



UNIVERSITÀ
DI PAVIA

DIPARTIMENTO DI SCIENZE POLITICHE E SOCIALI

CORSO DI LAUREA MAGISTRALE IN COMUNICAZIONE DIGITALE

**Intelligenza artificiale e nuove diseguaglianze
sociali nella transizione algoritmica: visibilità, lavoro e
potere.**

Relatore:

Chiar.mo Dott./Prof. Flavio Antonio Ceravolo

Correlatore:

Chiar.mo Dott./Prof. Andrea Fontana

Tesi di laurea di

Matteo Bucci

Matricola n.

561772

ANNO ACCADEMICO 2025/26

Sommario

Introduzione	4
Capitolo 1 Vecchie e nuove diseguaglianze.....	8
1.1 Introduzione	8
1.2 La sociologia di fronte alle diseguaglianze.....	9
1.3 Marx e Weber: la tradizione classica.....	13
1.3.1 Marx e la classe come rapporto di produzione.....	13
1.3.2 Weber: la triade classe-status-partito	15
1.3.3 Il confronto: convergenze e divergenze	17
1.4 La stratificazione nella sociologia contemporanea	18
1.4.1 Le due eredità: neo-marxismo e neo-weberismo.....	19
1.4.2 La mobilità sociale: aperture e chiusure.....	20
1.4.3 La teoria della frammentazione.....	22
1.5 Le diseguaglianze oltre la classe: genere ed etnia.....	23
1.5.1 La stratificazione di genere.....	24
1.5.2 La stratificazione etnico-razziale.....	25
1.6 La riproduzione delle diseguaglianze: dal capitale culturale alle credenziali educative	27
1.6.1 Bourdieu e il capitale culturale	28
1.6.2 Le credenziali educative e l'inflazione del titolo	29
1.7 Dalla stratificazione sociale alla soglia del digitale	31
Capitolo 2 Dal digital divide all'AI divide.....	34
2.1 Introduzione	34
2.2 Il digital divide: origine del concetto e critica della metafora.....	35
2.3 Il modello multidimensionale di van Dijk	40
2.4 I tre livelli del divario digitale.....	46

2.4.1 Primo livello: l'accesso fisico.....	47
2.4.2 Secondo livello: competenze e modalità d'uso.....	49
2.4.3 Terzo livello: i benefici concreti.....	52
2.5 Dalla svolta algoritmica all'algorithmic divide.....	55
2.6 L'AI divide	62
2.6.1 Continuità e discontinuità tra digital divide, algorithmic divide e AI divide.....	63
2.6.2 La definizione operativa di Wang et al. (2025).....	66
2.6.3 Uso consapevole e uso subito: verso il capitale algoritmico	68
2.7 Una traiettoria di lungo periodo, dall'accesso all'algoritmo	71
Capitolo 3. Le tre dimensioni dell'AI divide: visibilità, lavoro e potere..	76
3.1 Introduzione	76
3.2 Visibilità: chi appare, chi scompare nei flussi informativi.....	79
3.2.1 Filtri algoritmici e personalizzazione: la filter bubble	81
3.2.2 Echo chambers e polarizzazione: la dimensione sociale e cognitiva	83
3.2.3 Il caso italiano: dispercezioni e meccanismi cognitivi.....	87
3.2.4 Visibilità nell'AI generativa: dati di training e amplificazione dei bias	90
3.3 Lavoro: chi è esposto, chi guadagna, chi perde nei mercati algoritmici93	
3.3.1 Il quadro internazionale: l'AI come frontiera della trasformazione del lavoro	94
3.3.2 L'esposizione italiana all'AI: tra rischio di sostituzione e potenziale di complementarità	97
3.3.3 L'impatto distributivo sui salari italiani: la polarizzazione del mercato del lavoro	101
3.3.4 Lavoro nell'AI generativa: vincitori, perdenti e shift in expertise	103
3.4 Potere: chi decide, chi è deciso, chi è regolato dagli algoritmi	106

3.4.1 Surveillance capitalism e Big Other: la formulazione fondativa di Zuboff.....	108
3.4.2 Discriminazione algoritmica e diritti umani: il framework di Aizenberg e van den Hoven.....	110
3.4.3 La risposta regolatoria italiana: la prassi AGCOM ricostruita da Giannelli	114
3.4.4 Potere nell'AI generativa: dal framework di Aizenberg al regime post-ChatGPT.....	117
3.5 Tre dimensioni, un solo meccanismo	119
Capitolo 4. Dalla diagnosi alla proposta: l'AI divide nel contesto italiano ed europeo.....	124
4.1 Introduzione e metodo.....	124
4.2 Il divario esiste e cresce: l'evidenza dei dati.....	127
4.3 Chi avanza e chi resta indietro: la lettura degli esperti.....	130
4.4 Cosa si può fare: il dibattito sulle risposte	133
4.5 Una direzione possibile: dalle competenze al territorio	136
Conclusioni.....	141
Appendice delle interviste.....	151
Bibliografia.....	179
Sitografia	184

Introduzione

Nel giro di pochi anni l'intelligenza artificiale ha smesso di essere una materia per specialisti, ed è entrata nell'esperienza quotidiana di milioni di persone. La incontriamo nei sistemi che selezionano le informazioni che leggiamo, negli strumenti che assistono i nostri ambienti di lavoro, e ora anche nei dispositivi che generano testi e immagini su richiesta. Una diffusione così rapida ha alimentato un dibattito pubblico fortemente polarizzato. Per alcuni si tratta di una tecnologia capace di allargare le opportunità, e di mettere a disposizione di molti strumenti un tempo riservati a pochi. Per altri rappresenta una minaccia per l'occupazione e per l'autonomia delle persone, destinata a concentrare vantaggi e potere in poche mani. Entrambe le letture, prese da sole, risultano insoddisfacenti, perché trattano l'intelligenza artificiale come un agente che produce effetti uniformi sull'intera società, siano essi benefici o dannosi. La prospettiva adottata in questo lavoro è diversa. La domanda che orienta la ricerca non riguarda la natura intrinseca della tecnologia, ma il modo in cui ne vengono distribuiti gli effetti. Si tratta di capire chi si trovi avvantaggiato dalla nuova ondata e chi penalizzato, lungo quali linee questa distinzione passi, e in che rapporto la nuova distribuzione stia con le disuguaglianze che già attraversano la società. È una domanda di natura sociologica, e gli strumenti per affrontarla sono quelli delle scienze sociali.

L'ipotesi della tesi muove da qui. Quando l'intelligenza artificiale entra nei contesti sociali, opera su asimmetrie che esistono già: le riprende, le amplifica, e le riconfigura sotto forme nuove. Da questa premessa muove un lavoro di critica sociale, che intende mettere in dialogo due tradizioni di studio. Da un lato c'è la riflessione sociologica sulle disuguaglianze, che fornisce le categorie di fondo: la distribuzione diseguale delle risorse, i meccanismi di riproduzione del vantaggio attraverso le generazioni, e le tre dimensioni di struttura, cultura e potere

lungo cui la disciplina ha imparato a leggere la stratificazione sociale. C'è poi la ricerca sul divario digitale, e in particolare il modello multidimensionale di Jan van Dijk. Il modello ha mostrato come la distanza tra utenti che beneficiano delle tecnologie e utenti che ne restano esclusi non sia mai stata, semplicemente, una questione di accesso fisico. Si tratta invece del prodotto di una catena causale che lega le risorse sociali alle competenze e ai benefici concreti, in un rapporto di reciproca alimentazione con le disuguaglianze di partenza. Dall'incontro di queste due tradizioni prende forma la categoria attorno a cui ruota il lavoro. Si tratta del concetto di capitale algoritmico, ossia l'insieme delle conoscenze, delle competenze e delle disposizioni che permettono a un individuo di interagire con i sistemi algoritmici e di intelligenza artificiale in modo consapevole e strategico. La sua distribuzione diseguale è il filo che collega l'analisi teorica del lavoro alla sua verifica empirica.

All'interno di questa cornice, la ricerca si articola attorno a quattro domande, che reggono l'intero percorso e a cui le considerazioni conclusive torneranno per offrire una risposta complessiva. La prima domanda riguarda la natura del fenomeno: *l'AI divide costituisce una frattura inedita rispetto alle disuguaglianze del passato, oppure rappresenta la continuazione, in forma nuova, di dinamiche di esclusione già note?* Sulle manifestazioni concrete del divario verte la seconda domanda: *attraverso quali dimensioni della vita sociale si manifesta l'AI divide?* Il nodo distributivo, che tocca più da vicino il tema delle disuguaglianze, è l'oggetto della terza interrogazione: *nella transizione algoritmica, chi avanza e chi resta indietro, e con quali criteri si ridisegnano le posizioni nella società?* Lo sguardo si sposta infine dalla diagnosi alle risposte concrete: *di fronte all'AI divide, quali interventi sono possibili sul piano delle politiche pubbliche, e la sola regolazione è sufficiente a contrastarlo?* Si tratta di domande poste in forma aperta,

senza presupporre la risposta, e il percorso che segue ha il compito di costruirla passo dopo passo.

Per rispondere a questi interrogativi il lavoro intreccia due piani. Sul piano della ricostruzione teorica si ripercorre la letteratura sociologica sulle diseguaglianze e quella ormai venticinquennale sul divario digitale, fino agli studi più recenti dedicati all'intelligenza artificiale. Sul piano empirico, di cui si dirà tra un momento, la ricerca si concentra soprattutto sul contesto italiano. La tesi si appoggia qui a una pluralità di dati già esistenti. Ci sono rilevazioni internazionali e nazionali sull'esposizione delle professioni all'automazione, e ricerche specifiche sull'impatto distributivo dell'AI sui salari. Ho fatto poi ricorso alle statistiche ufficiali sull'adozione dell'intelligenza artificiale nelle imprese italiane, ad alcune indagini europee sulla percezione del rischio occupazionale, e infine agli studi più recenti sui meccanismi di selezione algoritmica dei contenuti. A questa base oggettiva si affianca un'indagine qualitativa originale, costruita attorno a una serie di interviste a testimoni privilegiati. Sono testimoni scelti da posizioni diverse e complementari: formatori manageriali, consulenti per le imprese, imprenditori del settore tecnologico, ricercatori in scienze sociali. Al ciclo di interviste si aggiunge poi un questionario esplorativo che, dato il numero contenuto di risposte raccolte, vale come primo sondaggio di tendenze più che come rilevazione rappresentativa. L'impianto empirico risulta dunque prevalentemente qualitativo, e impiega il dato statistico come sfondo di contesto su cui leggere le testimonianze raccolte.

Il percorso si sviluppa in quattro capitoli, che muovono progressivamente dal generale al particolare. Il primo capitolo definisce la cornice sociologica del lavoro, e ripercorre il modo in cui le scienze sociali hanno pensato le diseguaglianze. Si parte dalla tradizione classica di Marx e Weber, si passa agli studi sulla stratificazione e sulla mobilità sociale, e

si arriva alla riflessione sulla riproduzione del vantaggio attraverso il capitale culturale e le credenziali educative. È in questo capitolo che vengono isolate le tre dimensioni di struttura, cultura e potere, destinate a riconfigurarsi nei capitoli successivi. Quelle categorie vengono poi applicate al divario digitale nel secondo capitolo, che ricostruisce in forma genealogica venticinque anni di ricerca. Si parte dal modello di van Dijk e dai tre livelli classici dell'accesso, delle competenze e dei benefici, per arrivare a mostrare come ogni ondata tecnologica abbia riposizionato il divario senza superarlo, fino all'emergere *dell'algorithmic divide* e *dell'AI divide* e alla formulazione del concetto di capitale algoritmico. Il terzo capitolo si occupa delle forme concrete che il divario assume, articolando la distribuzione diseguale del capitale algoritmico lungo tre dimensioni interconnesse. Una è quella dei flussi informativi, dove si pone la questione della visibilità degli utenti al loro interno. C'è poi la posizione che ciascuno occupa nei mercati del lavoro trasformati dall'AI. E c'è infine la dimensione del controllo, ossia il potere effettivo di intervento sui sistemi algoritmici. Le tre dimensioni non operano in modo separato fra loro: si rinforzano a vicenda lungo le linee della stratificazione sociale di partenza. Il quarto capitolo, infine, porta l'analisi sul terreno empirico del contesto italiano ed europeo. Mette alla prova il quadro teorico costruito nei capitoli precedenti attraverso i dati e le voci degli esperti, e apre alla parte propositiva del lavoro. Le considerazioni conclusive raccolgono le fila dell'intero percorso, rispondono alle quattro domande iniziali, e avanzano una proposta articolata sul piano delle politiche pubbliche.

Capitolo 1 Vecchie e nuove diseguaglianze

1.1 Introduzione

Le diseguaglianze attraversano la sociologia fin dalle sue origini. Come una società distribuisca risorse, opportunità, riconoscimento e potere tra i propri membri è una questione che la disciplina si è posta sin dai suoi primi sviluppi, e che continua a riformularsi a ogni nuova fase storica. Spiegare perché certe posizioni sociali si trasmettono attraverso le generazioni, perché i vantaggi tendono a concentrarsi, perché lo svantaggio si riproduce, significa entrare nel campo di indagine specifico della sociologia, distinto dalle letture individuali o naturalistiche che attribuiscono il successo o il fallimento esclusivamente alle qualità dei singoli.

