	104
5.4.3 Telecomunicazioni	104
5.4.4 Ricerca scientifica e education.....	105
5.4.5 Servizi professionali.....	106
5.5 Benefici e rischi delle partnership	109
5.6 Prospettiva regolatoria e antitrust.....	111
6 Conclusioni.....	111
Bibliografia.....	114
Sitografia	121

1 Introduzione

L'intelligenza artificiale negli ultimi anni ha mostrato uno sviluppo notevole segnando una nuova fase di trasformazione tecnologica e industriale.

L'avvento di sistemi di apprendimento su larga scala, chiamati *foundation models* (FM), oltre a essere in grado di generare contenuti, testi, immagini o codice, ha anche ridefinito le possibilità di automazione, creatività e analisi dei dati, portandola ad essere considerata come una tecnologia a uso generale al pari di Internet e dell'elettricità.

I FM hanno permesso la rapida diffusione dell'AI generativa, portando ad accelerare la convergenza tra scienza, mercati e settori industriali, stimolando la ridefinizione delle catene del valore globali e dell'assetto competitivo delle imprese

Dal punto di vista economico, le stime di Goldman Sachs (2023) indicano che l'intelligenza artificiale potrebbe contribuire fino al 7% del PIL mondiale nel prossimo decennio; tale impatto deriva dall'aumento della produttività, dall'emergere di nuovi mercati e dai modelli di business *data-driven*. Il rischio di questa innovazione è che si creino concentrazioni di mercato, portando ad un controllo sulle risorse critiche come dati, potenza di calcolo e infrastrutture cloud da parte di pochi attori.

Da un punto di vista tecnologico, la catena del valore dell'AI presenta una struttura complessa e fortemente interdipendente; le grandi imprese digitali controllano segmenti chiave della filiera, aspetto che analizzeremo nei capitoli successivi. La Competition and Markets Authority (CMA, 2024) e la CE Policy Brief (2024) hanno rilevato che la presenza di un numero ristretto di attori dominanti in più stadi della catena rischia di generare *bottleneck effects*, restringendo l'accesso ai fattori produttivi essenziali e ostacolando la nascita di un ecosistema realmente competitivo e aperto.

Parallelamente, sta emergendo sempre con più forza la dinamica delle partnership in cui sviluppatori di modelli e fornitori di infrastrutture digitali stringono accordi, come avvenuto tra Microsoft e OpenAI, o Google e Anthropic. Il fenomeno menzionato mostra la tendenza di un sistema che evita le fusioni e che predilige le alleanze strategiche producendo effetti di controllo sostanziale lungo la catena del valore, anche in settori non direttamente legati all'AI generativa, ma dove essa viene incorporata come nei servizi professionali, nella sanità e nell'educazione. L'intelligenza artificiale si configura come un paradigma socio-tecnico destinato a incidere simultaneamente su produttività, organizzazione del lavoro, istruzione, innovazione e struttura dei mercati.

Da un punto di vista metodologico, il presente lavoro adotta una prospettiva interdisciplinare, per comprendere le dinamiche di valore, concorrenza e governance generate dall'AI. L'obiettivo è di esplorare come si articola la catena del valore dell'intelligenza artificiale, quali relazioni intercorrono tra i diversi livelli dello *stack* e come le partnership strategiche influenzano gli equilibri competitivi e le politiche regolatorie.

Nel complesso, la tesi intende offrire una lettura sistemica del fenomeno, mostrando come la combinazione di scala, dati e calcolo stia creando un nuovo ordine industriale centrato sull'intelligenza artificiale, in cui la competizione economica, la governance tecnologica e le politiche pubbliche risultano strettamente interconnesse. Solo una visione d'insieme, che tenga conto della catena del valore, delle relazioni verticali e delle implicazioni sociali, può consentire di comprendere appieno portata della rivoluzione in atto e di orientarla verso uno sviluppo sostenibile, aperto e inclusivo.

La tesi è articolata in cinque capitoli principali: il primo di stampo introduttivo sul lavoro svolto; il secondo si concentra sull'analizzare la catena del valore dell'AI; il terzo capitolo approfondisce gli effetti dell'AI sull'economia e sulla società. Il quarto capitolo tratta la regolamentazione e la policy sull'AI concentrandosi sul contesto europeo; l'ultimo capitolo è dedicato alle partnership settoriali a quelle *cross-industry*, sottolineando i problemi di concentrazione che si generano. La stesura del secondo, terzo e quarto capitolo è stata svolta con l'ausilio della collega Elena Zuccolo, mentre la restante parte è stata svolta 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.

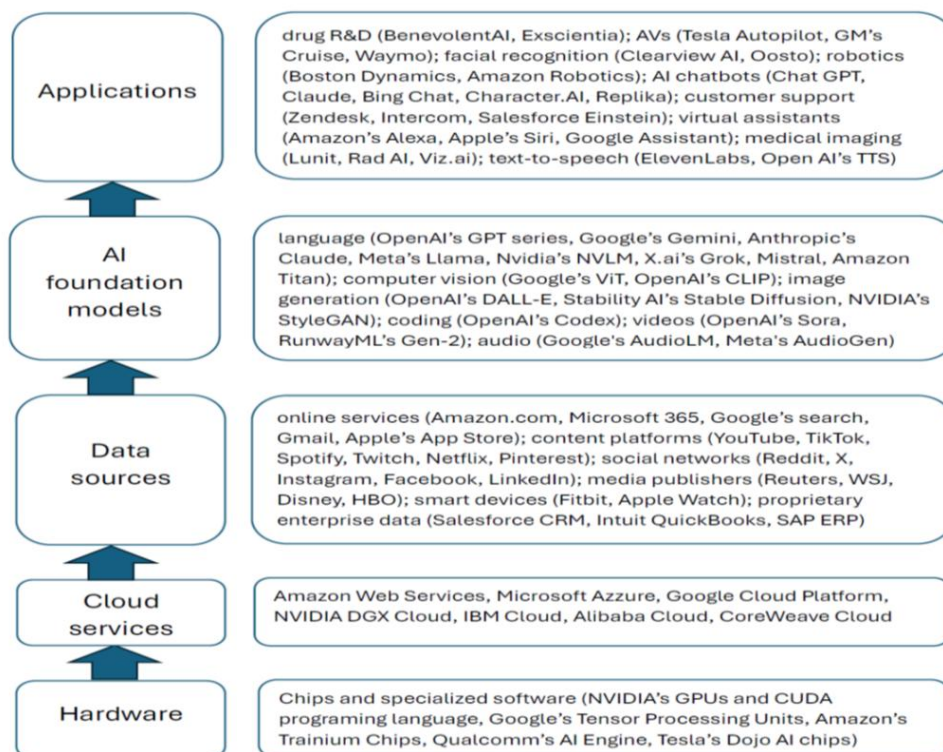


Figura 1: Rappresentazione dei livelli dell'AI stack.

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.

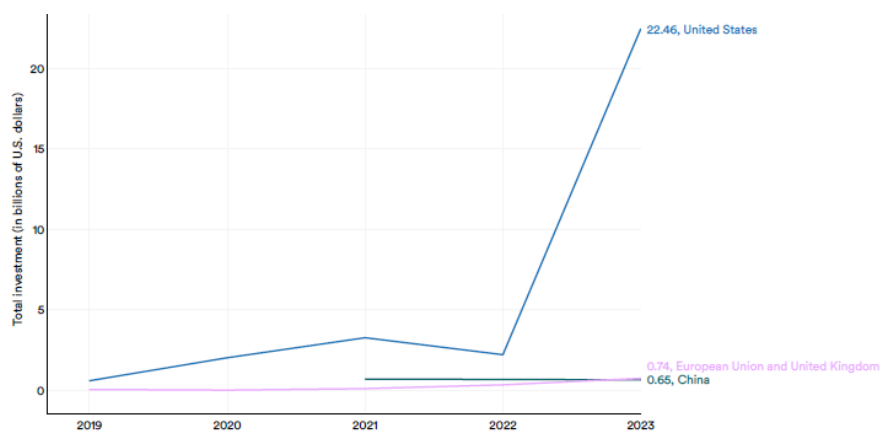


Figure 4.3.11

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,	Pubblico	Completo	Open-source

		comunità scientifica internazionale			
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.

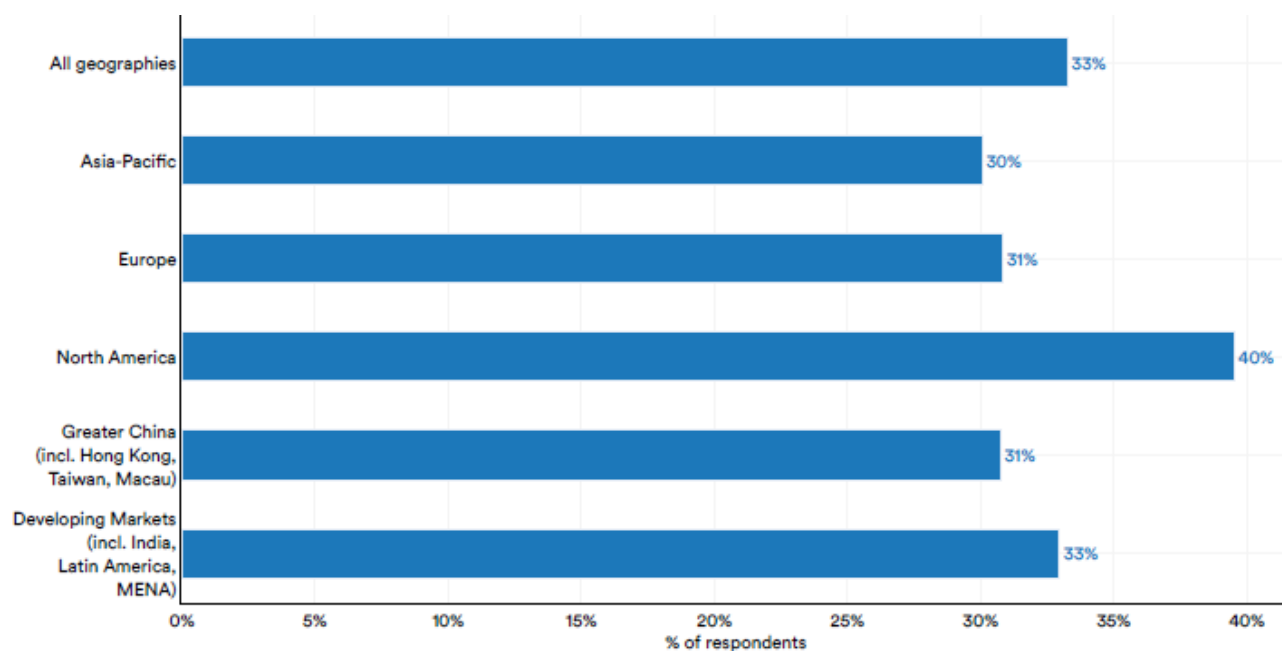


Figura 3: Tasso di adozione dell'AI generativa dalle organizzazioni nel mondo nel 2023

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.

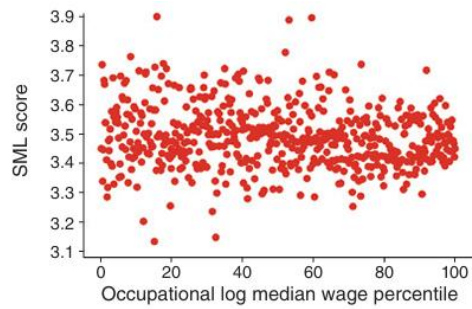


Figura 4: punteggio SML rispetto al percentile del salario mediano logaritmico occupazionale

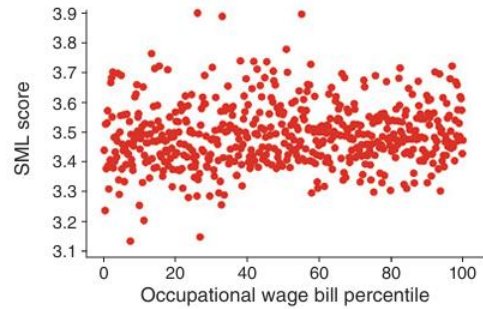


Figura 5: punteggio SML rispetto al percentile della massa salariale professionale

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.

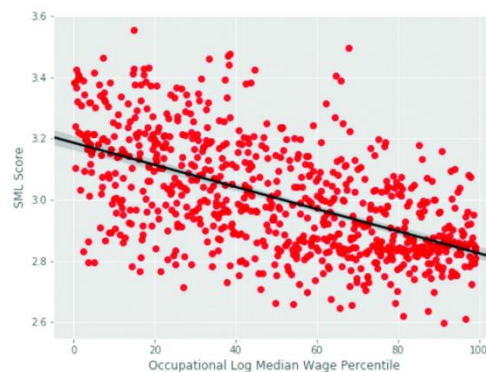


Figura 6: punteggio SML rispetto al percentile del salario mediano nel 2016 (coefficiente di regressione: -0.0034)

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 tre 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.

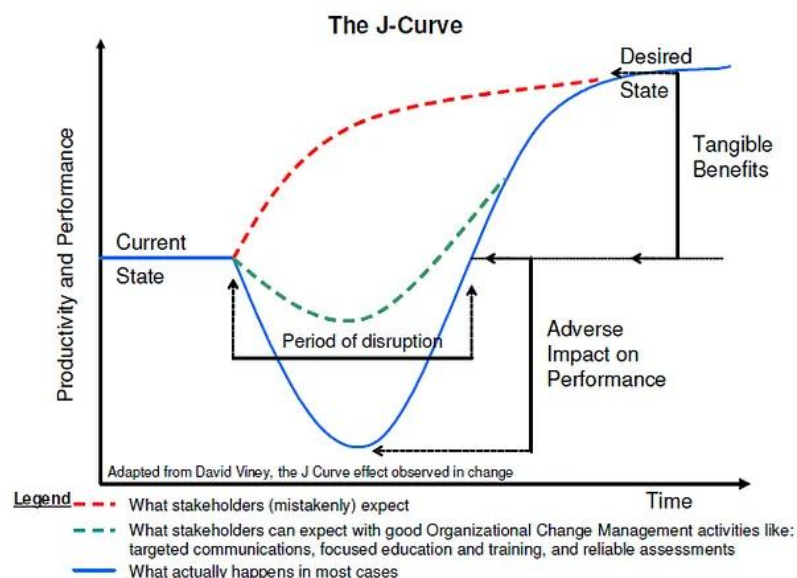


Figura 7: Andamento J-curve

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 (Hagiu & 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 alle 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 (Ezrahi & 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 frenare 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.