L'obiettivo di questo capitolo è circoscritto, si tratta di isolare le coordinate concettuali fondamentali con cui la disciplina ha affrontato il tema, e di mostrare in che senso questo patrimonio resti utile per leggere le forme di disparità contemporanee. La scelta risponde a un'esigenza precisa, le diseguaglianze digitali e algoritmiche che verranno discusse nei capitoli successivi si innestano su disparità sociali preesistenti, e ne riproducono sotto forme nuove le linee di frattura. Senza una cornice di riferimento, l'indagine sulle diseguaglianze di nuova generazione peccherebbe di astrattezza. Si finirebbe per avallare l'idea ingannevole secondo cui i sistemi tecnologici generino le disparità dal nulla, ignorando come, nei fatti, si limitino a sussumere e dilatare asimmetrie già ampiamente storicizzate.

Per questo motivo, l'ossatura del capitolo si articola su uno sviluppo congiunto, di tipo storico e concettuale. In via preliminare, il paragrafo 1.2 mette a fuoco il problema chiarendo i confini semantici di tre costrutti

guida: *stratificazione, risorse e status*. I successivi paragrafi 1.3 e 1.4 operano una sintesi del patrimonio classico, muovendo dalle prospettive di Marx e Weber per approdare agli studi novecenteschi in materia di mobilità sociale. Chiude il quadro il paragrafo 1.5, pensato per integrare una dimensione spesso trascurata dall'ortodossia sociologica del passato: la stratificazione non si esaurisce nella variabile di classe, ma si ibrida costantemente con il genere, le coordinate etniche e le appartenenze territoriali. Una configurazione complessa, insomma, che la ricerca odierna ha ormai codificato e imparato a trattare in termini di pluralità di assi di stratificazione.

Due sono i testi di riferimento: per l'inquadramento contemporaneo del dibattito utilizzeremo *Sociologia generale. Teorie, metodo, concetti* di David Croteau e William Hoynes, nella terza edizione italiana del 2022, per la ricostruzione del pensiero classico ci appoggeremo a *Teorie sociologiche* di Randall Collins (1992). I due manuali si integrano a vicenda: il primo offre l'articolazione attuale del dibattito sociologico sulle disuguaglianze, con particolare attenzione al contesto italiano e globale; il secondo, la profondità storica delle tradizioni teoriche da cui diviene quel dibattito.

1.2 La sociologia di fronte alle disuguaglianze

Prima di affrontare la ricostruzione storica del dibattito conviene fissare alcuni termini. Croteau e Hoynes definiscono la disuguaglianza sociale come «distribuzione ineguale di risorse economiche, sociali, politiche e culturali all'interno di un determinato contesto sociale» (Croteau & Hoynes, 2022, p. 313). Il punto qualificante di questa definizione è il carattere strutturato della disuguaglianza, che permette di distinguere dalle semplici differenze personali. Non tutte le differenze tra individui producono gerarchie: gli orientamenti politici, i gusti culturali, i

percorsi biografici variano da persona a persona senza generare strutture di subordinazione. La diseguaglianza esiste quando queste differenze si traducono in posizioni gerarchiche stabili, sostenute da norme e pratiche che le rendono durature, collocando gli individui «in una gerarchia di gruppi che ricevono risorse disuguali» (Croteau & Hoynes, 2022, p. 313). Proprio attorno a questa matrice strutturale la sociologia ha saputo rivendicare la propria autonomia scientifica. Una presa di distanza netta, in sostanza, sia dai determinismi di matrice biologica sia da quelle interpretazioni riduzionistiche che tendono a derubricare le disparità a una mera sommatoria di traiettorie individuali.

Per inquadrare una simile architettura gerarchica, il lessico disciplinare ricorre al concetto di *stratificazione*. L'impiego di un'immagine di chiara derivazione geologica aiuta a focalizzare il punto: in modo non dissimile dalle formazioni rocciose costituite per sedimentazione, anche gli aggregati umani appaiono ripartiti in fasce sovrapposte, ciascuna caratterizzata da un grado diseguale di accesso ai beni collettivi. Lo scarto essenziale rispetto al mondo naturale risiede, ovviamente, nell'assenza di staticità. Gli strati sociali non sono mai dati una volta per tutte; al contrario, prendono forma attraverso un incessante processo di riproduzione alimentato dall'agire quotidiano, dal peso degli apparati istituzionali e dai codici normativi che ne assicurano la tenuta.

Spingendo lo sguardo sul piano storico, Croteau e Hoynes rintracciano le prime evidenze di un assetto stratificato nel Neolitico, grosso modo ottomila anni fa, in perfetta concomitanza con il consolidarsi dell'agricoltura stanziale e l'affermarsi dei primi nuclei urbani (Croteau & Hoynes, 2022, p. 312). Non si tratta di un semplice dettaglio cronologico, ma di uno snodo teorico cruciale. Questo dato dimostra infatti come la diseguaglianza non rappresenti affatto un corollario ineluttabile della condizione umana, quanto piuttosto l'esito di particolarissime configurazioni organizzative. A riprova di ciò, e per quanto i riscontri

archeologici consentano di ricostruire, le antecedenti comunità di cacciatori-raccoglitori sembravano del tutto sprovviste di intelaiature gerarchiche assimilabili a quelle esplose con le società complesse successive.

Alla luce di questa multidimensionalità, l'idea di leggere le asimmetrie servendosi di una lente puramente economicistica finirebbe inevitabilmente, secondo gli autori, per comprometterne la reale portata euristica. Croteau e Hoynes (2022, p. 314) mappano infatti sette distinte tipologie di risorse la cui allocazione asimmetrica configura l'ossatura dei moderni sistemi di stratificazione:

- Le risorse *economiche*: che racchiudono la ricchezza monetaria, i beni immobili e la proprietà terriera;
- Le risorse *umane*: connesse ai percorsi di scolarizzazione, all'addestramento e alle competenze tecnico-professionali;
- Le risorse *culturali*: intese come quel bagaglio di cognizioni implicite e competenze informali trasmesse mediante i processi di socializzazione;
- Le risorse *sociali*: riferite al capitale relazionale e all'inserimento in network strategici capaci di veicolare informazioni, tutele e opportunità;
- Le risorse *di status*: legate al prestigio, all'onorabilità e al gradimento sociale;
- Le risorse *civili*: inerenti all'esercizio dei diritti di proprietà, alla tutela contrattuale, all'elettorato e all'accesso agli istituti di giustizia;
- Le risorse *politiche*: concernenti l'esercizio del potere formale e la capacità di orientare i processi decisionali collettivi.

L'interazione sinergica di questi sette vettori determina l'esatta collocazione di un soggetto nello spazio sociale. Ciò delucida il motivo per cui due individui provvisti del medesimo reddito possano esperire

condizioni di vita radicalmente asimmetriche: il possesso di un cospicuo capitale sociale o culturale, sebbene in costanza di entrate finanziarie modeste, garantisce strumenti di emancipazione e di azione preclusi ad altri.

A questa tassonomia delle risorse si affianca una seconda differenziazione teorica, focalizzata sui meccanismi di mobilità e di reclutamento attraverso cui i singoli soggetti si collocano in una determinata posizione entro la morfologia sociale. La sociologia distingue tra *status conseguito* e *status ascritto*. Il primo è la «posizione che si ottiene in gran parte ai propri sforzi» (Croteau & Hoynes, 2022, p. 315): ruoli professionali, titoli di studio, posizioni raggiunte attraverso scelte e azioni individuali. Lo status ascritto è invece assegnato indipendentemente dalla volontà o dalle capacità del singolo. Dipende da caratteristiche come la nascita, il genere, l'etnia o la famiglia di origine. Su questa distinzione si fonda la classificazione dei sistemi di stratificazione in *aperti* e *chiusi*. Nei sistemi aperti gli status conseguiti pesano in modo significativo, e la mobilità sociale è praticabile. Nei sistemi chiusi prevalgono gli status ascritti, e le posizioni si trasmettono in modo pressoché meccanico attraverso le generazioni. Studiare la stratificazione di una società significa anche, e in larga parte, valutare quanto quella società sia aperta o chiusa, e quali meccanismi rendano possibile o ostacolino il movimento tra le posizioni.

Croteau e Hoynes dichiarano in modo esplicito l'architettura concettuale con cui organizzano la propria trattazione delle disuguaglianze, articolata attorno a tre dimensioni: *struttura*, *cultura* e *potere* (Croteau & Hoynes, 2022, p. 311). La dimensione strutturale riguarda il fatto che l'azione individuale si svolge sempre «all'interno di un contesto sociale e culturale che influenza sia le opzioni reali di cui dispone sia i fattori economici, culturali e sociali ereditati» (Croteau & Hoynes, 2022, p. 347). La dimensione culturale rinvia al fatto che le

differenze di classe non si fermano alla ricchezza materiale, ma comprendono stili di vita e valori trasmessi dalla socializzazione. La dimensione del potere riguarda invece le asimmetrie nella capacità di influenzare le scelte collettive e nell'accesso ai luoghi della decisione. Questa griglia analitica orienta il presente capitolo: le stesse tre dimensioni che la sociologia usa per leggere le diseguaglianze tradizionali si ritroveranno, riconfigurate, nell'analisi delle diseguaglianze prodotte dall'intelligenza artificiale.

1.3 Marx e Weber: la tradizione classica

I modelli teorici di cui la sociologia si serve per analizzare le diseguaglianze derivano in larga parte dai contributi di Karl Marx e Max Weber. Sebbene abbiano operato a breve distanza temporale l'uno dall'altro, i due autori hanno sviluppato prospettive diverse, capaci di fissare il vocabolario fondamentale della disciplina e di orientare il dibattito scientifico per oltre un secolo. Questa sezione ricostruisce l'approccio di entrambi i pensatori, proponendo in chiusura un confronto volto a evidenziare i punti di convergenza e di distacco che, ancora oggi, alimentano due distinte tradizioni di ricerca.

1.3.1 Marx e la classe come rapporto di produzione

Per Karl Marx (1818-1883) lo studio della società ruota attorno al concetto di *classe*, intesa come la posizione oggettiva che gli individui occupano all'interno dei *rapporti di produzione*. In quest'ottica, la faglia principale che struttura la società capitalistica divide chi possiede i mezzi di produzione da chi ne è escluso. La stratificazione non viene quindi ridotta a una semplice differenza di reddito o di consumi, ma è vista come un rapporto di forza intrinsecamente conflittuale. Croteau e Hoynes sintetizzano così la tesi di Marx: «La struttura fondamentale della società

è stata sempre la stessa: una netta divisione tra chi possiede i mezzi di produzione e chi non li possiede [...] La risorsa principale non è più la terra ma il capitale» (Croteau & Hoynes, 2022, p. 322). Il *capitale*, nel vocabolario teorico di Marx, si configura come il «denaro da investire in fabbriche, terreni e altre imprese» (Croteau & Hoynes, 2022, p. 322); l'oligopolio su questa risorsa definisce la *classe capitalista*, o *borghesia*. Al polo opposto, quanti dispongono unicamente della propria *forza lavoro* e si trovano costretti ad alienarla in cambio di un salario compongono il *proletariato*. Si tratta di un modello dicotomico che semplifica programmaticamente la complessità empirica al fine di far emergere la legge fondamentale del sistema. Sebbene Marx fosse pienamente consapevole dell'esistenza di strati intermedi, egli riteneva che le tendenze di lungo periodo del modo di produzione capitalistico avrebbero finito per fagocitarli entro i due poli principali. Tale dinamica prende il nome di *proletarizzazione*, processo che descrive la progressiva perdita di autonomia economica da parte dei ceti medi, declassati a vendere le proprie prestazioni lavorative alle medesime condizioni del proletariato industriale. L'analisi di Marx non si presenta come neutrale rispetto ai valori. Marx interpreta il legame tra le classi come strutturalmente antagonista, poiché gli interessi materiali della borghesia e quelli del proletariato risultano inconciliabili. Questa conflittualità non è l'esito delle intenzioni psicologiche degli individui, ma deriva dalla loro collocazione strutturale nei processi produttivi. Muovendo da tali premesse, Marx prefigura un modello alternativo di organizzazione sociale: il *socialismo*, inteso come un sistema economico «in cui lo Stato detiene i grandi mezzi di produzione per conto dei lavoratori, abolendo così le distinzioni di classe che si basano sulla proprietà privata» (Croteau & Hoynes, 2022, p. 322). Al di là delle valutazioni sulle implicazioni politiche di tale dottrina, ciò che preme evidenziare in questa sede è il nucleo sociologico del contributo marxiano: l'evidenza che le disuguaglianze non costituiscano

dati naturali o contingenze accidentali, bensì il prodotto di specifiche strutture sociali storicamente determinate e, di conseguenza, modificabili.

1.3.2 Weber: la triade classe-status-partito

Max Weber (1864-1920) costruisce la propria teoria della stratificazione partendo dalle categorie marxiane, di cui contesta però l'impianto unidimensionale. Sebbene la sfera economica mantenga una rilevanza centrale, non basta a spiegare per Weber la totalità dei fenomeni descritti. La sua proposta teorica si articola di conseguenza su tre assi analitici distinti: la *classe*, il *gruppo di status* (o ceto) e il *partito* (Croteau & Hoynes, 2022, pp. 322-324). Queste tre dimensioni presentano un elevato grado di autonomia reciproca, consentendo traiettorie di variazione indipendenti: un individuo può occupare una posizione apicale su un asse e trovarsi in una condizione di marginalità su un altro. Un imprenditore di successo (elevata posizione di classe) può risultare privo di prestigio sociale (basso *status*); per contro, un intellettuale accreditato (alto *status*) può disporre di risorse materiali decisamente esigue (bassa classe). Questo potenziale *disallineamento multidimensionale* rappresenta uno dei contributi più originali e fecondi offerti da Weber alla sociologia delle diseguaglianze. Collins (1992, p. 192) riprende la definizione weberiana di classe come *situazione di mercato*, precisando che la posizione di classe di un individuo dipende dalle condizioni in cui partecipa allo scambio economico. Weber distingue tre tipi di mercato che determinano altrettante linee di stratificazione economica (Collins, 1992, p. 193):

- il *mercato del lavoro*, che separa chi vende la propria forza lavoro da chi la compra;
- il *mercato del credito*, che divide debitori e creditori;
- il *mercato delle merci*, che oppone acquirenti e venditori di beni.