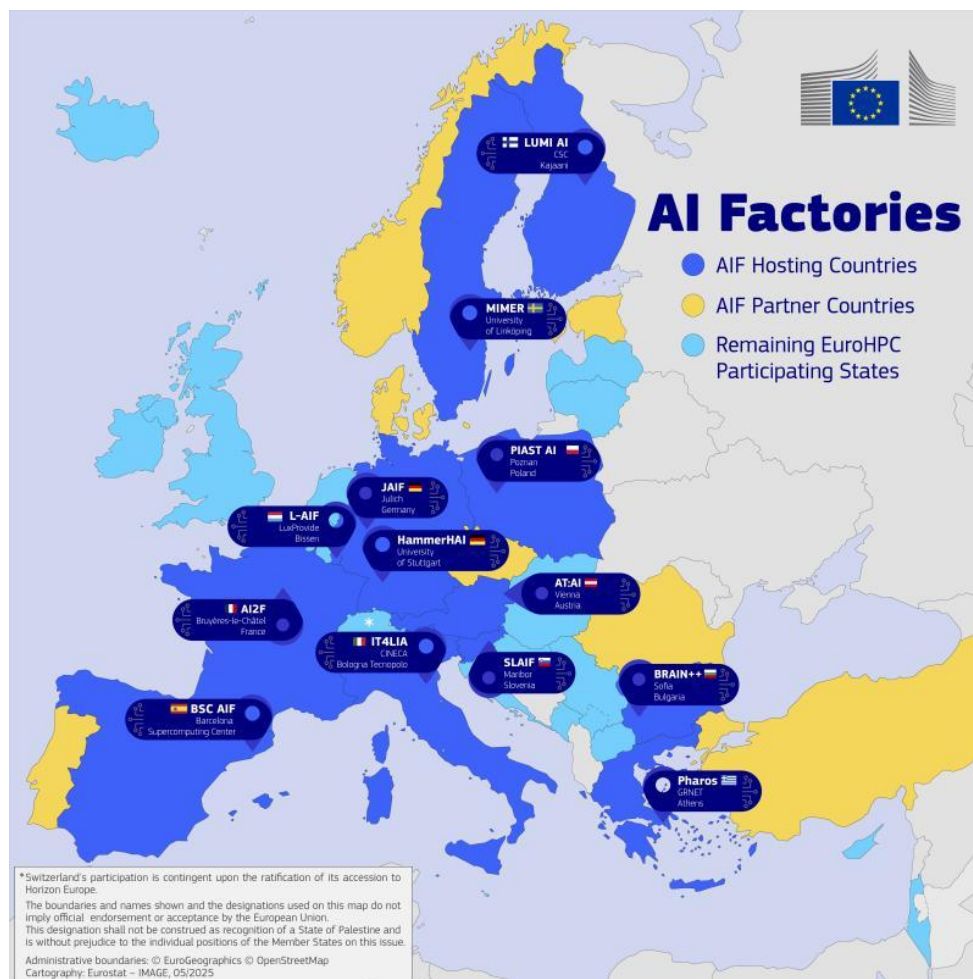


Figura 9: elenco AI Factories dell'UE

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 Partnership

Le partnership strategiche svolgono un ruolo cruciale nello sviluppo dell'intelligenza artificiale generativa, portando numerosi progressi. Queste consistono in forme di cooperazione tra grandi imprese tecnologiche, startup e, in misura crescente, tra attori istituzionali e finanziari. Il CMA nel report del 2024 ha identificato più di 90 partnerships tra le aziende GAMMAN² e le startup di intelligenza artificiale. La ragione che spinge questi soggetti a stringere accordi è la necessità di risorse eterogenee e altamente costose, richieste dal settore AI per realizzare e implementare i modelli di frontiera che nessun singolo attore riuscirebbe a sostenere in maniera efficiente da solo.

In questo contesto le partnership diventano il meccanismo principale che caratterizza la catena del valore dell'AI generativa. Il primo accordo ufficiale siglato tra due importanti realtà in questo ambito è stato nel 2016 con Microsoft e OpenAI, a dimostrazione che il modello basato solo sulle capacità interne non fosse sufficiente a sostenere la traiettoria di crescita dell'AI; quindi si è iniziato a cooperare con logiche di complementarità, dando vita a forme ibride che uniscono elementi di collaborazione e di rivalità. Questo fenomeno sottolinea come le industrie ad alta intensità tecnologica, avendo un tipo di innovazione cumulativa, abbiano la necessità di accedere tempestivamente a *know-how*, dataset e infrastrutture critiche, a causa anche del bisogno di ridurre il cosiddetto *time to market*, cioè il tempo di arrivo sul mercato, soprattutto per quanto riguarda i modelli e le sue applicazioni.

La centralità delle partnership nell'AI si spiega anche osservando la natura stessa dei *foundation models*; questi, essendo *general purpose technologies* e applicabili a molteplici contesti, richiedono un adattamento e una specializzazione che è raramente possibile nei confini di una singola impresa. Per tale motivo, gli sviluppatori di modelli hanno stretto alleanze sia verticali, con i grandi cloud provider che offrono potenza computazionale e capacità di distribuzione globale, sia orizzontali, con altri sviluppatori o con realtà open source. Un altro tipo di accordo, sempre più frequente, consiste nel cosiddetto *cross-industry*, che permette di integrare l'AI in filiere tradizionali.

Un ulteriore elemento che spiega la diffusione di queste forme collaborative è l'elevata incertezza che ancora caratterizza il settore. Del resto, la traiettoria tecnologica dell'AI generativa è sì rapida, ma allo stesso tempo non è ancora definito quale sarà lo standard che si affermerà come riferimento su larga scala. In questo contesto la partnership svolge una funzione

² L'acronimo GAMMAN indica le big tech Google, Amazon, Meta, Microsoft, Apple, e Nvidia.

di condivisione del rischio: investire in comune, piuttosto che acquisire o sviluppare tutto internamente, riduce l'esposizione finanziaria e aumenta la flessibilità strategica.

Le alleanze che si vanno a creare assumono diverse forme di accordi: equity parziale, licenze reciproche, joint ventures che permettono di costruire cooperazioni dinamiche senza dover ricorrere a fusioni o acquisizioni complete.

Per questo motivo le autorità antitrust, come la CMA britannica e la FTC (*Federal Trade Commission*), hanno definito queste partnership come assimilabili a fusioni: operazioni che, pur non configurandosi giuridicamente come fusioni o acquisizioni, hanno effetti economici analoghi in termini di potere di mercato e dinamiche concorrenziali tipiche delle concentrazioni. In definitiva, le partnership strategiche consentono l'accesso alle risorse fondamentali, accelerano i processi di innovazione e contribuiscono a plasmare la struttura stessa della concorrenza nei mercati. La loro analisi permette di comprendere le dinamiche che stanno ridisegnando i rapporti di forza tra attori tecnologici, startup emergenti e governi, e le implicazioni che emergeranno in materia di regolazione nei mercati digitali.

5.1 Il ruolo delle partnership

Lo studio delle partnership in economia industriale e in strategia aziendale permette di comprendere le alleanze che si hanno nel settore dell'intelligenza artificiale generativa.

A differenza delle relazioni contrattuali ordinarie, basate su scambi di fornitura dietro a compenso, le partnership tecnologiche implicano una condivisione strutturale di risorse, conoscenze e rischi con progetti che hanno un orizzonte temporale di lungo periodo. Per questo motivo, la seguente tipologia di accordi permette di affrontare al meglio l'incertezza e la forte intensità di capitale che contraddistinguono il settore AI.

Le partnership secondo l'autore DePamphilis (2023) possono assumere diverse configurazioni.

Le principali sono quattro:

- Le *equity partnership* si hanno quando le imprese acquisiscono partecipazioni reciproche rafforzando così la stabilità e l'impegno comune.
- Le *joint ventures* sono caratterizzate dalla separazione dalle imprese di origine, con la conseguente creazione di una nuova entità giuridica, e sono spesso utilizzate per lo sviluppo di una tecnologia o di un prodotto specifico.
- Il *licensing agreements*, che consiste nella concessione da parte di una società ad un'altra del diritto di utilizzare determinate tecnologie dietro a un compenso.

- Le *strategic alliances* rappresentano una categoria meno formalizzata, che si basa sulla complementarità degli attori mettendo in comune *know-how* e i rispettivi mercati, con l'obiettivo di cogliere opportunità che ciascun attore, da solo, non riuscirebbe a sfruttare.

Queste tipologie individuate sono pensate per contesti tradizionali, ma risultano adatte anche a descrivere le dinamiche nel settore AI. Ad esempio, la collaborazione tra Microsoft e OpenAI può essere considerata come una forma ibrida tra equity e alleanza strategica. Al contrario il progetto Stargate promosso da OpenAI insieme a Oracle e SoftBank, fondo per finanziare aziende innovative in tutto il mondo, corrisponde a una vera e propria joint venture, pensata per realizzare un supercomputer di nuova generazione.

Le ragioni che spingono le imprese a stipulare partnership tecnologiche sono molteplici. Un primo fattore riguarda l'accesso a risorse complementari. Se si prende in esame il caso Nvidia si nota come questi accordi diano la possibilità di accedere ad asset che mancano internamente come applicazioni verticali o reti distributive.

Un secondo elemento è la condivisione dei rischi e dei costi: secondo Ansari, Gupta & Urmetzer (2025), le startup di AI si legano a grandi piattaforme per ridurre l'incertezza tecnologica e finanziaria, mentre le Big Tech investono in più progetti contemporaneamente per diversificare le proprie scommesse. Un ulteriore incentivo è legato alla reputazione dell'azienda stessa, perché associarsi a un attore consolidato nel settore fornisce un segnale di credibilità sul mercato e facilita l'attrazione di capitale e talenti.

Le partnership tecnologiche, una volta avviate, consentono di ottenere numerosi benefici. Tra i più rilevanti vi è la possibilità di sfruttare economie di scala e di scopo, condividendo infrastrutture costose come data center e GPU. Esse permettono di accelerare i processi di ricerca e sviluppo, grazie allo scambio di competenze e alla possibilità di accedere a dataset esclusivi. Un ulteriore vantaggio riguarda la diffusione delle tecnologie in settori industriali diversi: si pensi alle alleanze *cross-industry* come quella tra Nvidia e Mercedes nell'automotive o alla collaborazione tra AWS, Anthropic e SK Telecom, compagnia coreana, per applicazioni nel settore delle telecomunicazioni.

Accanto ai benefici, queste alleanze comportano anche rischi significativi. Uno dei principali è il *lock-in* tecnologico, poiché può accadere che un partner più debole resti vincolato all'infrastruttura dell'attore dominante, come avvenuto nel caso dei crediti cloud di AWS verso le startup AI o dell'esclusiva di OpenAI con Azure di Microsoft, casi che verranno analizzati nel dettaglio nei paragrafi successivi. Questo tema si ricollega agli squilibri contrattuali, dove le Big Tech possono far valere il proprio potere dominante inserendo clausole che consolidano

la dipendenza delle startup. Un ulteriore rischio è rappresentato dall'appropriazione della proprietà intellettuale, poiché la condivisione di *know-how* può tradursi in perdita innovativa. Infine, più da un punto di vista di governance queste alleanze, con più attori in gioco che richiedono coordinamento e unione di intenti, potrebbero sfociare in conflitti di obiettivi, rallentando l'innovazione anziché accelerarla.

Per analizzare gli impatti e gli esiti delle partnership bisogna concentrarsi su due aspetti. Il primo è la *polarity management theory* elaborata da Ansari et al. (2025), che descrive le partnership come un equilibrio dinamico tra poli opposti: dominanza e dipendenza, apertura e chiusura, innovazione e appropriazione; in questo quadro, il successo dipende dalla capacità di bilanciare tensioni contrapposte senza che uno dei poli prevalga definitivamente.

Il secondo è la teoria della complementarità degli asset presentata da Teece (1986), che sottolinea come la capacità di innovare non dipenda soltanto dalle risorse tecnologiche interne, ma dall'accesso a un insieme più ampio di asset complementari, che possono essere acquisiti o resi disponibili proprio tramite partnership e alleanze.

5.2 Analisi comparativa delle partnership

Nell'analizzare le partnership all'interno della catena del valore dell'AI generativa bisogna distinguere quelle verticali da quelle orizzontali, oltre alle alleanze *cross-industry* e le iniziative pubblico-private, ciascuna con caratteristiche e implicazioni specifiche.

Le partnership verticali connettono fornitori di infrastrutture cloud con startup specializzate nello sviluppo di *foundation models*. Le più rilevanti sono le intese tra Microsoft e OpenAI, Amazon e Anthropic, o Google e Anthropic, nelle quali i grandi provider di cloud mettono a disposizione risorse computazionali e capitali che le startup non potrebbero procurarsi da sole, ottenendo in cambio clausole che vincolano l'uso esclusivo o preferenziale delle proprie infrastrutture e, in alcuni casi, i diritti di licenza sulla proprietà intellettuale.

L'effetto di queste partnership è duplice: da un lato garantiscono l'accesso a risorse fondamentali e accelerano l'innovazione, dall'altro creano il fenomeno del *lock-in* che rafforza la dipendenza delle startup dagli *hyperscaler* e consolida le barriere all'ingresso.

Di contro le partnership orizzontali uniscono attori operanti nello stesso livello della catena del valore. In questo caso le principali alleanze hanno l'obiettivo di creare un ecosistema open-source. Le più importanti sono l'accordo tra Meta, Hugging Face e Scaleway per la creazione di un acceleratore europeo per startup di AI e la collaborazione tra IBM e Meta per promuovere standard aperti.

In questi casi la logica è opposta alla visione delle partnership verticali, poiché si punta non all'esclusiva o al controllo, ma alla condivisione e alla complementarità orizzontale. In tale modo si favorisce la circolazione del *know-how*, si amplia la pluralità dell'ecosistema e si sostiene l'alternativa ai modelli chiusi e proprietari. Tuttavia, l'incertezza di queste iniziative rimane la sostenibilità economica, essendo prive delle rendite monopolistiche tipiche delle grandi piattaforme.

Una terza categoria è rappresentata dalle partnership *cross-industry*, che collegano il settore dell'AI ad altre industrie, favorendo la contaminazione tecnologica e l'espansione in nuovi mercati. In queste alleanze sono presenti attori predominati nel campo dell'intelligenza artificiale come Nvidia che collabora con Mercedes nell'integrazione di modelli AI nella guida autonoma e nei sistemi di *infotainment*, oppure con Caltech, università californiana, per la ricerca avanzata. Altre realtà come AWS e Anthropic hanno formato alleanze con aziende attive nel campo delle telecomunicazioni, come SK Telecom, per portare i modelli generativi nell'ecosistema 5G e nell'*edge computing*.

Lo scopo di queste scelte strategiche non è quella del controllo degli asset di base, ma della diffusione applicativa. Infatti, questa categoria di partnership apre nuovi sbocchi industriali all'AI, con il rischio che l'influenza dei grandi cloud provider si estenda sempre di più a valle della catena.

Infine, assumono rilevanza crescente le partnership pubblico-private già analizzate nei capitoli precedenti, soprattutto in Europa, dove il tema della sovranità tecnologica ha spinto istituzioni e governi a sviluppare iniziative mirate. Programmi come EuroHPC, con la realizzazione di supercomputer come LUMI, Leonardo e MareNostrum5, o strumenti come il Chips Act e le AI Factories finanziate dalla Commissione Europea, vanno proprio in questa direzione.

L'obiettivo è garantire accesso equo alle risorse computazionali, ridurre la dipendenza da fornitori extraeuropei e promuovere la resilienza del sistema produttivo. In questo caso il motore dell'alleanza non è la ricerca di profitti immediati, ma la tutela di interessi collettivi e strategici.

Nel loro insieme, queste diverse forme di partnership offrono un quadro articolato delle dinamiche in corso. Le alleanze verticali tendono ad accrescere la concentrazione e a rafforzare il *lock-in*; quelle orizzontali promuovono apertura e pluralità; quelle *cross-industry* estendono l'AI in nuovi mercati; mentre le pubblico-private mirano a rafforzare la sovranità tecnologica e a riequilibrare le asimmetrie globali. La comparazione mette così in luce come le partnership non siano un fenomeno omogeneo, ma un mosaico di strategie differenti che, pur seguendo

logiche diverse, contribuiscono tutte a ridisegnare la struttura competitiva dell'AI a livello globale.

5.3 Evoluzione delle principali partnership verticali e orizzontali

Un aspetto importante da considerare sono le diverse partnership che si sono susseguite nel corso degli ultimi dieci anni: si è passati da accordi bilaterali limitati a iniziative miliardarie con strutture contrattuali complesse e di rilevanza geopolitica. Questa evoluzione mostra come il principale strumento per l'innovazione, l'accesso alle risorse essenziali e per consolidare le posizioni di mercato acquisite siano le alleanze.

Il primo accordo rilevante risale al 2016 con la collaborazione tra Microsoft e OpenAI, che ai tempi era una piccola organizzazione no-profit per lo sviluppo dell'intelligenza artificiale. In particolare, nel 2019 Microsoft ha investito un miliardo di dollari per ottenere due clausole chiave. La prima comprendeva l'uso esclusivo del proprio cloud Azure per quello che riguardava l'addestramento e il *deployment* dei modelli di OpenAI. La seconda permetteva di appropriarsi la licenza esclusiva sulle tecnologie AI di OpenAI; infatti, altri soggetti non potevano avvalersi della stessa licenza con gli stessi diritti a meno di un rilascio pubblico. Questo accordo ha segnato l'inizio di una partnership ibrida, tra equity e strategica, che ha permesso a Microsoft di integrare nei propri servizi aziendali l'AI di OpenAI, come mostra la progressione negli anni nella *tabella 13*, nella quale vengono riassunte anche le altre partnership attive della Big Tech.

Tra il 2021 e il 2023 ha portato ad ulteriori investimenti fino a un totale di circa 13 miliardi di dollari, accompagnati dal lancio di prodotti Microsoft con ChatGPT, tra cui GitHub Copilot, Bing Chat e Office. In questa fase si ha evidenza del *lock-in* infrastrutturale instauratosi, poiché, per allenare il proprio modello, l'azienda utilizza esclusivamente supercomputer su Azure con un'integrazione stretta dei propri modelli nell'ecosistema della Big-Tech.

Nel 2021 è stato lanciato GitHub Copilot, come prodotto co-sviluppato da Microsoft e OpenAI, che si basava su GPT-3 e successivamente su GPT-4 e GPT-5, e che viene utilizzato per generare suggerimenti di codice e completamenti automatici per i programmatori in fase di scrittura del codice, aumentando così la produttività degli sviluppatori. Questo ha fornito a Microsoft il vantaggio di essere stato il primo ad integrare GPT nei suoi prodotti di sviluppo software. In questo caso si genera un *lock-in* che deriva dall'utilizzo di Copilot integrato, come già detto, con i modelli GPT di OpenAI. Infatti, gli sviluppatori che adottano Copilot si trovano vincolati all'uso combinato di GitHub, Visual Studio Code e Azure, con costi elevati di switching verso piattaforme concorrenti.

Un meccanismo simile si instaura con Bing Chat integrato all'interno del motore di ricerca di Microsoft Bing dal 2023, che successivamente è stato ampliato a Microsoft Edge. Bing Chat utilizza GPT-5 di OpenAI per generare risposte intelligenti e pertinenti alle domande degli utenti, simile a un chatbot avanzato. In questo modo gli utenti che preferiscono l'interazione AI su Bing sono incoraggiati a rimanere nell'ecosistema nel suo complesso, utilizzando Azure, Edge, e Windows, e andando così a rappresentare un altro esempio di lock-in, dove l'integrazione di GPT con il motore di ricerca e il browser spinge gli utenti a rimanere con Microsoft per un'esperienza AI completa.

Tabella 13: Partnership di Microsoft.

Anno	Partnership	Contributo delle parti	Rischio concorrenziale
2016	Microsoft - OpenAI	Microsoft fornisce infrastrutture cloud iniziali; OpenAI modelli sperimentali	Limitato, ma segna l'inizio di una dipendenza infrastrutturale
2019	Microsoft – OpenAI (\$1 mld)	Microsoft: Azure esclusivo con licenza esclusiva; OpenAI: know-how e modelli	Lock-in infrastrutturale e blocco licenze che limita l'accesso di altri player
2021	Microsoft – OpenAI (Copilot)	Microsoft: GitHub e Visual Studio Code; OpenAI: GPT-3 e modelli successivi da implementare in Copilot	<i>Lock-in</i> per sviluppatori legati a GitHub e Azure
2023	Microsoft – OpenAI (Bing Chat / Copilot su Edge)	Microsoft: motore di ricerca Bing e browser Edge; OpenAI: GPT-4 e successivi da implementare all'interno	Ecosistema chiuso che lega search, browser e sistema operativi

2024	Microsoft – Hugging Face (Azure AI Foundry, con Meta)	Microsoft: Azure marketplace; HF: modelli open; Meta: LLaMA	Rischio medio: Microsoft integra open nel proprio ecosistema chiuso, rafforzando posizione centrale
2025	Microsoft – Anthropic (Claude in Microsoft 365)	Microsoft: suite 365 e Copilot; Anthropic: modello Claude	Rischio medio: multi-modello attenua <i>lock-in</i> , ma rafforza la centralità di Microsoft

Nonostante queste evidenze il 15 aprile 2025 il CMA ha concluso che Microsoft esercita “*material influence*”, cioè ha la capacità di influenzare in modo sostanziale la direzione strategica di OpenAI, senza possederne la maggioranza o il controllo formale, quindi non esercita un controllo pieno.

In questo contesto altri *hyperscaler* hanno stretto accordi con realtà differenti. Amazon, come riportato nella *tabella 14*, nel 2023 ha annunciato un investimento fino a quattro miliardi di dollari in Anthropic, esteso a un totale di otto miliardi di dollari nel 2024, legando la startup e il suo modello Claude all’uso del cloud AWS per l’addestramento, della piattaforma Bedrock per la distribuzione e al co-sviluppo dei propri chip Trainium e Inferentia, con l’obiettivo di ridurre la dipendenza da Nvidia. Questo lega Anthropic non solo dal punto di vista tecnico, ma anche strategico, rafforzando il ruolo di AWS come cloud provider leader nel settore AI.

Tabella 14: Partnership Amazon.

Anno	Partnership	Contributo delle parti	Rischio concorrenziale
2023	Amazon – Anthropic (\$8 mld)	Amazon: AWS cloud e i chip proprietari Trainium/Inferentia; Anthropic: modello Claude	Dipendenza tecnica e strategica da AWS; rafforzamento <i>hyperscaler</i>

Nello stesso periodo Google, come mostrato nella *tabella 15*, ha siglato una partnership con la stessa azienda, investendo più di due miliardi di dollari, ottenendo una situazione di *equity non-voting*, cioè un diritto a una quota di utili senza diritto di voto con la possibilità di convertire il

debito in equity. Questa struttura permette che la partnership non venga classificata come fusione/acquisizione, come si evince dal CMA del 24 dicembre 2024. Allo stesso tempo consente a Google di mantenere un legame stretto con Anthropic e di avere un accesso privilegiato alla distribuzione dei suoi modelli. In aggiunta l'accordo prevede l'utilizzo da parte di Anthropic delle TPU, rafforzando l'ecosistema hardware di Google e la distribuzione di quei modelli su Vertex AI.

Queste strategie aumentano la concentrazione, ma allo stesso tempo segnano una competizione triangolare tra queste grandi realtà per assicurarsi partnership esclusive con i principali sviluppatori di *foundation models*.

Tabella 15: Partnership Google.

Anno	Partnership	Contributo delle parti	Rischio concorrenziale
2023	Google-Anthropic (> \$2 mld)	Google: TPU e Vertex AI; Anthropic: accesso modelli Claude oltre che la possibilità di avere equity non-voting	Struttura simile a M&A con rischi sulla concentrazione di potere di mercato
2025	Meta – Google Cloud (\$10mld per 6 anni)	Meta: modelli LLaMA; Google: cloud TPU	Rafforza Google Cloud come provider critico, aumenta concentrazione compute
2025	OpenAI – Google Cloud	OpenAI: modelli GPT; Google: data center e TPU	Rischio: OpenAI diversifica ma aumenta la dipendenza dagli hyperscaler
2025 (eventuale)	Apple – Google (Gemini in Siri)	Apple: Siri locale; Google: Gemini per query complesse/web	Rischio: collaborazione fra concorrenti diretti → concentrazione nell'assistenza vocale

Un'altra partnership rilevante, come indicato nella *tabella 16*, è stata annunciata da Apple nel 2024, nella quale si è accordata con OpenAI per integrare ChatGPT su iOS, rendendolo accessibile direttamente da Siri, l'assistente virtuale vocale di Apple. L'accordo non prevede un investimento diretto, ma uno scambio di visibilità e distribuzione: OpenAI ottiene l'accesso a centinaia di milioni di utenti iPhone, mentre Apple rafforza l'attrattiva dei propri device. Nel

complesso è un'alleanza di natura distributiva senza nessun trasferimento di capitale, con la possibilità da parte di Apple di entrare nell'ecosistema AI senza sviluppare necessariamente un proprio modello.

Tabella 16: Partnership Apple.

Anno	Partnership	Contributo delle parti	Rischio concorrenziale
2024	Apple – OpenAI (ChatGPT su iOS/Siri)	Apple: distribuzione su iPhone; OpenAI: ChatGPT	Rischio di lock-in distributivo su centinaia di milioni di device
2025 (eventuale)	Apple/Samsung – AI Perplexity	Apple/Samsung: smartphone come canale; AI Perplexity: modello di search AI	Potenziabile rischio: se sostituisce Google Search, ristruttura mercato, ma aumenta dipendenza dal nuovo attore
2025 (eventuale)	Apple – Google (Gemini in Siri)	Apple: Siri locale; Google: Gemini per query complesse/web	Rischio: collaborazione fra concorrenti diretti → concentrazione nell'assistenza vocale

Nel biennio 2023-2024, come si nota dalla *tabella 17*, Meta, Hugging Face e il cloud provider francese Scaleway hanno lanciato un acceleratore europeo per le startup di intelligenza artificiale open source, con l'obiettivo dichiarato di rafforzare l'ecosistema europeo e sostenere la diffusione di modelli aperti come la famiglia LLaMA. Questa partnership si distingue dalle altre per tre ragioni principali. In primo luogo, non è legata a investimenti diretti di capitale né a clausole di esclusiva, ma a un programma di supporto imprenditoriale che fornisce accesso gratuito o agevolato a risorse computazionali, mentoring e strumenti software. La seconda ragione ha una dimensione geopolitica, poiché mira a ridurre la dipendenza europea dagli *hyperscaler* americani. Infine, rappresenta un esempio concreto di come il paradigma open source possa essere utilizzato come leva competitiva.