La pluralizzazione delle linee di stratificazione economica è un primo allargamento rispetto allo schema marxiano, dove la divisione fondamentale era una sola, quella tra possessori e non possessori dei mezzi di produzione. La seconda dimensione, il *gruppo di status*, riguarda il prestigio sociale e gli stili di vita. Per Weber un *gruppo di status* si fonda su una comunità di onore, di consumi e di pratiche culturali che marca una distanza simbolica dagli altri gruppi (Croteau & Hoynes, 2022, p. 323). Il gruppo etnico, ad esempio, può attraversare diverse classi economiche pur mantenendo una propria identità di *status*. La terza dimensione è il *partito*, che Collins riprende dalla definizione weberiana: i partiti «vivono nella "casa del potere", egli vuol dire che essi sono gruppi specificamente politici, fazioni organizzate per lottare per il potere» (Collins, 1992, p. 195). Sotto la categoria del *partito* rientrano dunque tutti i raggruppamenti che si organizzano in vista di una finalità politica, indipendentemente dalla loro base sociale. Il contributo che lega le tre dimensioni weberiane è il concetto di *chance di vita*, definito da Croteau e Hoynes come «possibilità di accedere a risorse economiche e culturali apprezzate» (Croteau & Hoynes, 2022, p. 324). La posizione che un individuo occupa nelle tre dimensioni della stratificazione determina le sue *chance di vita*, ovvero la probabilità statistica di accedere a beni materiali, opportunità professionali, riconoscimento sociale, salute e istruzione. Non si tratta di un determinismo: le *chance* non sono certezze, ma probabilità differenziate. Il concetto è importante perché traduce in termini analitici il legame tra struttura sociale e biografia individuale: una persona nasce in una determinata configurazione di *chance*, e quella configurazione condiziona, senza determinarlo meccanicamente, il corso della sua vita.

1.3.3 Il confronto: convergenze e divergenze

Marx e Weber convergono su un punto fondamentale, segnalato in modo esplicito da Croteau e Hoynes: «Per Marx e Weber, la disuguaglianza tra classi era strettamente interconnessa con le lotte per la conquista del potere all'interno della società» (Croteau & Hoynes, 2022, p. 325). Per entrambi, in altri termini, la disuguaglianza non è un fenomeno isolato che può essere analizzato in sé: è connessa a dinamiche di potere e di conflitto, e va letta come parte di una configurazione più ampia. Su questo punto i due autori si distaccano nettamente dalle letture funzionaliste che, soprattutto a metà Novecento, avrebbero presentato la stratificazione come un meccanismo neutrale di allocazione delle risorse, finalizzato al funzionamento efficiente della società. Per Marx e Weber l'ottica è diametralmente opposta: la stratificazione genera di per sé conflitto e sta in piedi solo perché i gruppi al potere riescono a blindare le proprie posizioni.

Poi le strade si dividono, in particolare su un punto: il numero dei fattori in gioco. La lettura marxiana schiaccia la disparità su un unico asse fondante, la dimensione economico-produttiva, trattando qualsiasi altra asimmetria come una semplice conseguenza. Weber fa un'operazione completamente diversa. Riconosce una pluralità di dimensioni e, aspetto ancora più importante, ammette esplicitamente che queste sfere possano slegarsi l'una dall'altra e risultare del tutto disallineate. La seconda discrepanza teorica investe la concezione stessa di classe. Collins (1992, p. 194) richiama in proposito una distinzione di origine marxiana che Weber rielabora e che segnerà la sociologia successiva: quella tra *classe an sich* (classe in sé) e *classe für sich* (classe per sé). La prima dimensione è di natura oggettiva e ricalca la collocazione dell'individuo all'interno dei rapporti di produzione; la seconda ha invece carattere soggettivo e si riferisce alla maturazione di una coscienza d'identità che rende il gruppo capace di mobilitazione e azione collettiva. Se per Marx la transizione

verso la *classe per sé* costituiva uno sbocco storico pressoché inevitabile per il proletariato, per Weber e per il filone weberiano questo passaggio rimane problematico. Esso è subordinato a variabili culturali, organizzative e contingenti che possono agevolarlo oppure bloccarlo.

Da queste premesse prendono forma due percorsi di ricerca separati. La linea marxiana alimenta una tradizione focalizzata sulle trasformazioni del lavoro, sul conflitto strutturale e sulle dinamiche dello sfruttamento capitalistico. La tradizione weberiana sposta invece l'obiettivo sulla natura multidimensionale del problema. Qui l'indagine si concentra sul peso dello *status*, sul prestigio e su quelle asimmetrie di potere che viaggiano su binari indipendenti rispetto alla pura logica economica. È proprio dall'attrito e dal dialogo continuo tra queste due impostazioni che, lungo tutto il Novecento, hanno preso forma i modelli neo-marxiani e neo-weberiani su cui si regge l'attuale sociologia della stratificazione. Il punto è che non ci troviamo di fronte a una banale disputa teorica. Le conseguenze atterrano direttamente sul metodo: è a partire da questo bivio che la ricerca empirica decide quali variabili andare a misurare sul campo, in che modo isolare i gruppi sociali e come fabbricare, all'atto pratico, i propri strumenti di analisi.

1.4 La stratificazione nella sociologia contemporanea

Le coordinate fissate da Marx e Weber sono state riprese, articolate e in parte trasformate dalla sociologia novecentesca, che ha tentato di tradurre le categorie classiche in strumenti capaci di analizzare le società complesse del secondo dopoguerra. Il paragrafo presenta tre snodi di questo lavoro: gli sviluppi neo-marxiani e neo-weberiani che hanno operativizzato il concetto di classe per la ricerca empirica, l'analisi della mobilità sociale come indicatore del grado di apertura di una società, la teoria della frammentazione che ha messo in discussione la centralità

stessa della categoria di classe nelle società post-industriali. Tre prospettive che, pur con accenti diversi, condividono un'eredità comune e contribuiscono a definire il quadro entro cui collocare il dibattito contemporaneo.

1.4.1 Le due eredità: neo-marxismo e neo-weberismo

Nel corso del Novecento, il dibattito Marx-Weber si è prolungato in due famiglie di ricerca distinte ma in costante dialogo. La prima, di matrice neo-marxiana, ha cercato di ripensare il concetto di classe per renderlo operativo nelle società capitalistiche avanzate, dove la struttura binaria borghesia-proletariato del marxismo classico appariva insufficiente a cogliere la complessità della stratificazione. La seconda, di matrice neo-weberiana, ha lavorato per costruire mappe della stratificazione utilizzabili nella ricerca empirica, articolando le dimensioni della classe attorno a indicatori misurabili. Croteau e Hoynes presentano entrambe le tradizioni come parte integrante della cassetta degli attrezzi della sociologia contemporanea, sottolineando come ciascuna offra strumenti adeguati ad analizzare aspetti diversi della stessa realtà sociale.

L'autore più rappresentativo dell'approccio neo-marxiano è Erik Olin Wright, il quale nel 1985 ha proposto una rielaborazione del concetto marxiano di classe centrata sul concetto di *controllo* (Croteau & Hoynes, 2022, pp. 334-335). Wright distingue tre forme di controllo che insieme definiscono la posizione di classe: il *controllo sui mezzi fisici di produzione*, ovvero fabbriche, uffici, terreni; il *controllo sui mezzi monetari*, ovvero capitale e finanziamenti; il *controllo sulla forza lavoro altrui*. La combinazione di queste tre forme produce un quadro più articolato di quello marxiano classico, in cui trovano spazio le *classi contraddittorie*: posizioni intermedie come quelle dei quadri e dei manager, che esercitano un controllo parziale e si trovano in una

condizione ambigua, simultaneamente sopra il proletariato e sotto la borghesia. Wright si inserisce in una linea di pensiero che dialoga con Gramsci (2014) e Poulantzas (1978), citati dal manuale come riferimenti dell'approccio neo-marxiano alla stratificazione.

L'approccio neo-weberiano, sviluppato in particolare da John Goldthorpe nei lavori del 1983 e del 1987, articola la stratificazione su tre dimensioni: la *situazione di mercato*, che riguarda le condizioni di acquisto e vendita della forza lavoro; la *situazione di lavoro*, che riguarda il controllo gerarchico esercitato o subito nei processi produttivi; lo *status di consumo*, che riguarda gli stili di vita e le pratiche culturali che marcano la posizione sociale (Croteau & Hoynes, 2022, pp. 335-336). Su questa base Goldthorpe ha costruito una mappa empirica delle società contemporanee articolata in *sette classi sociali* che è divenuta uno strumento standard nelle indagini comparative sulla stratificazione in Europa. La forza dell'impostazione neo-weberiana sta nella sua operatività: le tre dimensioni si traducono in indicatori misurabili nelle survey, e le sette classi possono essere applicate a contesti nazionali diversi per produrre confronti sistematici.

1.4.2 La mobilità sociale: aperture e chiusure

Una società non si descrive soltanto guardando alla sua stratificazione in un dato momento, ma anche analizzando il movimento degli individui tra le posizioni che la compongono. Questo movimento è ciò che la sociologia chiama *mobilità sociale*. La distinzione fondamentale è tra mobilità *intergenerazionale*, che confronta la posizione di un individuo con quella dei suoi genitori, e mobilità *intragenerazionale*, che riguarda i cambiamenti nel corso della vita di una stessa persona. A questa si aggiunge la distinzione tra mobilità *ascendente* e *discendente*, a seconda che il movimento porti verso posizioni più elevate o più basse della

struttura sociale. Misurare la mobilità di una società significa misurarne, in larga parte, il grado di apertura.

Il Novecento ha conosciuto, nelle società occidentali, una fase di mobilità ascendente particolarmente intensa. Croteau e Hoynes richiamano l'analisi di Bauman e Schizzerotto sul caso italiano (Croteau & Hoynes, 2022, p. 327): con la transizione verso il modello industriale, masse enormi di contadini e braccianti sono state materialmente riconvertite, finendo per essere assorbite nelle fabbriche o nel lavoro autonomo. Da lì, moltissimi dei loro figli hanno poi fatto il salto verso i ceti medi impiegatizi. Il punto centrale è che questo scivolamento verso l'alto non dipendeva affatto dai talenti o dalle strategie di pianificazione delle singole famiglie. A trainare il processo era la modernizzazione dello scheletro economico nel suo complesso. L'onda dell'industria, il boom dei servizi e le porte aperte della scuola pubblica hanno funzionato come un gigantesco ascensore collettivo. Si generavano fisicamente posizioni inedite nella struttura occupazionale, e questo ha permesso a fette larghissime della popolazione di sperimentare una vera e propria *mobilità ascendente*.

La situazione contemporanea presenta un quadro diverso, almeno per quanto riguarda il caso italiano. Croteau e Hoynes, riprendendo Schizzerotto (2012), segnalano che le nuove generazioni, pur dotate di livelli di istruzione più elevati rispetto ai genitori, non riescono spesso a raggiungere posizioni lavorative stabili, e si trovano in molti casi in una condizione sociale peggiore di quella della generazione precedente (Croteau & Hoynes, 2022, p. 327). Il dato sociologico che emerge è cruciale. L'automatismo tipico del dopoguerra, per cui accumulare anni di studio si traduceva dritto in una sistemazione lavorativa migliore, si è ormai inceppato. La grande spinta della *mobilità ascendente* di massa si è esaurita, lasciando il posto, in sempre più casi, a veri e propri scivolamenti verso il basso da una generazione all'altra. La sociologia italiana ha

inquadrato questo cortocircuito parlando di *immobilità* delle società *formalmente aperte*. Il paradosso è evidente: sulla carta l'ascensore sociale è a disposizione di tutti, ma all'atto pratico le occasioni di salita finiscono per restare quasi sempre appannaggio di chi parte già da posizioni di forte vantaggio.

1.4.3 La teoria della frammentazione

Un terzo filone della sociologia contemporanea sostiene che la categoria di classe ha progressivamente perso consistenza nelle società post-industriali (Croteau & Hoynes, 2022, pp. 336-337). La tesi è stata articolata, con accenti diversi, da una serie di autori che il manuale richiama: Daniel Bell, Alain Touraine, Richard Sennett, Ulrich Beck. L'argomento di fondo è che l'organizzazione del lavoro post-fordista, la flessibilità delle carriere, l'individualizzazione dei percorsi biografici, la pluralizzazione degli stili di consumo abbiano eroso le grandi appartenenze collettive che strutturavano la società industriale. Le identità sociali, in questa lettura, si frammentano in posizioni più mobili e indistinte, in cui la classe come grande categoria condivisa cede il passo a una molteplicità di stili di vita, identità multiple, traiettorie personalizzate.

La tesi della frammentazione coglie una trasformazione reale, ma rischia di essere fraintesa se interpretata come tesi della scomparsa delle diseguaglianze. I dati sulla distribuzione della ricchezza, sulla mobilità sociale, sull'accesso alle risorse continuano a documentare disparità sistematiche e in alcuni casi crescenti, come gli indicatori globali richiamati dal manuale stesso confermano. La gerarchia sociale non è sparita. A cambiare è la sua manifestazione nella coscienza delle persone. Assistiamo oggi a una profonda frammentazione delle identità. Le posizioni oggettive restano però del tutto bloccate. Chi nasce in un contesto di vantaggio conserverà il proprio status con estrema facilità. La

differenza vera risiede nel vissuto interiore. Il singolo interpreta i traguardi raggiunti come tappe di un percorso strettamente personale. Scompare il senso di appartenenza a una classe definita. L'analisi sociologica odierna parte proprio da questo presupposto. Il processo di *individualizzazione* non azzerava le disuguaglianze. Ne attesta piuttosto una profonda trasformazione. I meccanismi di riproduzione del privilegio lavorano ancora. Agiscono solo in forme molto meno visibili.

Le tre tradizioni discusse in questo paragrafo, pur muovendo da prospettive diverse, convergono nel mostrare che la stratificazione contemporanea non si lascia ridurre a una sola dimensione. La disuguaglianza si articola su molteplici linee, alcune economiche, altre culturali, altre relazionali, e va analizzata con strumenti capaci di coglierne la complessità. Resta tuttavia una dimensione che il pensiero classico aveva lasciato in ombra e che la sociologia contemporanea ha portato al centro del dibattito: il fatto che le disuguaglianze non sono mai soltanto di classe, ma si intrecciano con il genere, con l'etnia, con la collocazione geografica, in configurazioni che producono forme di disparità non riducibili alla somma delle loro componenti. Il prossimo paragrafo affronta questa pluralità delle linee di disuguaglianza, soffermandosi sulle due dimensioni che la sociologia contemporanea ha posto accanto alla classe: il genere e l'etnia.

1.5 Le disuguaglianze oltre la classe: genere ed etnia

Concentrare l'analisi sulla sola classe lascia in ombra una parte rilevante delle disparità che attraversano le società contemporanee: il genere e l'origine etnica producono asimmetrie che non si lasciano ridurre alla posizione economica, e che hanno una storia, una struttura, una logica propria. Croteau e Hoynes le trattano in modo essenziale, collocandole nel quadro più ampio dei sistemi di stratificazione storici. Collins, invece, le

affronta in modo più analitico, dedicando una sezione del proprio capitolo sulla stratificazione alla questione del genere.

1.5.1 La stratificazione di genere

Croteau e Hoynes collocano la stratificazione di genere all'interno dell'evoluzione storica dei sistemi di disuguaglianza premoderni, in parallelo a modelli quali la schiavitù, i sistemi castali e gli ordini feudali. Il *patriarcato* viene qui inteso come un «sistema di stratificazione basato sul primato assoluto del pater familias sugli altri membri della comunità» (Croteau & Hoynes, 2022, p. 317). In questa configurazione la subordinazione delle donne è sistematica, ed è legittimata e riprodotta nel tempo da norme, istituzioni e divisioni del lavoro rigide. Anche se nelle società contemporanee le forme classiche di patriarcato sono state formalmente superate dall'uguaglianza giuridica, alcune di quelle asimmetrie persistono in forme nuove. Tra queste, gli autori menzionano il cosiddetto *tetto di cristallo*, termine che indica la barriera invisibile che ostacola l'ascesa delle donne verso le posizioni apicali delle organizzazioni, anche a fronte di percorsi di carriera formalmente accessibili (Croteau & Hoynes, 2022, p. 317).

L'analisi proposta da Collins (1992) adotta un impianto diverso, classificando la riflessione sulla stratificazione di genere all'interno di tre principali filoni teorici. Il primo, di matrice neo-freudiana, fa capo a Nancy Chodorow (1978) e rintraccia l'origine dei diversi ruoli sociali nei meccanismi psicologici della socializzazione infantile. L'assunto di base è che l'allevamento della prole, storicamente affidato alle madri, generi esiti divergenti nei due sessi: mentre le figlie tendono a sviluppare un processo di identificazione continua con la figura materna, i figli maschi costruiscono la propria identità per differenziazione e distacco (Collins, 1992, p. 208). Il secondo approccio, di natura economico-marxiana, si rifà

al testo di Friedrich Engels del 1884, *L'origine della famiglia, della proprietà privata e dello Stato*. In quest'opera, la subordinazione femminile viene posta in relazione diretta con l'istituzione della proprietà privata e con la transizione dalle antiche società matrilineari a quelle patrilineari (Collins, 1992, p. 209). Il terzo orientamento coincide con la teoria di Rae Lesser Blumberg, secondo cui la stratificazione di genere dipende dall'andamento di alcune variabili macro-sociali, prime fra tutte il potere economico detenuto dalle donne, l'intensità della loro partecipazione al mercato del lavoro e il livello di controllo esercitato sulle risorse domestiche e sulla sfera riproduttiva (Collins, 1992, p. 211). A questo quadro Collins integra un'ipotesi personale che individua nella *politica sessuale* una dimensione cruciale per comprendere i rapporti di potere all'interno delle società premoderne (Collins, 1992, pp. 214-219). Lo studioso rintraccia lo snodo fondamentale verso gli assetti attuali nel modello culturale vittoriano del XIX secolo. In questa fase, il mutamento dei livelli di indipendenza economica femminile diventa il fattore principale di ridefinizione dello status delle donne (Collins, 1992, p. 219). L'ingresso graduale nel mercato del lavoro salariato e la disponibilità di redditi autonomi hanno alterato i tradizionali rapporti di forza sia all'interno delle mura domestiche sia nello spazio pubblico, pur senza azzerare del tutto le asimmetrie di fondo. Le diseguaglianze di genere nelle società contemporanee, suggerisce questa prospettiva, sono il prodotto dell'intersezione tra strutture economiche, eredità culturali e dinamiche di potere che non si dissolvono per il solo fatto che la legge ha riconosciuto la parità formale.

1.5.2 La stratificazione etnico-razziale

Le diseguaglianze etniche e razziali ricevono nei due manuali un trattamento meno sistematico di quello riservato al genere, e per buona

parte vengono presentate come dimensione storica delle stratificazioni premoderne. Croteau e Hoynes le collocano nel paragrafo dedicato ai sistemi storici di stratificazione, in particolare nella discussione della schiavitù, definita come «forma estrema di disuguaglianza, per cui alcuni individui sono oggetto di proprietà di altri (uomini liberi) e quindi privati di fatto e di diritto, di ogni autonomia personale» (Croteau & Hoynes, 2022, p. 316). La schiavitù è il caso limite di una stratificazione fondata su caratteristiche ascritte, in cui l'appartenenza etnica diventa il fondamento di un sistema di subordinazione formalizzato.

Sul piano teorico, il riferimento più significativo che i due manuali offrono è la trattazione weberiana del gruppo di status, ripresa da Collins. Il gruppo etnico, scrive Collins (1992, p. 194) riassumendo Weber, è una delle forme tipiche di gruppo di status: una comunità che condivide stili di vita, pratiche, simboli, e che può attraversare diverse classi economiche pur mantenendo una propria identità. Questa lettura ha un'implicazione importante. La stratificazione etnica non è riducibile alla stratificazione economica, perché un gruppo etnico può occupare posizioni di classe diverse al proprio interno e tuttavia continuare a essere riconosciuto, e a riconoscersi, come comunità distinta. Le disuguaglianze etnico-razziali si articolano dunque su una dimensione propria, che si aggiunge alla classe senza coincidere con essa.

Genere ed etnia costituiscono dunque due dimensioni della stratificazione che si aggiungono alla classe senza coincidere con essa. Oggi la letteratura sociologica tende a incrociare queste dimensioni. Il posizionamento di un individuo nella gerarchia è di fatto il risultato di questa sovrapposizione. C'è però un problema strutturale, un aspetto che i classici avevano affrontato in modo piuttosto marginale. Si tratta della trasmissione delle disuguaglianze da una generazione all'altra. Il punto critico è capire come riescano a perpetuarsi nel tempo persino in società sulla carta meritocratiche. Contesti in cui le traiettorie di vita dovrebbero

teoricamente dipendere soltanto dal talento personale. Il paragrafo successivo affronta queste domande attraverso il contributo di Pierre Bourdieu sul capitale culturale e l'analisi che Collins dedica alle credenziali educative. È un passaggio che chiude la ricostruzione sociologica del capitolo e apre verso le diseguaglianze digitali trattate nel Capitolo 2.

1.6 La riproduzione delle diseguaglianze: dal capitale culturale alle credenziali educative

Le teorie discusse nei paragrafi precedenti hanno mostrato come la stratificazione si articoli su dimensioni molteplici, ciascuna con una propria logica. Il nodo centrale della questione, in realtà, riguarda proprio la trasmissione materiale di queste gerarchie. Assumendo di muoverci in sistemi che tutelano l'uguaglianza formale e promettono libero accesso all'istruzione, la tenuta intergenerazionale del privilegio sfugge alle logiche del semplice travaso di ricchezza economica. C'è inevitabilmente dell'altro sotto la superficie. Per smontare un simile meccanismo, l'indagine si appoggia in gran parte all'intelaiatura teorica fornita da due studiosi. Pierre Bourdieu porta in dote la categoria di *capitale culturale*, arrivando a smascherare la scuola stessa come un vero e proprio apparato di riproduzione dei divari. Su un binario parallelo si muove Randall Collins, che ancora la persistenza di quelle stesse asimmetrie al peso esercitato dalle credenziali educative. Le due prospettive convergono nel mostrare che il sistema formativo, lungi dall'essere un dispositivo meritocratico neutrale, è uno dei meccanismi attraverso cui le disparità si rinnovano da una generazione all'altra.

1.6.1 Bourdieu e il capitale culturale

Pierre Bourdieu è uno degli autori che hanno contribuito in modo più articolato all'analisi della riproduzione delle disuguaglianze nelle società contemporanee. Collins lo richiama insieme a Jean-Claude Passeron in riferimento al lavoro del 1970, dedicato alla scuola come istituzione che riproduce le disuguaglianze attraverso il *capitale culturale* delle classi dominanti (Collins, 1992, p. 228). Il perno del ragionamento ribalta l'illusione della neutralità didattica. Più che agire come un setaccio neutro del talento, la macchina scolastica tende fatalmente a ricompensare chi si allinea ai blocchi di partenza avendo già piena padronanza dei codici richiesti. Crescere all'interno di un tessuto familiare denso di *capitale culturale* significa assorbire, quasi per osmosi, l'abitudine alla lettura, una sintassi elaborata e la consuetudine con i consumi artistici. Di conseguenza, questi soggetti varcano la soglia della classe disponendo esattamente di quel corredo di grammatiche implicite che l'istituzione è programmata per valutare e certificare come eccellente. Chi parte da famiglie con minori risorse culturali deve invece acquisire a scuola ciò che gli altri possiedono già, e affronta un percorso strutturalmente più lungo e più incerto.

Il concetto di capitale culturale trova un parallelo significativo nella tassonomia delle risorse proposta da Croteau e Hoynes, dove le *risorse culturali* sono definite come l'insieme delle «conoscenze implicite e abilità informali» che gli individui acquisiscono attraverso la socializzazione (Croteau & Hoynes, 2022, p. 314). Sono risorse che non si misurano in reddito o in patrimonio, ma che agiscono concretamente sulle traiettorie individuali. Una persona che sappia muoversi nei contesti formali, decifrare un documento amministrativo, riconoscere i codici di un ambiente professionale dispone di una leva che altri non possiedono, anche a parità di reddito o di titolo di studio formale. La famiglia, ancor prima della scuola, è il primo luogo in cui queste risorse vengono

trasmesse, e proprio per questo le diseguaglianze culturali tendono a mantenersi anche in società che hanno eliminato le barriere formali all'accesso all'istruzione.

1.6.2 Le credenziali educative e l'inflazione del titolo

L'analisi di Collins sulle credenziali educative parte da una tesi che mette in discussione una delle convinzioni più diffuse sull'istruzione nelle società contemporanee. «La forma più ovvia di stratificazione nelle società industriali della fine del Ventesimo secolo è quella secondo l'istruzione» (Collins, 1992, p. 220): la posizione sociale è in larga parte determinata dal livello di titolo di studio conseguito. Il senso comune legge la questione in un'ottica puramente meritocratica. La scuola, in pratica, selezionerebbe il talento naturale. Il mercato occupazionale andrebbe poi a ricompensare i profili migliori. Collins rifiuta in blocco questa prospettiva, il nesso tra il possesso di un titolo accademico e le reali capacità della persona risulta assai debole. Sicuramente molto più labile di quanto l'ideologia del merito sia disposta ad ammettere (Collins, 1992, pp. 220-223).