Nel dicembre del 2023 IBM e Meta annunciano la nascita della AI Alliance, un'alleanza “pro-open” che punta a promuovere *open innovation* e *open science* nell'AI. Il progetto non include leader del settore come Microsoft, OpenAI o Google e si pone in contrapposizione rispetto agli attori citati per quanto riguarda la visione di apertura, la sicurezza e i profitti dell'intelligenza

artificiale. Tra i membri, oltre le due aziende, prendono parte all'iniziativa imprese come AMD che si occupa di sviluppo di GPU e CPU, Intel attiva nel campo dei microchip e dei semiconduttori, Dell nel settore dei computer e dei server, e infine Sony che si concentra sull'elettronica e videogiochi. Affianco a queste realtà si pongono altri attori come le università e le startup; in totale l'adesione iniziale viene stimata intorno a 45 organizzazioni.

L'idea, formando questa alleanza, è di spingere anche gli enti regolatori ad appoggiare, attraverso le norme, l'approccio aperto condiviso da queste aziende, le quali vedono l'AI come rappresentativa della conoscenza e della cultura umana e criticano i rischi di *regulatory capture* da parte del fronte chiuso.

La struttura operativa dell'AI Alliance vede dei gruppi di lavoro, un board per la governance e un organo di governance tecnico che si focalizzano su vari aspetti che riguardano le metriche di *trust & validation*, hardware e infrastrutture per il training, modelli e framework *open source*, standard e linee guida da condividere con iniziative governative e della società civile.

Nel 2024, dopo un anno dal lancio, si ha avuto una crescita fino a 140 membri e gli obiettivi che si sono posti per il 2025 erano due. Il primo di creare un Open Trusted Data Initiative (OTDI), che consiste nel poter dare accesso a dati aperti e *permissive-license* con la possibilità di avere la tracciabilità completa della vita di un dato attraverso strumenti e governance ad hoc. Il secondo un Trust & Safety Evaluation Initiative (TSEI) che permette di creare uno *stack* di valutazione aperto e di testare in maniera trasparente i modelli di intelligenza artificiale, in particolare i *foundation models* e i sistemi generativi. In questo modo si può valutare un modello in modo indipendente e trasparente, definire dei *benchmark* condivisi su standard aperti, sicurezza, robustezza, bias e performance. Inoltre, favorisce la riproducibilità, dando l'opportunità a chiunque di replicare i test e verificare i risultati. Lo *stack*, in particolare, include metriche standardizzate per misurare i rischi che possono essere allucinazioni, bias, contenuti tossici, dataset di valutazione con collezioni di set curati per categorie, *tool open source* con software per eseguire i test in maniera uniforme, e infine linee guida di reporting per modelli di documentazione pubblica.

Ad oggi OTDI e TSEI vanno dritti ai nodi oggi più critici, creando standard e *commons* aperti per poter definire provenienza dei dati, metriche di valutazioni riutilizzabili e prerequisiti per portare modelli open in produzione in modo affidabile. Sono considerate “infrastrutture soft” che mancavano e che possono ridurre asimmetrie rispetto alle piattaforme chiuse.

In aggiunta riescono a superare il problema di aziende come OpenAI o Google che usano spesso metriche proprietarie, con le quali è difficile confrontare modelli diversi. Altri due aspetti che

rendono questa iniziativa importante sono l'allineamento con gli enti regolatori, che chiedono maggiore trasparenza e responsabilità, la replicabilità e l'affidabilità delle valutazioni.

Per quanto riguarda i due grandi attori che hanno dato il via a questo progetto: IBM, già in passato creatore di prodotti open, come il sistema operativo aperto Linux, e volto all'*enterprise-grade*, ha l'occasione con l'AI Alliance di spingere strumenti e pratiche aperte e verificabili da usare nel mercato business, in concorrenza indiretta con offerte proprietarie più chiuse.

Meta, invece, utilizza l'alleanza per dare legittimazione e massa critica all'approccio LLaMA e, più in generale, ai modelli a pesi aperti. Sul piano reputazionale, la AI Alliance rende l'apertura una coalizione strategica e non solo una scelta isolata di Meta; dal punto di vista operativo, i dataset affidabili, gli *stack* di valutazione, gli hardware *enablement*, cioè utilizzabili e ottimizzati per l'AI, e la formazione creano beni pubblici complementari che abbassano i costi di adozione dei modelli aperti nei contesti industriali.

Emergono due criticità in questo piano: in primo luogo si nota la sovrapposizione con iniziative preesistenti, cioè alcune aziende aderenti hanno allo stesso tempo in essere accordi con altre imprese, come ad esempio Meta. La seconda criticità è rappresentata dalle assenze di *hyperscaler*, come Google Nvidia o Microsoft, che possono limitarne l'impatto nel breve termine. Resta quindi la domanda sul valore aggiunto concreto e sulla capacità di tradurre la retorica dell'apertura in adozioni industriali misurabili.

Tabella 17: Partnership Meta.

Anno	Partnership	Contributo delle parti	Rischio concorrenziale
2023	IBM – Meta → AI Alliance (45 membri diventati 140 nel 2024)	IBM: ecosistema open enterprise; Meta: LLaMA.	Rischio minore: modello aperto, ma frammentazione e assenza <i>hyperscaler</i> riducono l'impatto
2023–2024	Meta – Hugging Face – Scaleway (acceleratore europeo)	Meta: LLaMA open source; HF: hub modelli; Scaleway: cloud europeo	Basso rischio, orientato a pluralismo e riduzione dipendenza USA
2025	Meta – Google Cloud (\$10mld per 6 anni)	Meta: modelli LLaMA; Google: cloud TPU	Rafforza Google Cloud come provider critico, aumenta concentrazione compute
2025	AI Alliance (OTDI + TSEI)	Membri: IBM, Meta, AMD, Intel, Dell, Sony, università, startup. Sviluppano Open Trusted Data Initiative (dati open tracciabili) e Trust & Safety Evaluation Initiative (<i>stack</i> di valutazione aperto)	Rischio minore: crea <i>commons</i> aperti, ma se non adottati universalmente restano inefficaci contro i modelli chiusi

Nel 2024 Microsoft e Hugging Face, con il supporto di Meta, hanno rafforzato la loro collaborazione all'interno dell'ecosistema Azure AI Foundry, che rappresenta una piattaforma per sviluppare, testare e distribuire applicazioni basate sui *foundation models*.

L'obiettivo è aumentare la disponibilità di modelli open source, tra cui la famiglia LLaMA di Meta, direttamente all'interno della piattaforma cloud di Microsoft, rendendoli facilmente accessibili a sviluppatori e imprese. Attraverso questa intesa, Hugging Face riesce a consolidare sempre di più il proprio ruolo come punto di riferimento per la comunità *open source*, mentre Microsoft amplia la propria offerta sul *marketplace* di Azure, includendo oltre ai modelli proprietari, come GPT di OpenAI, anche modelli aperti e personalizzabili. La partnership, quindi, consente a Microsoft di posizionarsi come una piattaforma capace di integrare diversi approcci all'AI generativa, e in aggiunta si ha una diffusione maggiore di un paradigma alternativo rispetto al modello chiuso. In questa collaborazione Meta ha un duplice vantaggio:

favorisce la diffusione dei propri modelli LLaMA in contesti *enterprise* e accademici, e contribuisce a rafforzare la legittimazione dell'approccio *open source*.

Nel maggio 2024 OpenAI e Reddit hanno annunciato una partnership strategica che permette di utilizzare i contenuti pubblici della piattaforma social nell'addestramento e nel *fine-tuning* dei modelli linguistici di OpenAI. L'accordo prevede di poter usare i modelli addestrati con quei dati sia lato consumatori con ChatGPT, sia lato business tramite le API *enterprise* che OpenAI vende a banche, e-commerce e pubbliche amministrazioni.

Reddit è nato nel 2005 come piattaforma di social news e discussione organizzata in *subreddit* tematici, dove gli utenti possono condividere e commentare contenuti, formando delle vere e proprie community. Il social è utilizzato giornalmente da oltre 70 milioni di utenti e nel complesso si contano miliardi di conversazioni indicizzate, rappresentando una delle fonti più vaste e dinamiche di dati testuali su internet.

Questa partnership lato OpenAI consente di arricchire e aggiornare i propri modelli con dati nuovi e diversificati, aumentando capacità di generare risposte informate e contestuali. Dall'altra parte Reddit ha un vantaggio di tipo economico, poiché monetizza l'accesso al proprio archivio di conversazioni e al tempo stesso può integrare strumenti basati su ChatGPT all'interno della propria piattaforma.

Dal punto di vista competitivo, l'operazione ha ricadute significative. Innanzitutto rafforza la posizione di OpenAI nel mercato dei dati, assicurandole una risorsa unica e difficilmente replicabile; si osserva anche la riduzione del gap rispetto a Google, che già aveva siglato un accordo analogo con Reddit nei mesi precedenti per 60 milioni di dollari per addestrare Gemini. L'alleanza solleva tuttavia interrogativi regolatori poiché i contenuti su Reddit sono generati dagli utenti, e l'uso da parte di modelli di AI apre questioni legate a proprietà intellettuale, consenso e compensazione degli autori.

Quindi la partnership OpenAI–Reddit oltre ad essere un'intesa commerciale, segna un passo verso la costruzione di un ecosistema dati-centrico in cui le piattaforme social diventano fornitori strategici di contenuti per i modelli di frontiera.

Nel 2025 lo scenario si è ampliato con una serie di operazioni. Due aziende che fanno parte delle GAMMAN come Meta e Google, nonostante siano tradizionalmente rivali, hanno siglato un accordo da dieci miliardi di dollari per sei anni che prevede l'utilizzo del cloud di Google da parte di Meta per ospitare e addestrare i propri modelli. L'intento di Google è cercare di colmare il gap con AWS e Azure.

Sempre nello stesso anno ASML, il colosso olandese che detiene la produzione di macchine litografiche, ha deciso di investire 1,3 miliardi di euro in Mistral, inoltre in questa operazione

rientrano sia Nvidia sia il governo francese. Lo scopo della partnership è di natura industriale, pensata per rafforzare la sovranità tecnologica europea e ridurre la dipendenza dai player americani. Parallelamente ASML potrebbe utilizzare questi *foundation models* per migliorare i propri processi produttivi e rappresenta anche un'opportunità di spostarsi all'interno della catena del valore dell'AI, diventando un partner strategico di chi sviluppa i modelli. Facendo in modo di posizionarsi sia a livello hardware e software all'interno dell'*AI stack*.

Nel gennaio 2025 è stato promosso il progetto Stargate, che prevede una joint venture con l'obiettivo di costruire il supercomputer più potente al mondo dedicato all'AI di frontiera, con una capacità stimata di 4,5 GW, di molto superiore agli attuali data center. L'iniziativa nasce dalla partnership di quattro realtà: OpenAI, Oracle, SoftBank, MGX che puntano a garantire una capacità computazionale indipendente dai tradizionali *hyperscaler*. In questo contesto OpenAI fornisce il know-how scientifico e i modelli di AI come GPT, Oracle mette a disposizione la propria infrastruttura di cloud e data center, con competenze in *high-performance computing*, infine SoftBank e MGX fungono da investitori strategici che forniscono capitali e risorse per la governance del progetto. Questa joint venture permette di formare un'infrastruttura strategica a livello globale, garantisce risorse per l'AI di frontiera, dove modelli sempre più grandi richiedono capacità di calcolo enormi e superiori a quelle attuali. In aggiunta OpenAI ha la possibilità di non dipendere da un singolo provider, come Microsoft, ma di costruire un'infrastruttura in comune con altri attori.

Queste alleanze mostrano come il fenomeno non sia più limitato a partnership bilaterali startup–Big Tech, ma si allarghi a joint ventures multi-attore e a cross-over industriali con player di diverso tipo.

Nel giugno 2025 Reuters ha rivelato che OpenAI ha siglato un accordo senza precedenti con Google Cloud, nonostante la rivalità diretta tra le due aziende nel campo dell'AI generativa. L'intesa prevede che OpenAI utilizzi i data center e le TPU di Google per supportare l'addestramento e il funzionamento dei propri modelli di frontiera, accanto alle risorse già fornite da Microsoft tramite Azure. Dal punto di vista strategico segna una svolta perché il principale partner di OpenAI non è più esclusivamente Microsoft, ma entra in scena un altro *hyperscaler*. La scelta di questo accordo è dettata dalla sempre più crescente domanda di capacità computazionale, ed è la motivazione che ha portato l'azienda a decidere di diversificare i fornitori di cloud, come dimostra la *tabella 18*.

Questa alleanza strategica ha almeno tre implicazioni di rilievo. In primo luogo, segnala la centralità del calcolo ad alte prestazioni come risorsa critica, tanto da costringere OpenAI a collaborare con un competitor diretto pur di garantirsi accessibilità scalabile a GPU e TPU. In

secondo luogo, ridimensiona il rapporto privilegiato con Microsoft, aprendo scenari di tensione e potenziali rinegoziazioni contrattuali: Microsoft rimane partner strategico e investitore, ma non più unico canale per l'infrastruttura cloud. Infine, si consolida il ruolo di Google Cloud come attore imprescindibile nella fornitura di compute per l'AI, confermando come gli *hyperscaler* mantengano un potere strutturale anche nei confronti di sviluppatori di modelli di punta.

Nel settembre del 2025 OpenAI e Oracle hanno firmato un'intesa strategica stimata in circa 300 miliardi di dollari spalmati in cinque anni, finalizzata a garantire ulteriore capacità computazionale su larga scala per lo sviluppo dei modelli di frontiera. L'intesa prevede che OpenAI utilizzi i data center Oracle Cloud, per far fronte all'enorme fabbisogno di calcolo necessario ad addestrare i futuri sistemi di intelligenza artificiale generativa e, in prospettiva, l'Artificial General Intelligence (AGI), un'intelligenza artificiale capace di ragionare e agire in modo generale come o, meglio, di un essere umano.