Lo snodo teorico di questa critica risiede nella distinzione tra *requisiti di competenze* e *requisiti di credenziali* (Collins, 1992, p. 223). I primi identificano le abilità pratiche e tecniche indispensabili per l'esecuzione di una mansione; i secondi corrispondono ai titoli formali pretesi dalle aziende in fase di selezione. Secondo lo studioso, queste due sfere divergono. I dati empirici indicano infatti che i datori di lavoro tendono a esigere livelli di istruzione sovrastimati rispetto alle reali attività da svolgere, e che la correlazione tra scolarizzazione e produttività sul posto di lavoro si rivela debole. Il titolo di studio, di conseguenza, agisce meno come certificazione di un saper fare e più come un indicatore di *status*, ovvero come un contrassegno simbolico che inserisce l'individuo

in una specifica cerchia sociale. Questa premessa conduce Collins a formulare la tesi della *inflazione delle credenziali educative* (Collins, 1992, pp. 226-228). Nel corso del XX secolo, la crescita dei tassi di scolarizzazione nelle società industriali non è stata trainata da una reale necessità di competenze tecniche sempre più complesse. La spinta principale è derivata dalla competizione per lo *status*: l'acquisizione di titoli superiori è divenuta una strategia per differenziarsi dalla concorrenza in contesti occupazionali saturi. L'esito di questa dinamica è una svalutazione monetaria dei titoli: la credenziale richiesta per l'accesso a una determinata posizione lavorativa si è innalzata nel tempo, senza che a ciò corrispondesse una reale trasformazione delle mansioni. Richiamando esplicitamente le tesi di Bourdieu e Passeron (1970), Collins interpreta questo fenomeno come un dispositivo di riproduzione delle disparità: l'istituzione scolastica elargisce titoli in modo proporzionale al capitale culturale d'origine, permettendo così alle famiglie più dotate di risorse di tramandare ai figli un vantaggio competitivo che si converte in credenziali più elevate e, successivamente, in collocazioni occupazionali migliori (Collins, 1992, p. 228). L'argomentazione di Collins tocca qui un punto centrale per l'impianto della tesi. Analizzando l'ipertrofia del sistema formativo come surrogato delle politiche keynesiane contro la disoccupazione, l'autore osserva che i sistemi industriali avanzati, costretti a fare i conti con la contrazione del lavoro manifatturiero dovuta all'automazione, hanno sfruttato l'allungamento dei percorsi scolastici come ammortizzatore sociale per assorbire la manodopera in eccedenza (Collins, 1992, pp. 229-231). Già all'inizio degli anni Novanta, Collins anticipava una transizione oggi al centro della riflessione sociologica: «I robot stanno soppiantando il lavoro di fabbrica e altri lavori manuali; e l'intelligenza artificiale sta iniziando a impattare il lavoro impiegatizio. Poiché quest'ultimo è l'ambito principale dell'occupazione femminile possiamo anticipare che vi sarà una crisi crescente nella posizione

economica delle donne verso la fine del ventesimo secolo» (Collins, 1992, p. 230). La citazione, scritta più di trent'anni fa, anticipa con precisione il tema che attraversa l'intera tesi: la diffusione dell'intelligenza artificiale come fattore di trasformazione strutturale del lavoro, con effetti differenziati sui diversi segmenti della forza lavoro.

Le due prospettive discusse, quella di Bourdieu sul capitale culturale e quella di Collins sulle credenziali educative, convergono nel mostrare che le diseguaglianze nelle società contemporanee non si riducono alla trasmissione diretta di patrimonio economico. Esse si riproducono attraverso meccanismi culturali e istituzionali che operano in modo meno visibile ma altrettanto efficace, e che agiscono sul piano delle risorse simboliche, delle competenze, delle credenziali. È un'eredità teorica che diventa fondamentale per affrontare il tema del Capitolo 2. Le diseguaglianze digitali, come si vedrà, non si articolano principalmente sulla disponibilità materiale dei dispositivi: si articolano sulle competenze, sulle motivazioni, sugli usi che le persone fanno della tecnologia, e dunque sulle risorse culturali ereditate prima ancora che sull'accesso alla rete. Il modello multidimensionale di van Dijk, che il Capitolo 2 ricostruisce, raccoglie e rielabora proprio l'intuizione bourdieusiana e collinsiana: che le diseguaglianze più resistenti sono quelle che passano attraverso la cultura e la formazione, e che si trasmettono prima ancora di manifestarsi come differenze osservabili nei comportamenti.

1.7 Dalla stratificazione sociale alla soglia del digitale

Il percorso del capitolo è partito dall'impostazione sociologica del problema delle diseguaglianze, lavorando sulla distinzione tra differenza e disparità e sul carattere strutturato di quest'ultima. La definizione di Croteau e Hoynes ha tracciato il perimetro iniziale. A questa si sono aggiunte la classificazione delle sette risorse e la triade struttura, cultura,

potere. Questo impianto concettuale ha guidato l'intera ricostruzione. L'analisi è partita dalle riflessioni classiche di Marx e Weber, per approdare poi ai modelli operativi di Wright e Goldthorpe. Una transizione decisiva che dimostra come la ricerca attuale abbia tradotto categorie antiche in strumenti empirici concreti. Un passaggio ormai obbligato per decifrare le società avanzate che fa emergere sul fronte della mobilità sociale un dato netto, specialmente per l'Italia: la grande spinta ascendente del dopoguerra si è fermata. Oggi le nuove generazioni registrano un arretramento oggettivo rispetto ai genitori, un paradosso evidente, considerando i loro titoli di studio superiori.

L'orizzonte si è infine allargato oltre la pura variabile di classe, il genere e le radici etniche creano stratificazioni ulteriori. Forme di svantaggio che si sommano alla disuguaglianza economica senza però mai coincidere del tutto con essa. Il patriarcato, le teorie di Chodorow, Engels, Blumberg e Collins sulla politica sessuale, la trattazione weberiana del gruppo di status etnico hanno mostrato che la posizione sociale di un individuo dipende dall'intreccio di più dimensioni. L'ultimo passaggio ha riguardato la riproduzione delle disuguaglianze. Bourdieu, richiamato attraverso il lavoro del 1970 con Passeron sulla scuola come istituzione di riproduzione, e Collins, con la teoria delle credenziali educative e la distinzione tra requisiti di competenze e requisiti di credenziali, hanno offerto due prospettive convergenti su un punto fondamentale: nelle società contemporanee le disparità si trasmettono attraverso meccanismi culturali e istituzionali che agiscono prima ancora della distribuzione del reddito, e che operano in modo meno visibile della trasmissione diretta del patrimonio.

Le categorie elaborate dalla sociologia non costituiscono un patrimonio meramente storico, sono strumenti analitici che restano operativi per leggere le forme contemporanee della disparità, comprese quelle che la diffusione delle tecnologie digitali e dell'intelligenza

artificiale sta producendo. Il capitale culturale ereditato nella famiglia di origine, le competenze acquisite attraverso il sistema formativo, le credenziali che certificano una posizione, lo status che colloca gli individui in gerarchie simboliche oltre che economiche: sono dimensioni che ritornano, sotto forme nuove, anche quando l'oggetto dell'analisi si sposta dall'industria classica alle piattaforme digitali, dal lavoro manuale all'interazione con sistemi automatizzati. Le diseguaglianze prodotte dalla transizione algoritmica, come si vedrà nel capitolo successivo, si innestano su questa stratificazione preesistente e ne riproducono, in larga parte, le linee di frattura.

Un riferimento contenuto nelle pagine di Collins risulta particolarmente significativo come passaggio verso il prossimo capitolo. Discutendo l'espansione del sistema educativo come sostituto della disoccupazione strutturale, Collins anticipa nel 1992 una trasformazione che oggi è al centro del dibattito: «I robot stanno soppiantando il lavoro di fabbrica e altri lavori manuali; e l'intelligenza artificiale sta iniziando a impattare il lavoro impiegatizio» (Collins, 1992, p. 230). La previsione, formulata oltre trent'anni fa, mostra una continuità teorica forte tra la sociologia delle credenziali educative e le questioni che la transizione algoritmica solleva oggi. Il Capitolo 2 si occupa precisamente di questo passaggio, ricostruisce la genealogia delle diseguaglianze digitali, dalle prime asimmetrie nell'accesso fisico ai dispositivi fino all'emergere dell'AI divide, e mostra come il modello multidimensionale elaborato da Jan van Dijk raccolga e prolunghi l'eredità sociologica fin qui delineata, applicandola al terreno specifico della diffusione delle tecnologie digitali nella società contemporanea.

Capitolo 2 Dal digital divide all'AI divide

2.1 Introduzione

Il quadro sociologico ricostruito nel capitolo precedente ha messo a disposizione un insieme di categorie che restano utili per leggere le diseguaglianze contemporanee: il capitale culturale e le credenziali educative, ma anche il modo in cui le risorse si distribuiscono in maniera diseguale prima ancora di tradursi in esiti economici, e i meccanismi attraverso cui le disparità si trasmettono da una generazione all'altra. Le pagine che seguono applicano queste categorie a un terreno specifico, quello del rapporto tra le persone e le tecnologie digitali, e in particolare a un fenomeno che la ricerca sociologica ha cominciato a studiare a partire dalla seconda metà degli anni Novanta e che continua a riconfigurarsi a ogni nuova ondata tecnologica.

L'ipotesi che attraversa il capitolo può essere formulata in termini diretti: l'AI divide non rappresenta un fenomeno inedito, bensì l'ultimo stadio di una stratificazione progressiva che la ricerca sociologica ha cominciato a studiare a partire dalla seconda metà degli anni Novanta. Per mostrarlo si procede in quattro passaggi. Il paragrafo 2.2 ricostruisce la nascita del concetto di *digital divide* e i problemi interpretativi che la sua metafora originaria porta con sé. Si passa poi (paragrafo 2.3) al modello teorico multidimensionale elaborato da Jan van Dijk, riferimento più articolato per analizzare le diseguaglianze digitali nella loro complessità. La stratificazione storica della ricerca, organizzata dalla letteratura in tre livelli successivi (accesso, competenze, benefici concreti), è oggetto del paragrafo 2.4. Il passaggio alla dimensione algoritmica (paragrafo 2.5) e l'emergere dell'AI divide in senso stretto (paragrafo 2.6) chiudono il discorso prima delle conclusioni (paragrafo 2.7), che sintetizzano il

percorso e collocano l'Italia nel quadro europeo attraverso i dati più recenti della Commissione Europea.

Va precisato fin da subito che i diversi livelli del divario che verranno discussi non vanno intesi come fasi che si escludono a vicenda: l'emergere del dibattito sulle competenze digitali, negli anni Duemila, non ha coinciso con il superamento del problema dell'accesso fisico, così come l'attuale discussione sull'alfabetizzazione algoritmica coesiste con forme residue di esclusione che ancora riguardano segmenti consistenti della popolazione. Le fasi si sono accumulate, non sostituite. Questa compresenza rende necessario un approccio genealogico. Soltanto risalendo ai primi studi sul divario diventa possibile cogliere perché l'AI divide, come si proverà ad argomentare, sia l'esito coerente di una traiettoria di lungo periodo, e non una discontinuità teorica rispetto alle dinamiche di stratificazione sociale già ricostruite nel Capitolo 1.

2.2 Il digital divide: origine del concetto e critica della metafora

L'espressione *digital divide* entra nel lessico delle scienze sociali negli Stati Uniti a metà degli anni Novanta. Il suo atto di nascita istituzionale coincide con la pubblicazione del rapporto *Falling through the Net*, curato nel 1995 dalla National Telecommunications and Information Administration statunitense, nel quale veniva per la prima volta tematizzata una frattura tra *haves* e *have-nots* nell'accesso alle tecnologie dell'informazione (van Dijk, 2020, p. 1). La formulazione iniziale era semplice: chi possedeva un computer e una connessione a Internet si trovava da un lato della frattura, chi non li possedeva dall'altro. Nei primi anni di diffusione del personal computer e del web questa impostazione binaria appariva adeguata a rappresentare una realtà in cui,

effettivamente, la disponibilità materiale dello strumento costituiva il discrimine più evidente.

Nei decenni successivi la ricerca ha progressivamente complicato il quadro. Jan van Dijk, sociologo olandese che ha dedicato al tema l'intera carriera accademica, definisce oggi il digital divide come «a division between people who have access and use of digital media and those who do not» (van Dijk, 2020, p. 1), ma precisa immediatamente che la parola *access*, enfaticata nei primi anni di studio, ha progressivamente ceduto il passo alla parola *use*: una volta garantita la connessione, il problema si sposta su cosa le persone facciano effettivamente con gli strumenti di cui dispongono. Il termine stesso, peraltro, non è stato l'unico tentativo di nominare il fenomeno. La letteratura sociologica aveva già elaborato, nei decenni precedenti, concetti affini: dall'*information inequality* di Herbert Schiller al *knowledge gap* proposto da Tichenor e collaboratori alla fine degli anni Sessanta, ipotesi con cui si tentava di descrivere come i segmenti di popolazione con status socioeconomico più elevato tendessero ad acquisire informazione dai mass media a una velocità superiore rispetto ai segmenti meno istruiti (van Dijk, 2020, p. 1). Nel contesto statunitense si parlava inoltre di *democratic divide* e di *economic opportunity divide* per designare fenomeni che la successiva letteratura sul divario digitale avrebbe in parte inglobato.

Il successo del termine *digital divide* rispetto ad altre espressioni concorrenti si spiega probabilmente con la sua forza immediata, l'immagine del solco, della frattura che separa due gruppi, si lascia afferrare senza bisogno di mediazioni concettuali. Proprio questa immediatezza, tuttavia, porta con sé una serie di problemi interpretativi che van Dijk (2020, pp. 2-3) ha esplicitato in modo sistematico, identificando quattro malintesi indotti dalla metafora stessa.

Il primo è l'idea che esistano due popolazioni nettamente separate, una che ha accesso alla tecnologia e una che ne è esclusa. Nella realtà la variazione è continua: tra chi possiede una connessione a fibra ottica, un computer di ultima generazione e tempo a disposizione per utilizzarlo in modo sofisticato e chi accede a Internet soltanto attraverso uno smartphone datato con un piano dati limitato, esistono innumerevoli posizioni intermedie, ciascuna con implicazioni diverse in termini di opportunità concrete. La lettura dicotomica semplifica eccessivamente uno spettro che si configura piuttosto come una gradazione.

Un secondo aspetto è la presunta invalicabilità del divario. L'immagine del ponte che non può essere attraversato suggerisce una separazione strutturale, permanente, di fatto insuperabile. La ricerca empirica restituisce un quadro diverso: le distanze possono ridursi (nel corso degli ultimi vent'anni l'accesso fisico a Internet si è notevolmente diffuso anche in segmenti di popolazione inizialmente esclusi) pur con una precisazione importante (come verrà affrontata nel paragrafo 2.4), la riduzione di una disparità tende a coincidere con l'emergere di nuove disparità su altri piani.

C'è poi l'assolutezza della distinzione: parlare di inclusione ed esclusione come stati opposti induce a pensare la disuguaglianza digitale in termini binari, mentre le dimensioni che effettivamente contano (la motivazione e l'accesso materiale, ma anche le competenze, la qualità dell'uso, i benefici ottenuti) sono tutte distinzioni relative. Due persone formalmente incluse nel sistema digitale possono trovarsi in posizioni radicalmente diverse rispetto a ciò che riescono effettivamente a fare, e a ottenere, attraverso di esso.

Resta infine l'idea che esista un unico divario digitale. La metafora binaria, suggerendo una singola linea di frattura tra inclusi ed esclusi, occulta il fatto che la disuguaglianza digitale si articola in realtà su

molteplici piani, ciascuno con dinamiche proprie, e si lega alle disparità sociali, economiche e culturali già esistenti nella società (van Dijk, 2020, p. 3). Negli anni Duemila, quando il divario sull'accesso a Internet iniziava a ridursi, ne stava già emergendo un altro sulla qualità della connessione. Con la diffusione della banda larga il discrimine si è spostato sulle competenze d'uso, e oggi, mentre osservatori e policy maker dichiarano superate le questioni della prima ora, sono altre le asimmetrie a farsi avanti: quelle legate alla capacità di comprendere il funzionamento degli algoritmi e dei sistemi di intelligenza artificiale. Ogni nuova fase tecnologica apre una linea di frattura senza chiudere completamente le precedenti.

La critica alla metafora binaria non è peraltro un risultato isolato del lavoro di van Dijk, già nel 2002 la sociologa statunitense Eszter Hargittai, in un articolo destinato a diventare uno dei riferimenti più citati del settore, argomentava che, con la progressiva diffusione di Internet nella popolazione americana, continuare a guardare al divario digitale in termini di classificazione binaria (chi è online e chi no) diventava sempre meno utile. Occorreva piuttosto spostare l'attenzione sulle differenze nelle competenze con cui le persone si muovevano effettivamente nel medium, ciò che Hargittai (2002) denominerà *second-level digital divide*. Pur muovendosi da un angolo di osservazione diverso (empirico, basato su test di navigazione somministrati a un campione di utenti statunitensi) Hargittai giungeva a una conclusione convergente con quella di van Dijk: la metafora del confine invalicabile tra inclusi ed esclusi non coglieva ciò che la ricerca stava cominciando a documentare, ovvero che le differenze più rilevanti si collocavano dentro la popolazione degli utenti, non tra utenti e non utenti.

Da quest'analisi emerge una conclusione che van Dijk formula in modo particolarmente netto: «The term digital suggests that the digital

divide is a technical issue when, in fact, it is more of a social problem. Technical properties of digital media are important for access and use — but the causes and effects of (in)equality are social» (van Dijk, 2020, p.3). L'affermazione capovolge il presupposto più diffuso nel discorso pubblico e in parte della produzione divulgativa. L'autore non nega il ruolo della tecnologia: dispositivi e reti sono ovviamente la preconditione materiale dell'intero fenomeno. Quello che van Dijk vuole mettere in chiaro è che le disparità digitali hanno cause sociali e producono effetti sociali. Le condizioni in cui una persona accede alla tecnologia dipendono dalla sua classe d'origine; le competenze d'uso si sviluppano in modo differenziato lungo le linee del capitale culturale; la capacità di tradurre l'uso in benefici è correlata al contesto familiare e formativo. Anche gli effetti delle disparità ricalcano stratificazioni socio-demografiche già note, in cui pesano in primis reddito e istruzione, ma anche età, capitale culturale, collocazione geografica e posizione lavorativa. Nella stessa logica va letta un'altra affermazione dell'autore, altrettanto esplicita: «The digital divide is here to stay even when all such problems are overcome» (van Dijk, 2020, p. 3). Le disparità nella diffusione possono ridursi, ma quelle che riguardano il rapporto delle persone con la tecnologia tendono a permanere.

La rilevanza della critica di van Dijk alla metafora dicotomica non è soltanto di ordine storico-teorico. La sua attualità si coglie nel modo in cui il discorso pubblico contemporaneo, in particolare quello veicolato dalla stampa generalista, tende ancora oggi a rappresentare l'intelligenza artificiale come una tecnologia rispetto alla quale ci si può collocare soltanto dentro o fuori: come entusiasti che accettano le sue promesse senza riserve, oppure come oppositori che le contestano per ragioni etiche o culturali. Questo approccio binario, particolarmente visibile nel dibattito sull'uso dell'AI in ambito universitario, rischia di riprodurre esattamente

l'errore interpretativo che la ricerca sociologica ha da tempo denunciato, con conseguenze concrete sulle pratiche formative (questione su cui ritorneremo nel Capitolo 3). Se dunque il problema è sociale, e se le metafore binarie risultano inadeguate a rappresentarlo, diventa necessario disporre di un modello teorico capace di articolare il fenomeno nelle sue diverse dimensioni: è il lavoro che van Dijk ha condotto nel corso del suo itinerario di ricerca, costruendo progressivamente il framework multidimensionale che verrà presentato nel prossimo paragrafo.

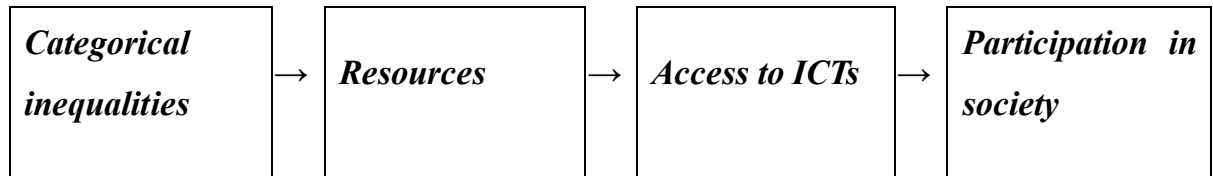
2.3 Il modello multidimensionale di van Dijk

La proposta più articolata per superare la lettura binaria del divario digitale è quella elaborata da van Dijk a partire dai primi anni Duemila. Il suo framework, nella versione consolidata del 2020, mette insieme cinque elementi: una teoria causale, una tassonomia delle risorse, una classificazione delle disuguaglianze sociali, una scansione delle fasi dell'accesso e un inventario delle competenze necessarie. È il sistema di coordinate concettuali entro cui si muoverà l'intera tesi, e per questo conviene esporlo con una certa cura.

Il cuore del framework è la cosiddetta *Resources and Appropriation Theory*, formulata da van Dijk nel volume *The Deepening Divide* del 2005 e ripresa con lievi aggiustamenti nell'edizione del 2020 qui di riferimento. La teoria si fonda su un'ipotesi causale: le disuguaglianze nell'accesso e nell'uso delle tecnologie digitali non possono essere analizzate come fenomeni isolati, ma vanno collegate alle disuguaglianze sociali preesistenti attraverso il concetto intermedio di risorse. In altri termini, le asimmetrie sociali (di classe e istruzione anzitutto, ma anche di età, di genere, di area geografica) si traducono in una distribuzione diseguale di risorse materiali e immateriali, e sono queste risorse a determinare, in

ultima istanza, chi può effettivamente accedere e trarre beneficio dalle tecnologie dell'informazione (van Dijk, 2020, pp. 30-33).

Lo schema causale che ne risulta può essere rappresentato in forma lineare attraverso quattro passaggi concatenati.



Schema 1: Il modello causale della Resources and Appropriation Theory (van Dijk, 2020, p. 33)

La catena, tuttavia, non si chiude al quarto passaggio. Uno degli elementi teorici più rilevanti del modello è la sua circolarità: la partecipazione diseguale alla società, prodotta dall'accesso diseguale alle tecnologie, retroagisce sulle disuguaglianze categoriali e sulla distribuzione delle risorse da cui il ciclo era partito, rinforzandole. Le cinque proposizioni in cui van Dijk sintetizza la teoria rendono esplicito questo meccanismo:

1. *Categorical inequalities in society produce an unequal distribution of resources.*
2. *An unequal distribution of resources causes unequal access to digital technologies.*
3. *Unequal access to digital technologies also depends on the characteristics of these technologies.*
4. *Unequal access to digital technologies brings about unequal participation in society.*
5. *Unequal participation in society reinforces categorical inequalities and unequal distribution of resources.*

Nella quinta proposizione si concentra la lettura teorica più importante del modello (van Dijk, 2020, p. 31). Il divario digitale, secondo van Dijk, non si supera iniettando tecnologia in una società altrimenti equa. Funziona piuttosto come un meccanismo che rinforza le asimmetrie sociali preesistenti. Chi parte da una posizione di vantaggio tende a consolidarla attraverso l'uso della tecnologia, mentre per chi parte svantaggiato la stessa tecnologia diventa un'ulteriore occasione di marginalizzazione. Questa lettura, che van Dijk (2020, p. 117) riconduce esplicitamente all'effetto Matteo formulato da Robert Merton nel 1968, ha implicazioni decisive per il modo in cui interpreteremo, nei prossimi capitoli, l'impatto dell'AI sulla stratificazione sociale.

Il secondo elemento costitutivo del modello è la tassonomia delle risorse. Van Dijk identifica cinque categorie che concorrono a determinare l'accesso e l'uso delle tecnologie digitali. La prima, le risorse temporali, riguarda il tempo di cui un individuo dispone per utilizzare i media digitali: una variabile spesso sottovalutata, ma particolarmente sensibile in relazione al carico lavorativo e alle responsabilità di cura. Vengono poi le risorse materiali, ossia il reddito e il possesso di beni, che determinano la possibilità di acquistare dispositivi di qualità e di sottoscrivere connessioni performanti. Le risorse mentali, terza categoria, comprendono le capacità cognitive e le abilità tecniche di base. La quarta categoria, le risorse sociali, è la rete di supporto a cui una persona può rivolgersi quando incontra difficoltà nell'uso della tecnologia: una dimensione che, come la ricerca empirica ha mostrato, pesa soprattutto sugli utenti meno esperti. Le risorse culturali, infine, comprendono tutto ciò che, attraverso lo stile di vita acquisito per via familiare e i segnali di status che lo accompagnano, rende l'uso delle tecnologie digitali più o meno naturale o ambito (van Dijk, 2020, pp. 27, 32).

Da questa classificazione emerge un punto importante. Le cinque categorie di risorse esistono indipendentemente dalla tecnologia: erano già distribuite in modo diseguale nella società ben prima della rivoluzione digitale, e continuano a esserlo oggi. Quando una nuova tecnologia si diffonde, queste risorse condizionano chi può accedervi e a quali condizioni. È un punto che torna, sotto forme diverse, in ogni passaggio dell'opera di van Dijk, e che costituisce la traduzione teorica dell'affermazione (richiamata nel paragrafo precedente) secondo cui il divario digitale è un problema sociale prima che tecnico. Si tratta peraltro di un punto pienamente in linea con il quadro sociologico ricostruito nel Capitolo 1: la tassonomia delle risorse di van Dijk riprende e specializza, applicandola al dominio digitale, la stessa logica che Croteau e Hoynes (2022, p. 314) avevano formulato in termini più generali con le sette categorie di risorse alla base dei sistemi di stratificazione contemporanei.

Le categorie di disuguaglianza su cui il modello opera si dividono, nell'articolazione di van Dijk, in personali e posizionali. Tra le prime van Dijk colloca l'etnia, lo stato di salute, l'età, il genere e i tratti di personalità; tra le seconde, l'istruzione, la posizione lavorativa, la composizione familiare, la collocazione nelle reti sociali e il contesto geografico, con le distinzioni fondamentali tra aree urbane e rurali e tra paesi sviluppati e in via di sviluppo (van Dijk, 2020, p. 32). La combinazione di queste variabili con la tassonomia delle risorse produce, per ciascun individuo, un profilo di accesso e di uso altamente specifico: è questa la ragione per cui, come si è visto nel paragrafo 2.2, la lettura binaria del divario non regge empiricamente.

Il terzo elemento del framework è la scansione delle fasi attraverso cui si realizza l'accesso alle tecnologie digitali. Van Dijk identifica quattro momenti che, pur essendo concettualmente distinti, si influenzano reciprocamente e non vanno letti in senso strettamente sequenziale. La

motivazione (o meglio l'atteggiamento rispetto alla tecnologia) costituisce la prima soglia: senza un interesse o una necessità percepita, l'accesso non viene nemmeno tentato. La seconda fase, quella dell'*accesso fisico*, comporta il possesso materiale dei dispositivi e la disponibilità di una connessione. Segue la fase delle *competenze digitali*, ossia delle capacità necessarie per utilizzare effettivamente gli strumenti. L'ultima fase è quella dell'*uso*, inteso come impiego concreto della tecnologia in attività specifiche (van Dijk, 2020, p. 14). Una persona può avere motivazione e accesso fisico ma mancare di competenze sufficienti, può avere accesso e competenze ma non utilizzare la tecnologia in modi che producono benefici concreti. Ciascuna fase è condizione necessaria ma non sufficiente per quella successiva.

La terza fase (quella delle competenze digitali) merita un approfondimento, perché rappresenta il passaggio su cui la ricerca ha concentrato maggiore attenzione a partire dalla metà degli anni Duemila. Costituisce il passaggio decisivo per il passaggio dal divario digitale classico ai divari più recenti, inclusi quelli algoritmici. Van Dijk, in collaborazione con Alexander van Deursen, ha elaborato una classificazione a sei competenze che distingue tra abilità relative al medium e abilità relative ai contenuti.