La seguente partnership segna l'ingresso di Oracle nel ristretto club di fornitori cloud per l'AI generativa fino ad allora dominato da Microsoft, Amazon e Google. Inoltre, permette a OpenAI di diversificare strategicamente ancora di più le fonti di calcolo, riducendo la dipendenza da Azure e rafforzando la resilienza infrastrutturale. Un ultimo aspetto concerne l'emergere dell'AI come driver centrale di domanda per il livello della catena rappresentato dal cloud, trasformando i rapporti di forza tra *hyperscaler* e sviluppatori di modelli: non più solo il software che si appoggia al cloud, ma dei cloud interamente riconfigurati intorno all'AI.

Dal punto di vista competitivo, l'operazione si intreccia con il progetto Stargate, portando a una convergenza tra queste due iniziative: si rafforza il posizionamento di Oracle come partner privilegiato di OpenAI sul fronte infrastrutturale, offrendo una base di calcolo indipendente. In prospettiva, questo accordo potrebbe quindi riequilibrare la mappa delle alleanze nell'AI, con Oracle che si potrebbe ritagliare un ruolo di primo piano come fornitore alternativo di capacità computazionale di frontiera.

Tabella 18: Partnership OpenAI.

Anno	Partnership	Contributo delle parti	Rischio concorrenziale
2016-2023	Microsoft-OpenAI	Microsoft fornisce infrastrutture cloud iniziali; OpenAI modelli nell'ecosistema Microsoft	Lock-in infrastrutturale e blocco licenze che limita l'accesso di altri player
2024	Apple – OpenAI (ChatGPT su iOS/Siri)	Apple: distribuzione su iPhone; OpenAI: ChatGPT	Rischio di lock-in distributivo su centinaia di milioni di device
2024	OpenAI – Reddit	OpenAI: training/fine-tuning modelli; Reddit: dati social unici	Rischio: vantaggio esclusivo sui dati, barriere all'ingresso per potenziali competitor
2025	OpenAI – Google Cloud	OpenAI: modelli GPT; Google: data center e TPU	Rischio: OpenAI diversifica ma aumenta la dipendenza dagli <i>hyperscaler</i>
2025	OpenAI – Oracle (\$300 mld per 5 anni)	OpenAI: modelli di frontiera; Oracle: cloud e data center	Rischio: potere enorme di pochi attori, crea dipendenza multi-decennale
2025	OpenAI – Oracle – SoftBank – MGX → “Stargate”	OpenAI: modelli GPT; Oracle: cloud HPC; SoftBank e MGX: capitale	Rischio concentrazione di compute in una JV ristretta; riduce diversità di attori
2025	OpenAI – Nvidia (\$100 mld)	Nvidia: GPU Blackwell e DGX; OpenAI: FM training e inferenza	Rischio alto: Nvidia diventa fornitore indispensabile, aumentando il potere sulla supply chain

A inizio settembre del 2025 Microsoft ha iniziato a valutare un accordo con Anthropic, finalizzato a ridurre la propria dipendenza strategica da OpenAI, in risposta anche alle strategie di diversificazione dell'azienda partner. In questo modo Microsoft esplora modelli alternativi per differenziare a sua volta i fornitori di *foundation models* e mitiga il rischio di *lock-in*.

Amazon e Google hanno stretto precedentemente accordi con Anthropic, andando a formare una partnership incrociata tra un *hyperscaler* e una startup già legata ad altri giganti tecnologici. L'operazione mostra come le rivalità verticali passano in secondo piano, rispetto alla necessità di accesso ai modelli e risorse computazionali.

Allo stato attuale Microsoft ha annunciato l'integrazione del modello Claude di Anthropic all'interno della suite Microsoft 365, dando la possibilità agli utenti di Copilot di utilizzarlo accanto a GPT di OpenAI. In tale modo, le imprese clienti di Microsoft possono scegliere tra diversi modelli generativi, secondo logiche di *multi-model e best-of-breed*.

Questa evoluzione segnala che nel settore AI le partnership diventano non solo strumenti di efficienza, ma vere e proprie mosse geopolitiche industriali, capaci di ridefinire equilibri di potere e sollevare l'attenzione dei regolatori.

Si aggiunge alle alleanze instaurate finora da OpenAI quella con Nvidia, che prevede una partnership strategica da cento miliardi di dollari, destinata a realizzare una nuova generazione di data center su scala globale. L'accordo prevede la costruzione di infrastrutture avanzate, equipaggiate con la nuova generazione di GPU Nvidia Blackwell e con i sistemi DGX, supercomputer progettati per il training e l'inferenza di modelli di deep learning su larga scala, ottimizzati per l'addestramento e l'esecuzione dei *foundation models* di OpenAI.

Per OpenAI, la partnership rappresenta una garanzia di accesso prioritario alle risorse computazionali più avanzate sul mercato, in un contesto in cui la disponibilità di GPU è diventata il vero collo di bottiglia dello sviluppo AI. Per Nvidia, l'intesa consolida il proprio ruolo come fornitore infrastrutturale indispensabile, non solo vendendo chip, ma diventando partner diretto nella realizzazione di data center di nuova generazione.

Un altro sviluppo recente riguarda le trattative tra Apple e Google per l'integrazione del modello Gemini all'interno di Siri. L'accordo, ancora in fase esplorativa, prevede da un lato che i modelli proprietari di Apple continuino a gestire i dati personali e le funzionalità locali, preservando così l'attenzione al tema della privacy che caratterizza il brand; dall'altro, Gemini sarebbe impiegato per rispondere a domande complesse, creative o legate al web, ampliando notevolmente le capacità di Siri. Questa partnership segnerebbe un passaggio di rilievo, portando Apple ad accordarsi con un proprio competitor diretto, con una strategia ibrida tra la

chiusura tipica del proprio ecosistema e la necessità di adottare asset esterni per mantenere la competitività nell'ambito dell'intelligenza artificiale generativa.

Sempre nel campo degli smartphone, come è evidenziato nella *tabella 19*, invece, AI Perplexity, startup che ha come obiettivo costruire un motore di ricerca “*AI-first*”, basato principalmente su modelli linguistici di grandi dimensioni, ha intrapreso accordi con Motorola per quanto riguarda l'integrazione dei propri servizi all'interno dei loro device. Allo stesso modo Apple e Samsung stanno cercando di ottenere un'intesa per sfidare direttamente Google Search, proponendolo come motore di ricerca all'interno dei propri smartphone.

Le partnership mostrano una crescente complessità delle operazioni, si può notare come si è passati da semplici alleanze bilaterali con investimenti e accordi di natura infrastrutturale ad accordi multi-attore e joint ventures che prevedono governance condivisa e investimenti miliardari. Aumentando ancora di più l'importanza strategica di questo strumento.

Parallelamente, emerge con chiarezza un effetto di concentrazione del mercato: negli ultimi anni questi accordi hanno rafforzato il ruolo dei grandi *hyperscaler*, che grazie alle partnership hanno consolidato il controllo sulle risorse critiche. La situazione in essere porta le startup indipendenti ad avere sempre più difficoltà nel competere senza legarsi a uno di questi attori dominanti, generando una dinamica di dipendenza strutturale che riduce la varietà dell'ecosistema.

Un esempio emblematico di questa ambivalenza è rappresentato dal progetto Stargate (OpenAI–Oracle–SoftBank–MGX). Da un lato, Stargate può essere interpretato come un fattore di concorrenza, poiché offre capacità di calcolo alternative e potenzialmente indipendenti rispetto al dominio tradizionale degli *hyperscaler*. Dall'altro lato, lo stesso progetto contribuisce alla concentrazione, poiché costituisce un blocco di risorse ingente nelle mani di pochi attori con potere finanziario e tecnologico enorme, escludendo di fatto startup e concorrenti minori dall'accesso a tali infrastrutture.

Tabella 19: Partnership ASML -Mistral e Perplexity – Motorola.

Anno	Partnership	Contributo delle parti	Rischio concorrenziale
2025	ASML – Mistral (con partecipazione di Nvidia e Governo francese) (€1,3 mld)	ASML: capitale e know-how industriale; Mistral: modelli; Nvidia: GPU; FR: supporto politico	Rischio basso-medio: alleanza industriale geopolitica EU, limita però il pluralismo delle startup
2025	Perplexity – Motorola	Perplexity: motore <i>search AI-first</i> ; Motorola: integrazione device	Rischio basso: nuova entrata, ma legata a device specifici

5.4 Partnership cross-industry

Le partnership *cross-industry* a differenza di quelle verticali e orizzontali non si sviluppano all'interno della catena del valore dell'AI, ma collegano le Big Tech e i fornitori di modelli a settori industriali tradizionali come sanità, automotive, telecomunicazioni, ricerca e nei servizi professionali. Queste alleanze hanno l'obiettivo di diffondere l'AI generativa in applicazioni concrete, creando tecnologie complementari e nuovi mercati, ma sollevano anche interrogativi legati a *lock-in*, standard settoriali e governance dei dati. Le partnership verranno riassunte alla fine del paragrafo nella *tabella 20*.

5.4.1 Sanità

Il primo settore che verrà analizzato è quello della sanità, in cui l'intelligenza artificiale generativa e i *foundation models* trovano diverse applicazioni. Questi accordi tra le Big Tech e gli attori sanitari mirano a combinare i dati clinici e le competenze provenienti dalle strutture pubbliche o private con la capacità di calcolo e i modelli di linguaggio avanzati.

Uno degli accordi più significativi riguarda la collaborazione tra Google DeepMind e il National Health Service (NHS) britannico, iniziata nel 2016, che nel corso degli anni si è evoluta in progetti sull'AI nel campo della diagnosi oculistica e nell'analisi delle immagini mediche. Durante questa partnership DeepMind ha utilizzato i dataset pubblici forniti dal NHS per sviluppare dei sistemi capaci di individuare patologie oculari a partire dalle scansioni retiniche, raggiungendo livelli di accuratezza paragonabili a quelli dei medici specialisti.

Tuttavia, l'utilità dei dati in queste applicazioni si scontra con i problemi di privacy e di gestione dei dati sensibili.

Una seconda alleanza è quella tra Microsoft e Nuance Communications, azienda leader nelle tecnologie di riconoscimento vocale nel campo sanitario. L'accordo sottoscritto nel 2019 inizialmente prevedeva lo sfruttamento del cloud Azure di Microsoft per sviluppare delle soluzioni migliorative per il riconoscimento vocale con l'intelligenza artificiale. In seguito, nel 2021, è avvenuta l'acquisizione da parte di Microsoft dell'azienda per più di 16 miliardi di dollari, dalla quale è nato DAX Copilot: un assistente basato su GPT che trascrive automaticamente le conversazioni tra medico e paziente, alleggerendo il carico burocratico e migliorando la qualità delle cure. In questo caso la partnership *cross-industry* si è trasformata in un'integrazione verticale tra l'infrastruttura cloud, Azure, i modelli generativi, OpenAI, e le applicazioni settoriali sviluppate da Nuance.

Allo stesso modo Nvidia ha stretto accordi in ambito sanitario con GE Healthcare, per trovare delle soluzioni di intelligenza artificiale che si potessero applicare ai sistemi radiologici ed ecografici. La collaborazione si pone l'obiettivo di riuscire a sviluppare strumenti diagnostici per l'elaborazione in tempo reale delle immagini radiologiche ed ecografiche. In modo da aumentare la rapidità, l'efficienza e la precisione nella rilevazione precoce delle patologie, riuscendo, così, a supportare i professionisti sanitari in contesti caratterizzati da carenza di personale e dall'aumento del carico di lavoro.

La partnership prevede l'integrazione delle GPU e delle librerie software, tra cui CUDA, di Nvidia all'interno delle piattaforme medicali di GE Healthcare. Il vantaggio strategico di questa scelta risiede nella complementarità delle aziende. Infatti, GE Healthcare fornisce l'esperienza clinica, la rete di distribuzione ospedaliera e l'accesso a grandi dataset sanitari, mentre Nvidia contribuisce con la potenza computazionale e i software.

Il rischio principale di questa iniziativa è la presenza del leader di settore dei chip che potrebbe creare una dipendenza da standard proprietari, limitando la possibilità di interoperabilità con altre piattaforme. Inoltre, l'uso di modelli AI in ambito clinico pone problemi cruciali sul piano della responsabilità medica e della trasparenza algoritmica, aspetti che richiedono una governance rigorosa per garantire la fiducia degli operatori sanitari e dei pazienti.

Negli esempi considerati si nota come le Big Tech mettano a disposizione risorse computazionali, cloud e modelli generativi, mentre ospedali, università e imprese sanitarie forniscano i dataset clinici, le competenze verticali e l'accesso al mercato. I benefici includono la riduzione dei costi di ricerca, il miglioramento della diagnostica e l'efficienza dei processi clinici. Ciononostante, la dipendenza da piattaforme proprietarie potrebbe portare a dei *lock-in*

anche nel settore sanitario, mentre l'uso di dati sensibili necessita di standard rigorosi di governance per evitare abusi e violazioni della privacy.

5.4.2 Automotive

Un altro settore caratterizzato dall'utilizzo di AI è quello dell'automotive, dove si hanno diverse applicazioni. In particolare, per la guida autonoma, i sistemi di infotainment e l'ottimizzazione della catena del valore. Le partnership *cross-industry* tra Big Tech e case automobilistiche mostrano come l'intelligenza artificiale generativa stia trasformando non solo il prodotto finale, l'automobile, ma anche i processi di ricerca e sviluppo.

Nel 2020 Nvidia e Mercedes-Benz hanno avviato un accordo, che gli ha permesso di integrare nei propri veicoli la piattaforma di calcolo Nvidia Drive, una suite hardware-software che integra GPU, sistemi operativi e strumenti di *deep learning* progettati per la guida autonoma e l'elaborazione in tempo reale di grandi volumi di dati provenienti da sensori e telecamere.

Lo scopo è di sviluppare sistemi di guida autonoma di livello avanzato, strumenti predittivi per la manutenzione, oltre all'uso dei modelli generativi, per gli assistenti vocali personalizzati e le interfacce uomo-macchina.