<i>Competenze medium-related</i>	<i>Competenze content-related</i>
• <i>operational</i>: uso di base dei dispositivi • <i>formal</i>: navigazione e orientamento negli ambienti digitali	• <i>information</i>: ricerca, selezione e valutazione di informazioni • <i>communication</i>: interazione online e gestione dell'identità digitale • <i>content-creation</i>: produzione di contenuti originali • <i>strategic</i>: uso dei media digitali per obiettivi personali o professionali

Schema 2: Le sei competenze digitali (van Dijk, 2020, p. 68)

Le prime due competenze (*operational* e *formal*) riguardano l'interazione con il medium digitale: dal saper accendere un dispositivo alla padronanza dei comandi di base, fino all'orientamento all'interno dei siti web e alla comprensione della logica ipertestuale. Le altre quattro attengono invece ai contenuti. Le competenze *information* permettono di cercare, selezionare e valutare criticamente le informazioni reperite in rete. Le competenze *communication* hanno a che fare con l'interazione tra utenti, la gestione della propria identità digitale e le dinamiche dell'attenzione online. Vi sono poi le competenze *content-creation*, che rinviano alla capacità di produrre contenuti originali (testi, immagini, video) destinati alla circolazione in rete. Le competenze *strategic*, infine, indicano la capacità di utilizzare i media digitali per raggiungere obiettivi personali o professionali concreti (van Dijk, 2020, p. 68).

La distinzione tra le prime due e le ultime quattro competenze non è soltanto descrittiva, la ricerca empirica (come si avrà modo di mostrare nel paragrafo 2.4) ha documentato che, mentre le competenze *medium-*

related tendono a distribuirsi in modo relativamente omogeneo nella popolazione connessa, le competenze *content-related*, e in particolare quelle strategiche, risultano correlate in misura molto più netta con il livello di istruzione e lo status socioeconomico. Su questo terreno si innesta la discussione più recente sull'*algorithmic divide* e sull'AI divide: l'interazione efficace con i sistemi algoritmici e con l'intelligenza artificiale richiede infatti competenze che si collocano ai vertici della tassonomia di van Dijk, quelle strategiche in particolare, e che sono distribuite in modo più diseguale di quanto la diffusione capillare degli smartphone lasci intuire.

Il modello così ricostruito fornisce il sistema di coordinate entro cui l'analisi del divario digitale può essere condotta con il grado di complessità che il fenomeno richiede. La teoria causale spiega perché il divario si produce e come si riproduce. Le tassonomie delle risorse e delle categorie di disuguaglianza specificano da dove derivano le asimmetrie, mentre la scansione delle fasi dell'accesso e la classificazione delle competenze mostrano attraverso quali passaggi le disuguaglianze si traducono in esclusione. Con questo framework è possibile affrontare ora, nel prossimo paragrafo, la ricostruzione storica dei tre livelli in cui la ricerca sul divario si è articolata, mostrando come ciascuno corrisponda a un diverso passaggio del modello appena presentato.

2.4 I tre livelli del divario digitale

Disporre del modello multidimensionale presentato nel paragrafo precedente permette ora di affrontare la ricostruzione storica che la letteratura ha elaborato per dare conto dell'evoluzione della ricerca sul divario digitale. Negli ultimi venticinque anni gli studi hanno progressivamente spostato il proprio baricentro, e la letteratura si è

organizzata intorno a una tripartizione divenuta convenzionale: il primo livello è centrato sull'accesso fisico, il secondo si occupa di competenze e modalità d'uso, il terzo dei benefici concreti che le persone riescono a ottenere dalle tecnologie. La corrispondenza con il modello di van Dijk è abbastanza precisa: l'accesso fisico è la seconda fase del modello, le competenze la terza, l'uso e gli *outcomes* la quarta. La periodizzazione, tuttavia, non va letta in senso strettamente sostitutivo: come si vedrà, i tre livelli si sono accumulati più che succeduti, e ciascuna delle fratture identificate nella prima fase persiste, sotto forme trasformate, nelle fasi successive.

2.4.1 Primo livello: l'accesso fisico

La prima fase della ricerca sul *digital divide* si estende, secondo la periodizzazione proposta da van Dijk (2020, pp. 6-8), tra il 1995 e il 2003 circa, è la fase in cui il termine nasce e si afferma, e in cui il problema viene letto essenzialmente in termini di possesso materiale: chi ha un computer e una connessione a Internet, e chi non li ha. La semplicità dell'impostazione rispecchiava la realtà di quegli anni, quando l'accesso alla rete era una condizione minoritaria che pesava economicamente sulle famiglie e si distribuiva in modo ineguale sul territorio. Per dare un riferimento numerico ai dati statunitensi richiamati da van Dijk (2020, p. 8): nel 1993 possedeva un computer poco più di un americano su dieci, e nello stesso anno la percentuale di utenti Internet era pari a circa un americano su tre. Otto anni dopo entrambe le percentuali avevano raggiunto il 50%.

La letteratura di questa fase è dominata da un'impostazione descrittiva, tesa a misurare la diffusione dei dispositivi e a correlarla con variabili sociodemografiche classiche: reddito e istruzione anzitutto, ma

anche età, etnia e area geografica. Un contributo significativo arriva da Pippa Norris, che nel volume *Digital Divide* del 2001 propone una distinzione divenuta canonica tra tre livelli del fenomeno: il *global divide*, ossia la frattura tra paesi sviluppati e in via di sviluppo nell'accesso alle ICT; il *social divide*, che riguarda le diseguaglianze interne a ciascun paese; il *democratic divide*, relativo alle differenze nella capacità di utilizzare la rete per la partecipazione politica (Norris, 2001, citata in van Dijk, 2020, p. 6).

Si tratta di una tripartizione che anticipa in parte la stratificazione successiva della ricerca, pur muovendosi ancora all'interno del paradigma dell'accesso.

Rilevante è uno snodo cruciale, richiamato da van Dijk in chiave quasi anedddotica ma rivelatore del clima politico del tempo, è la progressiva smobilitazione del termine sotto l'amministrazione Bush, quando il *digital divide* viene, nelle parole dell'autore, «officially buried» (van Dijk, 2020, p. 8). La dismissione non corrispondeva a una risoluzione del problema, ma al cambiamento delle priorità politiche e alla convinzione, diffusa in quegli anni, che il mercato avrebbe spontaneamente colmato le disparità. Nei fatti, il problema non scompariva: si spostava, mentre il tasso di penetrazione di Internet cresceva rapidamente nelle economie avanzate. Diventava sempre più evidente alla comunità scientifica che l'aver collegato una persona alla rete non equivaleva ad averne risolto il rapporto con la tecnologia, chi era online poteva usare Internet in modi radicalmente diversi, con esiti altrettanto differenziati. Questa consapevolezza segnerà il passaggio alla fase successiva della ricerca.

2.4.2 Secondo livello: competenze e modalità d'uso

La transizione verso quello che oggi viene chiamato *second-level digital divide* si colloca, convenzionalmente, nei primissimi anni Duemila. Il termine viene coniato da Eszter Hargittai (2002) in un articolo pubblicato su «First Monday», destinato a diventare uno dei riferimenti più citati del settore. Hargittai argomentava che, con la progressiva estensione dell'accesso a Internet, continuare a leggere il divario in termini dicotomici (chi è online e chi no) diventava sempre meno informativo: la ricerca doveva spostarsi, piuttosto, sulle differenze nelle competenze con cui le persone si muovevano effettivamente nella rete, documentando che tempi, strategie e capacità di reperire informazioni online variavano in modo sostanziale anche all'interno della popolazione già connessa.

Negli stessi anni altri autori convergono su conclusioni simili, Mark Warschauer nel 2003 propone un «rethinking» del divario digitale, mentre Neil Selwyn nel 2004 parla di «reconceptualizing» (van Dijk, 2020, p. 8). Anche van Dijk, insieme a Kenneth Hacker, già nel 2003 definisce il divario digitale come «a complex and dynamic phenomenon» (van Dijk, 2020, p. 8), contribuendo ad abbandonare le letture binarie. Il comune denominatore di questi interventi è il riconoscimento che l'accesso materiale, pur restando condizione necessaria, non spiega le asimmetrie che la diffusione di Internet stava producendo, occorre guardare, con strumenti analitici nuovi, alle modalità d'uso.

È in questo contesto che van Dijk elabora il concetto di *usage gap*, divenuto uno dei riferimenti più efficaci per descrivere il fenomeno. Il termine, introdotto dall'autore già alla fine degli anni Novanta e ripreso in lavori successivi, indica l'esistenza di un uso sistematicamente diverso di Internet da parte di classi sociali diverse: le persone con istruzione e reddito più elevati tendono a utilizzare la rete per attività *capital-enhancing* (informazione, apprendimento, attività lavorative,

partecipazione civica), mentre le persone con status socioeconomico più basso tendono a privilegiare attività legate all'intrattenimento, al consumo e alla comunicazione interpersonale semplice (van Dijk, 2020, pp. 92-94).

Van Dijk collega esplicitamente l'*usage gap* a una tradizione di ricerca più antica, quella del *knowledge gap* formulata da Phillip Tichenor, George Donohue e Clarice Olien nel 1970 in riferimento ai mass media tradizionali. L'ipotesi dei tre sociologi americani è che «as the diffusion of mass media information in a social system increases, segments of the population with a higher socio-economic status tend to acquire this information at a faster rate than the lower status segments» (Tichenor, Donohue & Olien, 1970, citato in van Dijk, 2020, p. 93). L'aumento della disponibilità di informazione, paradossalmente, tende a produrre una crescita più rapida della conoscenza nei segmenti più istruiti, ampliando anziché ridurre le distanze. L'*usage gap* di van Dijk estende questo meccanismo al dominio di Internet: la diffusione dello strumento, lungi dal garantire un accesso equo ai benefici, rischia di riprodurre nei ceti più svantaggiati usi meno sofisticati e meno capitalizzabili.

Sul piano dell'evidenza empirica, uno dei contributi più rigorosi della fase viene dallo studio di van Deursen e van Dijk (2011) pubblicato su «New Media & Society». I due autori sottopongono a un campione rappresentativo della popolazione olandese una serie di test pratici di esecuzione. La novità metodologica è importante: invece di chiedere ai partecipanti di autovalutare le proprie competenze (come faceva gran parte della letteratura precedente), gli autori somministrano compiti effettivi da svolgere online. I risultati restituiscono un quadro stratificato. Le competenze *operational* e *formal*, cioè l'uso di base dei dispositivi e l'orientamento negli ambienti digitali, sono relativamente diffuse nella popolazione. Le competenze *information* e quelle *strategic*, invece, indispensabili per valutare criticamente le fonti e per usare la rete in vista

di obiettivi concreti, si distribuiscono in modo molto più diseguale, con l'istruzione come variabile determinante. Il dato è cruciale per il prosieguo della discussione: come si vedrà nei prossimi paragrafi, l'interazione efficace con i sistemi algoritmici e con l'intelligenza artificiale richiede precisamente le competenze che questa fase di ricerca ha documentato essere meno equamente distribuite.

Una conseguenza diretta del secondo livello di analisi è la messa in discussione del cosiddetto mito dei *digital natives*, espressione divenuta popolare nel dibattito pubblico e nella letteratura divulgativa per designare le generazioni cresciute nell'ambiente digitale e ritenute, per questa ragione, naturalmente più competenti nell'uso delle tecnologie. Van Dijk, sulla base di una letteratura ormai consistente (da Hargittai a Helsper ed Eynon fino a Calvani e collaboratori) osserva che «the notion of "digital natives" and "net generation" is very wide-ranging» (van Dijk, 2020, p. 72) e che i risultati empirici sono tutt'altro che univoci: in numerose ricerche le generazioni più anziane si rivelano, contrariamente al senso comune, più abili nell'uso di Internet rispetto alle più giovani, specialmente nelle dimensioni *content-related* che richiedono valutazione critica dell'informazione. La presunta *nativeness* digitale si rivela così, più che un dato oggettivo, una rappresentazione sociale che rischia di oscurare le reali diseguaglianze di competenza all'interno delle stesse coorti giovanili.

Vale la pena una nota a margine, da osservazione diretta. Nel contesto italiano le politiche delle singole università sull'uso dell'AI da parte degli studenti sono oggi piuttosto eterogenee, e questa eterogeneità sembra riprodurre in tempo reale le dinamiche di *usage gap* descritte dalla letteratura. All'Università di Pavia, per esempio, gli studenti vengono incoraggiati a utilizzare i modelli generativi all'interno di specifici insegnamenti, con l'obiettivo di sviluppare una familiarità critica con

questi strumenti. In altri atenei, come l'Università di Firenze per quanto mi riferiscono colleghi, l'uso dell'AI è invece oggetto di un atteggiamento restrittivo, a favore di una didattica più tradizionale. Il punto non è giudicare le scelte di un ateneo rispetto a un altro, ma constatare che due studenti con profilo sociodemografico simile, iscritti a corsi affini in due università diverse, finiscono per maturare competenze differenziate verso la tecnologia non per risorse individuali ma per effetto del contesto istituzionale. È un *usage gap* generato a monte dalla struttura formativa, con ricadute sul capitale algoritmico di ciascuno studente (un tema su cui si tornerà nel Capitolo 3).

2.4.3 Terzo livello: i benefici concreti

A partire dalla fine della prima decade del nuovo millennio la ricerca ha progressivamente spostato il proprio focus su una dimensione ulteriore: i benefici concreti (nella letteratura anglofona *tangible outcomes*) che le persone riescono effettivamente a ottenere dall'uso di Internet nella vita quotidiana. Il termine *third-level digital divide* viene coniato nel 2010 dai sociologi spagnoli José Manuel Robles Morales, Cristóbal Torres Albero e Óscar Molina (van Dijk, 2020, p. 11), e viene ripreso negli anni successivi da una letteratura crescente che ne sviluppa l'impianto teorico ed empirico.