Questa alleanza consente all'azienda automobilistica di accedere a tecnologie di frontiera senza dover sviluppare in autonomia le costose infrastrutture hardware e software, mentre Nvidia rafforza la propria posizione come fornitore chiave di semiconduttori e piattaforme AI per il settore automotive, e raccoglie allo stesso tempo enormi dataset di guida reali.

Analogamente all'ambito sanitario la presenza dell'azienda di chip potrebbe causare un *lock-in* tecnologico a Mercedes, non potendo integrare in futuro soluzioni alternative

La logica delle partnership *cross-industry* è evidente: le aziende automobilistiche offrono dati di guida, competenze meccaniche e accesso al mercato; le Big Tech e i fornitori AI mettono a disposizione la capacità computazionale, gli algoritmi e i modelli generativi. I benefici che si percepiscono riguardano la velocità di innovazione, la personalizzazione dell'esperienza utente e l'avanzamento della guida autonoma. I rischi, invece, includono il potenziale *lock-in* e la creazione di standard proprietari che potrebbero frammentare l'ecosistema globale della mobilità intelligente.

5.4.3 Telecomunicazioni

Nel campo delle telecomunicazioni le partnership che riguardano l'AI generativa assumono maggior rilievo per l'integrazione con le reti 5G e l'*edge computing*. Gli operatori telco vedono nell'AI un'opportunità per differenziare l'offerta e aumentare il valore dei servizi, mentre le Big Tech ottengono accesso a infrastrutture di rete capillari e a basi di utenti globali.

Il caso più importante riguarda SK Telecom (STK), principale operatore mobile sudcoreano, che dal 2023 ha stretto una serie di accordi strategici con Anthropic e AWS per sviluppare soluzioni AI verticalizzate. I termini prevedono l'utilizzo da parte del modello Claude di Anthropic, per adattarlo alle esigenze specifiche delle telecomunicazioni, il quale prende il nome di Telco Claude. Il modello consente un'evoluzione del cloud computing per le reti di telecomunicazioni, virtualizzando le funzioni di rete e rendendole disponibili come software su infrastrutture standard e flessibili, invece che su un hardware proprietario.

L'obiettivo di questa integrazione con i sistemi di SKT è il potenziamento del customer care, a cui seguirà un miglioramento sia della qualità dei servizi digitali sia dell'ottimizzazione della gestione predittiva delle reti. Invece il collegamento con Amazon si realizza, poiché Anthropic stessa ha accordi in essere con Amazon, i quali permettono che anche questo modello venga distribuito su larga scala tramite la piattaforma Amazon Bedrock consentendo di addestrare e personalizzare modelli generativi con dati settoriali senza doverli ricostruire da zero.

Dal punto di vista strategico, la partnership riflette un modello di complementarità degli asset: SKT mette a disposizione dati di rete e accesso al mercato, Anthropic fornisce i *foundation models*, mentre AWS garantisce la potenza computazionale e l'infrastruttura cloud.

5.4.4 Ricerca scientifica e education

Nella ricerca scientifica le collaborazioni, a differenza di altri ambiti, non hanno solo una valenza commerciale, ma puntano a rispondere all'esigenza di sviluppare innovazione per la comunità, attraverso dataset, strumenti e modelli aperti che possano accelerare la produzione di conoscenza.

Nuovamente Nvidia in questo ambito ha instaurato delle partnership, la più significativa è l'accordo con il California Institute of Technology (Caltech), e che più in generale è collegata a iniziative promosse dallo Stato della California per diffondere risorse di intelligenza artificiale nei college e nei centri di ricerca.

Nel 2024 la California ha lanciato, insieme a Nvidia, una collaborazione definita come la prima nel suo genere, che si pone l'obiettivo di mettere a disposizione delle università pubbliche e private la potenza computazionale, le infrastrutture AI e dei programmi formativi basati sulle piattaforme Nvidia. La partnership include la fornitura di GPU Nvidia di ultima generazione e i software come CUDA e Omniverse. Quest'ultima è una piattaforma collaborativa che in tempo reale consente ai creatori di progettare, costruire e collaborare in ambienti 3D e metaversi.

Parallelamente, Nvidia collabora direttamente con istituzioni come Caltech e laboratori indipendenti per applicare l'AI generativa a sfide scientifiche di frontiera, dalla modellazione

proteica alle simulazioni molecolari. Queste partnership non hanno solo un obiettivo tecnologico, ma intendono creare un ecosistema accademico-industriale che rafforzi il ruolo della California come hub globale per l'AI di frontiera. Il rischio è sempre quello di dipendere da un unico standard proprietario, diventando difficile la ricerca di fornitori alternativi di tecnologia.

Un'altra compagnia che ha esteso la propria strategia in ambito AI è IBM. L'azienda si è concentrata sull'*education* e sulla ricerca scientifica attraverso lo sviluppo della piattaforma Watsonx che integra *foundation models*, strumenti di gestione dei dati e soluzioni di AI *enterprise*. La piattaforma collabora con le università, gli enti di ricerca e i partner tecnologici per offrire dataset certificati, tool di valutazione e infrastrutture cloud accessibili anche in contesti non industriali. Un esempio concreto riguarda la collaborazione con Academic Software, che mira a mettere Watsonx a disposizione di studenti e ricercatori in Europa. La partnership permette alle università di colmare il divario tra accademia e industria attraverso nuove competenze supportate dall'AI di livello *enterprise*. La dipendenza da queste piattaforme porta sempre con sé l'ombra del *lock-in* tecnologico e la riduzione dell'autonomia accademica oltre che lo spostamento del baricentro dalla scienza verso logiche di mercato, con il vantaggio però di creare la formazione di competenze nel panorama dell'AI.

5.4.5 Servizi professionali

Nei servizi professionali le partnership *cross-industry* permettono di trovare soluzioni che uniscano le capacità analitiche avanzate con l'automazione documentale e l'utilizzo di strumenti di compliance normativa. Gli accordi in questo ambito, solitamente, vengono stipulati con startup AI che integrano i propri *foundation models* nei processi quotidiani che sono altamente *data-driven*.

Un esempio di questa dinamica è l'alleanza tra PwC, Harvey AI e Lefebvre. Harvey AI è una startup fondata da ex avvocati e ricercatori di Stanford, specializzata nello sviluppo di *foundation models* adattati al diritto e alla consulenza legale. In questo accordo svolge il ruolo di integrazione dei propri modelli all'interno della divisione Legal Business Solutions di PwC, posizionando la società di consulenza tra i pionieri nell'uso dell'AI generativa per i servizi legali. Per integrare il tutto è stato necessario sviluppare una piattaforma che consentisse di automatizzare compiti complessi di analisi contrattuale, ricerca giurisprudenziale e redazione di pareri legali, che ha portato al beneficio di ridurre drasticamente i tempi e i costi delle attività ripetitive. Invece la partnership con Lefebvre, un editore giuridico europeo, ha dato la

possibilità di integrare banche dati normative e contenuti legali certificati, fondamentali per addestrare e validare i modelli utilizzati.

In sintesi, Harvey fornisce il modello generativo adattato al settore legale, Lefebvre contribuisce con contenuti giuridici proprietari e know-how editoriale, e PWC porta la propria base clienti globale e la capacità di consulenza integrata. Oltre ai benefici già descritti questi accordi e la piattaforma portano a dei rischi che riguardano la responsabilità professionale e la trasparenza delle decisioni algoritmiche.

Tabella 20: Partnership cross-industry.

Settore	Attori principali	Contributi	Benefici	Rischi
Sanità	Google DeepMind – NHS (UK)	NHS: dataset clinici; DeepMind: modelli diagnostici	Miglioramento diagnostica oculistica, riduzione tempi diagnosi	Problemi privacy e gestione dati sensibili
	Microsoft – Nuance (2019 → acquisizione 2021)	Microsoft: cloud Azure e GPT OpenAI; Nuance: competenze sanitarie vocali	DAX Copilot: trascrizione automatica colloqui medico-paziente, riduzione burocrazia	<i>Lock-in</i> infrastrutturale e uso dati sensibili
	Nvidia – GE Healthcare	GE: know-how clinico, rete ospedali; Nvidia: GPU, CUDA	Diagnostica radiologica ed ecografica più veloce ed efficiente	Dipendenza da standard proprietari Nvidia, responsabilità medica e trasparenza algoritmi

Automotive	Nvidia – Mercedes-Benz (2020)	Mercedes: dati di guida e mercato; Nvidia: piattaforma Drive (GPU, DL, OS)	Sviluppo guida autonoma, manutenzione predittiva, assistenti vocali	<i>Lock-in</i> tecnologico, difficoltà interoperabilità
Telecomunicazioni	SK Telecom – Anthropic – AWS (2023)	SKT: dati rete e utenti; Anthropic: modello Claude adattato (Telco Claude); AWS: cloud/compute	Customer care avanzato, reti virtualizzate, gestione predittiva	Dipendenza da <i>hyperscaler</i> , <i>lock-in</i> infrastrutturale
Ricerca ed education	Stato California – Nvidia – Caltech (2024)	Nvidia: GPU, CUDA, Omniverse; università: ricerca scientifica	Potenziamento calcolo per università, applicazioni scientifiche di frontiera	Dipendenza da tecnologia Nvidia, <i>lock-in</i> standard
	IBM – Watsonx e Academic Software	IBM: <i>foundation models enterprise</i> ; università: competenze e dati	Accesso a dataset certificati e tool AI per studenti/ricercatori	Rischio <i>lock-in</i> accademico, spostamento verso logiche di mercato
Servizi professionali	PwC – Harvey AI – Lefebvre	PwC: clientela e consulenza; Harvey: AI	Automazione analisi contratti, ricerche	Trasparenza algoritmica,

		legale; Lefebvre: banche dati giuridiche	giuridiche, riduzione costi	responsabilità professionale
--	--	---	--------------------------------	---------------------------------

5.5 Benefici e rischi delle partnership

Alla luce dei vari accordi riportati nei paragrafi 5.3 e 5.4 si nota come la collaborazione non si esaurisca in ambito tecnico, ma influenzi in modo determinante la competizione e l'innovazione nel settore, portando a dei benefici e dei rischi sistematici.

Per quanto riguarda i benefici si hanno tre effetti principali: innovazione, sfruttamento della complementarità degli asset e diffusione dell'AI in settori verticali. Il primo consiste nell'accelerazione dell'innovazione attraverso partnership con startup come OpenAI, Anthropic o Mistral. Infatti in questo modo hanno avuto l'opportunità di accedere a risorse computazionali e finanziarie che altrimenti sarebbero rimaste fuori portata, che hanno consentito di ridurre i tempi necessari per addestrare modelli di frontiera. Il secondo effetto mostra come le Big Tech offrano le infrastrutture cloud, hardware e di distribuzione, mentre le startup portino creatività, velocità di sviluppo e *know-how* specializzato. L'ultimo riguarda la possibilità di diffondere i modelli attraverso partnership *cross-industry*. Un esempio emblematico è l'integrazione di Claude in ambito telco, grazie agli accordi tra Anthropic, AWS e operatori come SK Telecom che permettono di declinare un modello generativo generalista in soluzioni verticali.

Sul versante opposto i rischi sono molteplici. In primo luogo, le partnership creano spesso *lock-in* e dipendenza strutturale, poiché l'uso di servizi cloud rende difficile per le startup svincolarsi dai provider, e questo genera delle barriere all'ingresso. Inoltre, è presente il rischio di *foreclosure*, per cui le startup indipendenti che non stringono alleanze con gli *hyperscaler* faticano ad accedere a risorse computazionali e mercati, rischiando di essere escluse dalla competizione.

In aggiunta la FTC nel 2025 ha messo in luce alcune clausole nascoste o poco trasparenti inserite negli accordi tra *hyperscaler* e sviluppatori di *foundation models*, che non emergono nei comunicati ufficiali, ma che condizionano in profondità la struttura del mercato.

Molti contratti prevedono una combinazione di equity e revenue sharing, in base alla quale i fornitori di cloud ottengono quote di capitale o diritti sui ricavi futuri, talvolta con opzioni che lasciano aperta la possibilità di acquisizioni complete.

Altri accordi hanno al loro interno diritti di consultazione e di influenza strategica, per cui è necessaria la presenza di osservatori nei board, obblighi di preavviso sulle decisioni rilevanti e, in alcuni casi, clausole di esclusiva o trattamenti preferenziali, come nel caso della partnership Microsoft–OpenAI: dal 2019 ha vincolato l’uso esclusivo del cloud Azure e ha garantito a Microsoft licenze privilegiate sui modelli di OpenAI.

Un ulteriore elemento ricorrente è la cosiddetta *circular spending* per cui le Big Tech investono miliardi in una startup, mostrando formalmente l’intenzione a compiere operazioni di finanziamento a sostegno all’innovazione, ma nella pratica lo stesso sviluppatore di modelli è obbligato a spendere miliardi di dollari in servizi cloud presso il partner. Questa clausola è chiamata *cloud commitment* e genera un meccanismo nel quale il provider rientra nei soldi spesi con il consumo vincolato delle proprie infrastrutture, riducendo al minimo il rischio industriale e finanziario. È il caso, ad esempio, di Anthropic, che con Amazon si è impegnata ad addestrare i propri modelli esclusivamente su AWS, utilizzando la piattaforma Bedrock per la distribuzione e integrando i chip Trainium e Inferentia.

Parallelamente, i contratti riconoscono ai fornitori di infrastrutture un accesso privilegiato a dati sensibili e proprietà intellettuale: non solo informazioni sulle performance e sui ricavi dei modelli, ma anche metodologie di *training*, output per test interni e persino la possibilità di co-sviluppare componenti hardware ottimizzati per i modelli del partner. In questa casistica rientra l’accordo tra Google e Anthropic. Infatti, una clausola inserita nel contratto prevede l’uso obbligatorio delle TPU di Google e la distribuzione dei modelli Claude tramite Vertex AI, oltre all’investimento in *equity non-voting*, assicurando a Google un canale privilegiato per consolidare il proprio ecosistema.

Un altro strumento di trasferimento tecnologico è lo scambio di personale tecnico, con ingegneri cosiddetti *embedded*: gli ingegneri Microsoft, Amazon o Google vengono “inseriti” nei team della startup per aiutare ad ottimizzare i modelli per funzionare al meglio con le infrastrutture proprietarie. In tale modo gli *hyperscaler* possono acquisire know-how e proprietà intellettuale.

Infine, alcuni accordi prevedono restrizioni all’autonomia commerciale degli sviluppatori: in determinati casi i modelli non possono essere lanciati o resi disponibili tramite API autonome prima di essere integrati e distribuiti attraverso il cloud del partner, rafforzando così dinamiche di *lock-in*.

Nel complesso, queste clausole dimostrano come le partnership rafforzino il potere strutturale degli *hyperscaler* senza la necessità di arrivare formalmente ad una fusione, ma utilizzando accordi che presentano al proprio interno accessi privilegiati. Si garantiscono così un controllo

diretto sugli input critici e limitano la contestabilità del mercato, con effetti che sollevano forti interrogativi dal punto di vista della concorrenza e della regolazione.

5.6 Prospettiva regolatoria e antitrust

Le partnership strategiche nel settore dell'AI generativa sono sempre più oggetto di attenzione da parte delle autorità antitrust, che vengono percepite come un potenziale strumento di concentrazione di mercato.

Per esempio, la CMA ha condotto, tra il 2023 e il 2025, come anticipati nei paragrafi precedenti, analisi approfondite in particolare su due casi: Microsoft–OpenAI e Google–Anthropic. Se si confrontano i due accordi emerge, grazie agli studi del CMA, che nel primo caso si ha un'influenza tale sull'altra azienda da costituire una concentrazione di fatto anche se formalmente non si ha una fusione: pur in assenza di acquisizione formale, l'accesso a *board rights* e clausole contrattuali conferisce a Microsoft un'influenza sostanziale sulle decisioni di OpenAI. Per quanto riguarda il secondo caso, anche se l'impatto concorrenziale è significativo, è stato classificato, sempre dal CMA, come una partnership non assimilabile ad una fusione.

A livello internazionale si hanno tendenzialmente due linee diverse di gestione di questa categoria di accordi. Negli Stati Uniti, attraverso la FTC, si privilegia la trasparenza delle clausole presenti nelle partnership. Infatti, l'obiettivo dei loro studi era mettere in luce dinamiche nascoste: *equity non-voting*, *consultation rights*, accesso preferenziale ai dati e *revenue sharing*.

Invece la Commissione Europea mira a evitare che queste partnership si tramutino in un ulteriore rischio sul controllo di input critici, quali i chip, i dati e il cloud, diventando la leva principale per consolidare il potere di mercato. In questo modo si concentra maggiormente su come questi accordi modifichino la contestabilità del mercato con l'obiettivo di raggiungere un accesso equo alle risorse.

6 Conclusioni

Nell'analisi condotta nel seguente elaborato emerge la portata del fenomeno dell'intelligenza artificiale, che si innesta come una tecnologica capace di impattare in modo significativo nel contesto economico e sociale. Questo è possibile essendo una tecnologia ad uso generale e alla sua natura LLM che permette una sua applicazione in svariati contesti, come dimostrano le partnership *cross-industry*.

Dal punto di vista degli effetti economici e sociali l'adozione dell'AI si può concludere che incide in modo significativo sulla produttività, sull'organizzazione del lavoro sulla formazione e la struttura dei mercati. Nonostante i benefici potenziali, la piena realizzazione dei guadagni

di efficienza dipende dalla capacità dei sistemi economici di investire in complementi intangibili come capitale umano, dati di qualità, infrastrutture digitali e fiducia istituzionale. L'AI può ampliare la produttività del lavoro e favorire l'innovazione, ma comporta anche la necessità di ripensare modelli di competenza e sistemi educativi. Come rilevato dall'AI Index Report (2024) e dall'OECD (2024), l'effetto netto sull'occupazione non sarà determinato dalla sostituzione uomo-macchina, bensì dal grado di complementarità tra abilità umane e capacità computazionali, e dalla rapidità con cui individui e imprese sapranno adattarsi alla trasformazione in corso.

All'interno dell'analisi si può notare che la formazione e le competenze saranno un punto cruciale per consentire a lavoratori e cittadini di interagire consapevolmente con sistemi intelligenti; poiché la diffusione dell'AI richiede non solo professionisti altamente specializzati, ma anche una diffusa *AI literacy*.

Dal punto di vista della catena del valore dell'AI il lavoro svolto permette di comprendere come i rapporti di forza tra le cosiddette GAMMAN e i nuovi entranti si stanno configurando sempre più sbilanciati verso i primi, che sfruttano le proprie conoscenze e la propria capacità di essere imprescindibili all'interno della catena, costituendo un collo di bottiglia competitivo, capace di influenzare l'intero ecosistema. Emerge così una tendenza da parte di chi predomina l'*AI stack* di allargare sempre di più la propria sfera di influenza andando o a sviluppare proprie soluzioni, come il caso dei FM, o servendosi della propria posizione di potere per stringere accordi con altri attori a monte o a valle della catena.

In questo contesto si può osservare che i dati e il modo di allenare i modelli avranno un'importanza che crescerà negli anni a venire, non a caso leader del settore come OpenAI hanno stretto accordi per attingere a dataset unici come Reddit.

Inoltre, dallo studio emerge che, nonostante ci siano progetti, come Stargate, che si pongono l'obiettivo di svincolarsi dalle aziende GAMMAN, sia imprescindibile allearsi con almeno una di queste per poter avere la possibilità di emergere nel contesto descritto.

In aggiunta appare evidente la dinamica per cui i leader in un determinato stadio della catena si alleino con altri leader in un altro, o nello stesso segmento, per escludere eventuali competitor emergenti e affermare la propria posizione dominante nel settore: lo strumento più utilizzato sono le partnership.

Le quali come indicato nell'elaborato si configurano per gli *hyperscaler* come un elemento di controllo su realtà emergenti nel medesimo settore, ne sono un esempio Microsoft-OpenAI,

Google-Anthropic e Amazon-Hugging Face che mostrano come la linea di confine tra cooperazione e controllo sia sempre più sottile.

Le autorità di regolazione, come la Competition and Markets Authority (2023-2024) e la Federal Trade Commission (2025), hanno rilevato come questa concentrazione possa compromettere la contendibilità dei mercati e richieda nuove forme di vigilanza per preservare la concorrenza lungo l'intera filiera.

Un'altra motivazione che spinge gli *hyperscler* ad allearsi con altre aziende è espandere la propria influenza in altre aree al di fuori del proprio ambito, con accordi *cross-industry* nel campo della sanità, della telecomunicazione e dell'automotive: da un lato accelerano l'innovazione, dall'altro creano dipendenze tecnologiche e strutturali da parte delle Big Tech.

Gli organismi di regolazione, come la CMA e la Commissione Europea, hanno avviato un monitoraggio costante di tali relazioni, introducendo principi volti a garantire trasparenza, interoperabilità e concorrenza effettiva.

Nel complesso, la tesi mostra come l'intelligenza artificiale non sia soltanto una tecnologia, ma un nuovo fattore di organizzazione economica e industriale. La sua evoluzione dipenderà dalla capacità di bilanciare tre elementi fondamentali: l'innovazione, per stimolare lo sviluppo scientifico e applicativo; la concorrenza, per evitare concentrazioni eccessive di potere economico; la regolazione, per garantire un utilizzo etico e inclusivo delle tecnologie.

La sfida, quindi, non consiste più nel domandarsi se l'intelligenza artificiale trasformerà l'economia, ma in quale direzione e con quali regole avverrà questa trasformazione.

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.

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

Badr, Mariam. *Unleashing the Power of AI: Exploring the Transformative Potential of Artificial Intelligence*. University of Massachusetts Global – School of Business, 2023.

Balland, Pierre-Alexandre, et al. *Generative AI and Foundation Models in the EU: Uptake, Opportunities, Challenges, and a Way Forward*. European Economic and Social Committee (EESC), 2025.

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.

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. *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. *Machine Learning, Labor Demand, and the Reorganization of Work*. MIT, 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.

CMA. *Full Text Decision in the Matter of Amazon / Anthropic Partnership*. Competition and Markets Authority (UK), 2025.

CMA. *Full Text Decision in the Matter of Google / Anthropic Partnership*. Competition and Markets Authority (UK), 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.

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.

Hagiu, 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.

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.

Team AWS. *Unlocking Value with Foundation Models – AWS Partner Series: Generative AI in Telco*. Amazon Web Services, Inc. (AWS), 2024.

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.

van der Vlist, Fernando, et al. *Big AI: Cloud Infrastructure Dependence and the Industrialisation of Artificial Intelligence*. 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

“Accordo OpenAI-Oracle, 300 Miliardi Di Investimenti in Potenza Di Calcolo in 5 Anni.” *Il Sole 24 ORE*, 11 Sept. 2025, www.ilsole24ore.com/art/accordo-openai-oracle-300-miliardi-investimenti-potenza-calcolo-5-anni-AHwb9RZC

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

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

“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.

“Annuncio Del Progetto Stargate.” *Openai.com*, 2025, <https://openai.com/it-IT/index/announcing-the-stargate-project/>

“Anthropic Partners with Google Cloud.” *Anthropic.com*, 2023, www.anthropic.com/news/anthropic-partners-with-google-cloud

“Asml Investe 1,3 Miliardi in Mistral Ai, Leader Europeo Nell’intelligenza Artificiale.” *Il Sole 24 ORE*, 9 Sept. 2025, www.ilsole24ore.com/art/ai-l-olandese-asml-investe-13-miliardi-francese-mistral-AHy0ApWC?refresh_ce=1

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

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

“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.

“Customer Story | SK Telecom | Claude.” *Claude.com*, 2025, <https://claude.com/customers/skt>.

“Da Rivali a Partner: Meta E Google Siglano Un Mega Accordo Da 10 Miliardi Sul Cloud.” *CorCom*, 22 Aug. 2025, www.corrierecomunicazioni.it/digital-economy/da-rivali-a-partner-meta-e-google-siglano-un-mega-accordo-da-10-miliardi-sul-cloud/.

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

“Entos (Now Iambic) in Collaboration with Caltech and NVIDIA: Advances in Geometric Deep Learning for Computational Chemistry in Proceedings of the National Academy of Science USA.” *Iambic.ai*, 2022, www.iambic.ai/post/pnas.

“FTC Issues Staff Report on AI Partnerships & Investments Study.” *Federal Trade Commission*, 17 Jan. 2025, www.ftc.gov/news-events/news/press-releases/2025/01/ftc-issues-staff-report-ai-partnerships-investments-study.

“GE HealthCare E NVIDIA Collaborano Allo Sviluppo Di Soluzioni Autonome per Sistemi Radiologici Ed Ecografici.” *Gehealthcare.it*, 2025, www.gehealthcare.it/insights/article/ge-healthcare-e-nvidia-collaborano-allo-sviluppo-di-soluzioni-autonome-per-sistemi-radiologici-ed-ecografici.

“IBM Watsonx Technology Partners.” *Ibm.com*, 2021, www.ibm.com/products/watsonx/partners.

“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/.

“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>.

“Meta, Accordo Cloud Da Oltre 10 Miliardi Di Dollari Con Google.” *Il Sole 24 ORE*, 22 Aug. 2025, www.ilsole24ore.com/art/meta-accordo-cloud-oltre-10-miliardi-dollari-google-AHH87DHC?refresh_ce=1.

“NVIDIA E Mercedes Benz Automotive Partner.” *NVIDIA*, 2022, www.nvidia.com/it/solutions/autonomous-vehicles/partners/mercedes/.

“OpenAI E NVIDIA Annunciano Una Partnership Strategica per l’Implementazione Di 10 Gigawatt Di Sistemi NVIDIA.” *Openai.com*, 6 Oct. 2025, <https://openai.com/it-IT/index/openai-nvidia-systems-partnership/>.

“OpenAI Sta Rinegoziando l’Accordo Con Microsoft in Vista Di Una Futura Ipo.” *Il Sole 24 ORE*, 11 May 2025, www.ilsole24ore.com/art/openai-sta-rinegoziando-l-accordo-microsoft-vista-una-futura-ipo-AHJ0Csh.

“Partnership with SK Marks a Turning Point for AI Innovation in Korea” – Prasad Kalyanaraman, vp of AWS Infrastructure Services | SK Telecom Newsroom.” *SK Telecom Newsroom*, 2025, <https://news.sktelecom.com/en/2074>

“Perché Apple E Meta Vogliono Perplexity, l’AI Nel Mirino Dei Big.” *Fastweb Plus*, 2025, www.fastweb.it/fastweb-plus/intelligenza-artificiale/perche-apple-e-meta-vogliono-perplexity-lai-nel-mirino-dei-big/.

“Powering the next Generation of AI Development with AWS.” *Anthropic.com*, 2024, www.anthropic.com/news/anthropic-amazon-trainium.

“PwC UK | Harvey.” *Harvey*, 2025, www.harvey.ai/customers/pwc-uk.

“Reddit Si è Accordato Con OpenAI per Permetterle Di Usare I Suoi Contenuti per Allenare I Modelli Di Intelligenza Artificiale.” *Il Post*, 17 May 2024, www.ilpost.it/2024/05/17/reddit-accordo-openai/.

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

“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.

“SK Group and AWS Team up to Build Cloud Computing Infrastructure to Support AI Innovation | SK.” *SK*, 2025, <https://eng.sk.com/news/sk-group-and-aws-team-up-to-build-cloud-computing-infrastructure-to-support-ai-innovation>.

“SK Telecom Improves Telco-Specific Q&a by Fine-Tuning Anthropic’s Claude Models in Amazon Bedrock | Amazon Web Services.” *Amazon Web Services*, 10 Oct. 2024, <https://aws.amazon.com/it/blogs/machine-learning/sk-telecom-improves-telco-specific-qa-by-fine-tuning-anthropics-claude-models-in-amazon-bedrock/>.

“SKT Partnership Announcement.” *Anthropic.com*, 2023, www.anthropic.com/news/skt-partnership-announcement.

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

“The AI Alliance: Our First Year | AI Alliance.” *Thealliance.ai*, 5 Dec. 2024, <https://thealliance.ai/blog/our-first-year>.

“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>.

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. “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. “Microsoft Investe 10 Miliardi in OpenAi.” *Wired Italia*, Wired Italia, 23 Jan. 2023, www.wired.it/article/microsoft-openai-chatgpt/

Bettya. “Meta Partners with Hugging Face & Scaleway to Support Open Source.” *Meta Newsroom*, 8 Nov. 2023, <https://about.fb.com/news/2023/11/meta-partners-with-hugging-face-scaleway-to-support-open-source/>

Bettya. “Meta, Hugging Face, and Scaleway Announce a New AI Accelerator Program for European Startups.” *Meta Newsroom*, 24 June 2024, <https://about.fb.com/news/2024/06/meta-hugging-face-and-scaleway-announce-a-new-ai-accelerator-program-for-european-startups/>

Cai, Kenrick, and Krystal Hu. “Exclusive: OpenAI Taps Google in Unprecedented Cloud Deal despite AI Rivalry, Sources Say.” *Reuters*, 10 June 2025, www.reuters.com/business/retail-consumer/openai-taps-google-unprecedented-cloud-deal-despite-ai-rivalry-sources-say-2025-06-10/.

Colletti, Riccardo. “Amazon Pronta a Un Nuovo Investimento Da 8 Miliardi Di Dollari in Anthropic (Financial Times).” *IGizmo.it*, 10 July 2025, www.igizmo.it/amazon-pronta-a-un-nuovo-investimento-da-8-miliardi-di-dollari-in-anthropic-financial-times/.

Congiu, Gabriele. “Nvidia Investirà Sino a 100 Miliardi Di Dollari in OpenAI: C’è L’accordo.” *HDblog.it*, 23 Sept. 2025, www.hdblog.it/nvidia/articoli/n632412/nvidia-investimento-openai-100-miliardi-dollari/.

Edwards, Benj, and Adamà Faye. “OpenAI E Nvidia, Cosa Prevede Il Mega-Accordo Da 100 Miliardi.” *Wired Italia*, Wired Italia, 23 Sept. 2025, www.wired.it/article/openai-nvidia-accordo-partnership-100-miliardi-data-center-intelligenza-artificiale/.

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

Forbes.it. “Microsoft Guarda Ad Anthropic per Ridurre La Dipendenza Da OpenAI.” *Forbes Italia*, 10 Sept. 2025, <https://forbes.it/2025/09/10/microsoft-guarda-anthropic-ridurre-dipendenza-openai>.

Frasso, Edoardo. “Accordo Tra La California E NVIDIA per Portare Risorse AI Nei College - AI News.” *AI News*, 12 Aug. 2024, <https://ainews.it/accordo-tra-la-california-e-nvidia-per-portare-risorse-lai-nei-college/>.

Frasso, Edoardo. “Il NYT Scopre Che Google Possiede Il 14% Di Anthropic: Le Aziende Sono Davvero Avversarie? - AI News.” *AI News*, 13 Mar. 2025, <https://ainews.it/il-nyt-scopre-che-google-possiede-il-14-di-anthropic-le-aziende-sono-davvero-avversarie/>.

Garay, Jorge, and Riccardo Piccolo. “Anthropic, Google Investe 2 Miliardi Nella Rivala Di OpenAI.” *Wired Italia*, Wired Italia, 2 Nov. 2023, www.wired.it/article/anthropic-google-investe-2-miliardi-openai-claude-2-chatgpt/.

Geyer, Mike. “Mercedes-Benz Prepares Its Digital Production System for Next-Gen Platform with NVIDIA Omniverse, MB.OS and Generative AI.” *NVIDIA Blog*, 20 Sept. 2023, <https://blogs.nvidia.com/blog/mercedes-benz-ev-nvidia-omniverse-generative-ai/>.

Goldman Sachs. “Generative AI Could Raise Global GDP by 7%.” *Www.goldmansachs.com*, 5 Apr. 2023, www.goldmansachs.com/insights/articles/generative-ai-could-raise-global-gdp-by-7-percent.html.

Guardian staff reporter. “Meta and IBM Launch “AI Alliance” to Promote Open-Source AI Development.” *The Guardian*, The Guardian, 5 Dec. 2023, www.theguardian.com/technology/2023/dec/05/open-source-ai-meta-ibm.

Karanikas, Dino. “In Recent Developments, OpenAI, the Organization behind the Pioneering AI Models such as ChatGPT, Has Announced a Strategic Partnership with Reddit, a Vast Network of Communities Where People Engage over Shared Interests, Hobbies, and Passions. This Collaboration Is Pivotal for Several Reasons and More.” *Linkedin.com*, 17 May 2024, www.linkedin.com/pulse/exploring-strategic-partnership-between-openai-reddit-dino-karanikas-3nste.

Knight, Will, and Adamà Faye. “Amazon E Anthropic Stanno Costruendo Un Supercomputer AI.” *Wired Italia*, Wired Italia, 4 Dec. 2024, www.wired.it/article/amazon-anthropic-supercomputer-intelligenza-artificiale/.

Knight, Will. “Microsoft Makes a \$16 Billion Entry into Health Care AI.” *WIRED*, 13 Apr. 2021, www.wired.com/story/microsoft-dollar16-billion-entry-health-care-ai/.

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](https://www.espressonline.it/presenza-futuro/notizie/intelligenza-artificiale-sbaglia-innocente-finisce-carcere-usa-f4ca2cb0-a9a4-11eb-8b01-83c2a483d7f5.shtml).

Luca Tremolada. “Blog | Apple Intelligence, La Privacy Nell’era Dell’Ai E l’Accordo Con OpenAi - Info Data.” *Info Data*, 17 June 2024, www.infodata.ilsole24ore.com/2024/06/17/apple-intelligence-la-privacy-nellera-dellai-e-laccordo-con-openai/.

Maso, Elena Dal. “Asml Investe 1,3 Miliardi in Mistral AI, Startup Francese. Fra I Soci, Nvidia E Il Governo.” *MF Milano Finanza*, Milano Finanza, 9 Sept. 2025, www.milanofinanza.it/news/asml-investe-1-3-miliardi-in-mistral-ai-startup-francese-fra-i-soci-nvidia-e-il-governo-202509090824288839.

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.

Mele, Pierluigi. “Dall’iPhone Alla Mente Digitale: L’audace Mossa Di Apple Su Perplexity AI.” *RaiNews*, 23 June 2025, www.rainews.it/articoli/2025/06/dalliphone-alla-mente-digitale-laudace-mossa-di-apple-su-perplexity-ai-5f185099-42bc-4112-90f5-5afea2e411ad.html.

mercedes.galan@lcpublishinggroup.com. “PwC, Harvey, and Lefebvre Launch AI-Powered Platform.” *Iberian Lawyer*, 3 Apr. 2025, <https://iberianlawyer.com/pwc-harvey-and-lefebvre-launch-ai-powered-platform/>.

Metz, Cade, et al. “Inside Google’s Investment in Anthropic.” *The New York Times*, 11 Mar. 2025, www.nytimes.com/2025/03/11/technology/google-investment-anthropic.html.

MF Milano Finanza. “OpenAI Si Affida a Oracle, Maxi Accordo Da 300 Miliardi Di Dollari per Il Futuro Di ChatGpt.” *MF Milano Finanza*, Milano Finanza, 9 Nov. 2025, www.milanofinanza.it/news/oracle-e-openai-firmano-un-accordo-da-300-miliardi-di-dollari-nel-settore-cloud-202509111133243109.

Microsoft Source. “Microsoft Accelerates Industry Cloud Strategy for Healthcare with the Acquisition of Nuance - Source.” *Source*, 12 Apr. 2021,

<https://news.microsoft.com/2021/04/12/microsoft-accelerates-industry-cloud-strategy-for-healthcare-with-the-acquisition-of-nuance/>.

Morgantini, Federico, and DigiTechNews - Digital | Technology | Learning. “Ecco in Cosa Consiste Il Nuovo Accordo Tra OpenAI E Reddit.” *DigiTech.News*, DigiTechNews - Digital | Technology | Learning, 17 May 2024, www.digitech.news/digital/17/05/2024/openai-reddit-accordo/.

Mosca, Giuditta. “Che Cos’è Reddit E Perché Se Ne Parla Così Tanto.” *La Repubblica*, 15 June 2023, www.repubblica.it/tecnologia/2023/06/15/news/reddit_come_funziona_tutorial_subreddit_api-404431601/.

Murgida, Rosario. “Guida Autonoma: Mercedes Utilizzerà La Piattaforma Nvidia | Quattroruote.it.” *Quattroruote.it*, Quattroruote, 24 June 2020, www.quattroruote.it/news/tecnologia/2020/06/24/guida_autonoma_dal_2024_la_mercedes_utilizzera_la_piattaforma_nvidia.html.

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

OpenAI. “OpenAI and Microsoft Extend Partnership.” *Openai.com*, 2023, <https://openai.com/index/openai-and-microsoft-extend-partnership/>.

OpenAI. “OpenAI and Apple Announce Partnership.” *Openai.com*, 10 June 2024, <https://openai.com/index/openai-and-apple-announce-partnership/>.

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.

Parvini, Sarah. “California Partners with Nvidia to Bring Artificial Intelligence Resources to Colleges.” *The Seattle Times*, 9 Aug. 2024, www.seattletimes.com/business/california-partners-with-nvidia-to-bring-artificial-intelligence-resources-to-colleges/.

Patella, Alessandro. “Amazon Completa Il Suoi Investimento Da 4 Miliardi in Anthropic.” *Wired Italia*, Wired Italia, 28 Mar. 2024, www.wired.it/article/amazon-anthropic-investimento/.

Patella, Alessandro. “Stargate, Cosa Sappiamo Del Gigante Dell’AI Che Trump Vuole Creare.” *Wired Italia*, Wired Italia, 22 Jan. 2025, www.wired.it/article/stargate-cos-e-intelligenza-artificiale-trump-oracle-open-ai-softbank/.

Pisa, Pier Luigi. “Microsoft E OpenAI Hanno Un Accordo Segreto Sulla Definizione Di AGI.” *La Repubblica*, 27 Dec. 2024, www.repubblica.it/tecnologia/dossier/intelligenza-artificiale/2024/12/27/news/microsoft_openai_agi_accordo_segreto_intelligenza_artificiale_generale-423908793/.

Pisa, Pier Luigi. “Perplexity, Trattative Con Samsung E Motorola per Integrare La Ricerca IA Negli Smartphone.” *La Repubblica*, 17 Apr. 2025, www.repubblica.it/tecnologia/2025/04/17/news/perplexity-ai-integrazione-smartphone-samsung-motorola-424133962/.

Press, Associated. “Amazon Pours Additional \$2.75bn into AI Startup Anthropic.” *The Guardian*, 27 Mar. 2024, www.theguardian.com/technology/2024/mar/27/anthropic-amazon-ai-startup.

PricewaterhouseCoopers. “PwC Announces Strategic Alliance with Harvey, Positioning PwC’s Legal Business Solutions at the Forefront of Legal Generative AI.” *PwC*, 15 Mar. 2023, www.pwc.com/gx/en/news-room/press-releases/2023/pwc-announces-strategic-alliance-with-harvey-positioning-pwcs-legal-business-solutions-at-the-forefront-of-legal-generative-ai.html.

Ruffilli, Bruno. “Apple Valuta Un Accordo Con Google per Portare Gemini Dentro Siri.” *La Repubblica*, 4 Sept. 2025, www.repubblica.it/tecnologia/2025/09/04/news/apple_valuta_una_partnership_con_google_gemini_per_rivoluzionare_l_ai_di_siri-424825711/.

Ruffilli, Bruno. “Per Avere ChatGPT Su iPhone, Apple Ha Pagato OpenAI in Visibilità.” *La Repubblica*, 14 June 2024, www.repubblica.it/tecnologia/2024/06/14/news/per_avere_chatgpt_su_iphone_apple_ha_pagato_openai_in_visibilita-423233676/.

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/.

Simonetta, Biagio. "Intelligenza Artificiale, Amazon Punta Altri 4 Miliardi Su Anthropic." *Il Sole 24 ORE*, 22 Nov. 2024, www.ilsole24ore.com/art/intelligenza-artificiale-amazon-punta-altri-4-miliardi-anthropic-AGZDE2KB.

ThasmikaGokal. "Microsoft and Hugging Face Deepen Generative AI Partnership." *TECHCOMMUNITY.MICROSOFT.COM*, 21 May 2024, <https://techcommunity.microsoft.com/blog/azure-ai-foundry-blog/microsoft-and-hugging-face-deepen-generative-ai-partnership/4144565>.

Venus. "Mercedes E Nvidia Insieme per l'Auto Del Futuro - Venus." *Venus*, 23 July 2020, www.venus-spa.it/mercedes-e-nvidia-insieme-per-lauto-del-futuro/.

Viner, Jonathan. "Academic Software Partners with IBM Watsonx: AI for Education." *Academicsoftware.com*, Lernova, 18 Dec. 2024, <https://academicsoftware.com/en-gb/news/ibm-watsonx-academic-software-partners>.

Wiggers, Kyle. "Meta and IBM Form an AI Alliance, but to What End? | TechCrunch." *TechCrunch*, 5 Dec. 2023, <https://techcrunch.com/2023/12/05/meta-and-ibm-form-an-ai-alliance-but-to-what-end/>.

Yanitsa Boyadzhieva. "Amazon Pumps Another \$2.75bn into Anthropic." *TelecomTV*, 28 Mar. 2024, www.telecomtv.com/content/digital-platforms-services/amazon-pumps-another-2-75bn-into-anthropic-50049/.