Il contributo più sistematico al terzo livello è quello di Helsper e van Dijk (2012), che propongono una teoria degli *outcomes* articolata in quattro domini: economico, culturale, sociale e benessere personale. Un utilizzo efficace di Internet, secondo gli autori, può tradursi in benefici economici (aumento del reddito, migliori opportunità lavorative, risparmi nei consumi), culturali (accesso a informazione e formazione di qualità), sociali (ampliamento e rafforzamento delle reti di relazione) e di benessere

personale (miglioramento della salute fisica e psicologica, aumento della soddisfazione di vita). Il punto decisivo sul piano teorico è che la traduzione dell'uso in *outcomes* non è automatica: persone con accesso e competenze simili possono ottenere benefici molto diversi a seconda delle risorse (nel senso che van Dijk attribuisce al termine) di cui dispongono nella vita offline.

La formulazione empirica più documentata di questa tesi si trova nello studio di van Deursen e Helsper (2015), pubblicato all'interno del *Communication and Information Technologies Annual* della Emerald. L'articolo, che reca nel titolo la domanda «Who Benefits Most from Being Online?», conclude la propria analisi della popolazione olandese con un risultato che ha ricevuto ampia eco nella letteratura: Internet risulta «more beneficial for those with higher social status, not in terms of how extensively they use the technology but in what they achieve as a result of this use» (van Deursen & Helsper, 2015, p. 29). Sulla distinzione vale la pena soffermarsi. A separare le classi sociali non è la quantità di tempo trascorso online o l'intensità d'uso, variabili su cui i ceti meno abbienti possono talvolta superare quelli superiori (passando davanti allo schermo persino più ore). La differenza sta semmai nel rendimento di quel tempo: in ciò che si riesce a ricavarne in termini di informazione utile e di opportunità professionali, di relazioni che contano e di capitale culturale accumulato. A questo terzo livello, in sostanza, ciò che pesa è quanto la rete renda a chi la usa.

La domanda di ricerca che orienta il terzo livello, esplicitata da van Dijk (2020, p. 12), può essere formulata in termini piuttosto netti: il divario digitale riduce le disuguaglianze sociali esistenti, o le rinforza? E, parallelamente, ne crea di nuove? La letteratura empirica orienta verso la seconda lettura. Come si è già anticipato nel paragrafo 2.3, il modello di van Dijk ipotizza una retroazione tra partecipazione diseguale e

disuguaglianze categoriali (la quinta proposizione della Resources and Appropriation Theory) ed è proprio al terzo livello che questa retroazione diventa osservabile empiricamente. Il divario digitale non si limita a riprodurre le disparità offline: le amplifica, producendo un effetto cumulativo che van Dijk riconduce esplicitamente all'effetto Matteo mertoniano.

La tripartizione che la letteratura ha elaborato nell'arco di venticinque anni può essere riassunta nello schema seguente, che mette in evidenza la continuità e insieme la distinzione dei tre livelli.

Livello	Periodo	Focus	Autori di riferimento
Primo	1995–2003	Accesso fisico: possesso di dispositivi e connessione	NTIA 1995; Norris 2001
Secondo	2004–oggi	Competenze e modalità d'uso	Hargittai 2002; Warschauer 2003; Selwyn 2004; van Deursen & van Dijk 2011
Terzo	2010–oggi	Benefici concreti (outcomes) nella vita offline	Robles Morales et al. 2010; Helsper & van Dijk 2012; van Deursen & Helsper 2015

Schema 3: I tre livelli del divario digitale: periodizzazione, focus e autori di riferimento

Lo schema va letto con una cautela già richiamata in apertura del paragrafo: le tre fasi non si succedono escludendosi, ma si stratificano. Ancora oggi il primo livello resta un problema concreto per segmenti di popolazione (gli anziani, le aree rurali, i gruppi a basso reddito, soprattutto

nei paesi in via di sviluppo ma anche in segmenti tutt'altro che marginali di quelli sviluppati) mentre contemporaneamente la ricerca discute i meccanismi del terzo livello in popolazioni quasi interamente connesse. Il quadro empirico è dunque quello di una coesistenza di fratture, più che di una sostituzione progressiva. A questa tripartizione, nell'ultimo decennio, si è aggiunta una quarta dimensione, legata alla diffusione dei sistemi algoritmici che mediano l'accesso ai contenuti e governano porzioni crescenti dei processi informativi, economici e lavorativi. Di questa trasformazione, e dei divari che essa produce, si occupa il prossimo paragrafo.

2.5 Dalla svolta algoritmica all'*algorithmic divide*

Il quadro a tre livelli ricostruito nel paragrafo precedente si è consolidato, nella letteratura, in un arco temporale che arriva fino alla prima metà dello scorso decennio. Negli anni più recenti, tuttavia, la ricerca ha dovuto confrontarsi con una trasformazione del panorama digitale che non si lasciava ricondurre agevolmente alle categorie esistenti. La diffusione pervasiva dei sistemi di raccomandazione, dei motori di ricerca personalizzati, dei meccanismi di filtraggio automatico dei contenuti ha introdotto nella vita quotidiana degli utenti una mediazione algoritmica che, pur operando spesso in modo invisibile, orienta in misura crescente l'accesso all'informazione, al consumo, alle relazioni sociali. Non si tratta più, soltanto, di chi ha accesso a Internet o di quali competenze possiede per usarlo, ma di come i contenuti che riceve vengono selezionati, ordinati, proposti da sistemi automatici che l'utente, nella maggior parte dei casi, non vede e non comprende. Questa trasformazione ha spinto la letteratura a elaborare un nuovo concetto, quello di *algorithmic divide*, per designare le disuguaglianze che si generano intorno a tale mediazione.

Lo stesso van Dijk, nelle pagine conclusive dell'edizione 2020 del suo volume, anticipa questa evoluzione individuando quattro direzioni lungo cui la tecnologia digitale andrà incontro nei prossimi anni. La moltiplicazione dei dispositivi, il bisogno crescente di comprendere i sistemi e non solo le interfacce, la miniaturizzazione dei *device* fino alla loro integrazione nel corpo, la fusione tra realtà fisica e realtà digitale: quattro tendenze che, secondo l'autore, produrranno un'ulteriore polarizzazione delle competenze richieste agli utenti (van Dijk, 2020, pp. 127-129). Il passaggio più rilevante per il discorso che si sta sviluppando riguarda la seconda direzione. Van Dijk osserva che fino ad oggi la ricerca sul divario digitale si è occupata prevalentemente dell'interfaccia (il rapporto tra l'utente e lo schermo, tra l'utente e l'applicazione), mentre negli anni a venire diventerà decisivo il rapporto degli utenti con i sistemi sottostanti: «Most users — perhaps only 1 or 2 per cent of the entire population on earth — do not even know how the Internet as a system works, and the complicated systems just mentioned are even more difficult to understand and be controlled by the average user» (van Dijk, 2020, p. 128).

La formulazione è significativa, perché mette a fuoco un elemento nuovo rispetto al framework classico: non soltanto la capacità di utilizzare una tecnologia, ma quella di comprenderne il funzionamento. Van Dijk esplicita la questione con riferimento diretto agli algoritmi: «These new digital media are directed by artificial intelligence and continually produce big data. Decisions are made not only by users but also automatically by algorithms, so users need to have the strategic skills to follow or override these decisions» (van Dijk, 2020, p. 128). Vale la pena fermarsi su questo passaggio. Nell'ambiente mediato dagli algoritmi l'utente non è il soggetto che prende la maggior parte delle decisioni rilevanti per la propria esperienza online: i contenuti che vede, l'ordine con cui gli vengono

proposti, la frequenza con cui ricompaiono nel suo feed dipendono dal sistema. La domanda da porre allora non riguarda soltanto la capacità d'uso della tecnologia. Riguarda se l'utente sia in grado di riconoscere quando una decisione è stata presa automaticamente, e se possieda gli strumenti per validarla, eventualmente, oppure per scavalcarla.

La riflessione di van Dijk può essere utilmente integrata con la cornice macro-sociologica elaborata da Manuel Castells, il cui lavoro sulla *network society* (avviato negli anni Novanta e articolato nella trilogia *The Information Age*) ha fornito alla sociologia dei media contemporanea una delle chiavi interpretative più influenti. In un articolo del 2007 pubblicato sull'*International Journal of Communication*, Castells riformula in termini sintetici la propria tesi centrale: nella società in rete il potere non si esercita più, primariamente, attraverso il controllo di territori o di istituzioni, ma attraverso il controllo dei flussi di informazione e comunicazione che organizzano l'economia, la politica e la cultura globali (Castells, 2007).

Al cuore di questa lettura Castells individua due figure: i *programmers* e gli *switchers*. I *programmers* definiscono le regole secondo cui le reti operano: stabiliscono come l'informazione circola, quali contenuti vengono spinti in alto nella gerarchia di visibilità, quali invece restano occultati. Gli *switchers*, invece, collegano fra loro reti che opererebbero altrimenti come sistemi separati, occupando posizioni di passaggio tra mondi informativi, economici, politici altrimenti scollegati. Si tratta di posizioni che conferiscono una forma di potere che Castells considera caratteristica della *network society*. Lo definisce un potere infrastrutturale: il suo esercizio non avviene tramite imposizione diretta, ma decidendo l'ambiente entro cui i restanti attori sono costretti a operare. Applicata al dominio algoritmico, questa lettura consente di inquadrare la mediazione automatizzata in termini non meramente tecnici. Gli algoritmi che governano le piattaforme digitali sono, nella sostanza, dispositivi che

stabiliscono le regole della visibilità e della circolazione dei contenuti, e le organizzazioni e gli individui che li progettano e mantengono occupano, rispetto alla vasta maggioranza degli utenti, una posizione strutturalmente asimmetrica. *L'algorithmic divide*, letto attraverso Castells, è quindi una vera e propria configurazione del potere, non solo una frattura di competenze individuali.

A livello di ricerca empirica, uno dei contributi più influenti al tema è l'articolo pubblicato nel 2021 su «Information, Communication & Society» da Gran, Booth e Bucher, il cui titolo stesso, *To be or not to be algorithm aware: a question of a new digital divide?*, pone la questione nei termini che interessano la presente discussione (Gran, Booth & Bucher, 2021). Gli autori introducono il concetto di *algorithm awareness*, traducibile come consapevolezza algoritmica, per indicare la capacità degli utenti di riconoscere che gli algoritmi esistono e influenzano i contenuti che ricevono su piattaforme come YouTube, Spotify, Meta, i motori di ricerca, le piattaforme di streaming di informazione. La domanda di ricerca è se il possesso o meno di tale consapevolezza corrisponda a un nuovo livello di divario digitale.

L'indagine è stata condotta su un campione rappresentativo della popolazione norvegese, scelto non a caso: la Norvegia è uno dei paesi al mondo con i tassi più elevati di connessione e di alfabetizzazione digitale di base, il che consente di isolare l'asimmetria che interessa gli autori al netto dei divari classici. I risultati mostrano che, anche in un contesto di questo tipo, la consapevolezza algoritmica si distribuisce in modo marcatamente diseguale: l'età, il livello di istruzione e il genere risultano variabili predittive significative, con gli utenti più giovani, più istruiti e di genere maschile che presentano livelli di *awareness* più elevati. Gli autori formulano, a partire da questi dati, una lettura che ha un peso non trascurabile per la discussione teorica: la consapevolezza algoritmica va

considerata una forza digitale, e la sua assenza ha conseguenze che eccedono la dimensione individuale, riguardando i processi democratici e l'autonomia informativa dei cittadini. Chi non sa che gli algoritmi esistono non è nelle condizioni di interrogarli né di chiedere conto delle loro scelte. Su una linea concettualmente vicina, ma più specifica, si colloca lo studio pubblicato nel 2024 su «Frontiers in Communication» da Eder e Sjøvaag (2024), dedicato in modo specifico al rapporto tra intelligenza artificiale e selezione delle notizie. Le autrici, in un'indagine rappresentativa condotta in Germania nel 2022 su un campione di 1.090 individui, introducono il concetto di *algorithmic literacy* (alfabetizzazione algoritmica) come categoria più articolata rispetto alla semplice *awareness*. Se l'*awareness* coincide con la semplice consapevolezza che gli algoritmi esistono, l'alfabetizzazione algoritmica è una capacità più articolata. Implica capire come funzionano i sistemi di raccomandazione e saper valutare gli effetti che producono nel proprio consumo informativo. Implica anche riconoscere quando un contenuto è arrivato all'utente tramite mediazione automatica anziché tramite una scelta umana esplicita. Il dominio in cui questa distinzione conta di più è quello delle notizie.

Una quota crescente del consumo informativo, soprattutto fra i giovani, passa oggi attraverso piattaforme social in cui l'ordinamento e la visibilità dei contenuti sono determinati dagli algoritmi. Chi non capisce come funzionino questi sistemi è esposto a una forma di asimmetria informativa che si aggiunge a quelle, ormai note nella letteratura, riguardanti l'accesso e la competenza. Anche lo studio di Eder e Sjøvaag documenta una distribuzione diseguale dell'alfabetizzazione algoritmica lungo variabili socio-demografiche classiche, confermando su un campione tedesco quanto emerso da Gran e collaboratori sul campione norvegese.

Se i lavori finora citati si concentrano sulla dimensione soggettiva del divario (ciò che gli utenti sanno o non sanno degli algoritmi) uno studio pubblicato nel 2024 su «Communication Research» da Shi e Li sposta lo sguardo sulla dimensione oggettiva, mostrando come gli algoritmi stessi producano esiti differenziati a seconda del profilo degli utenti. Gli autori hanno sottoposto l'algoritmo di raccomandazione di Douyin (la versione cinese di TikTok) a un test automatizzato, impiegando agenti software che simulavano utenti con caratteristiche controllate (Shi & Li, 2024). La variabile indipendente era il modello di smartphone utilizzato, scelto come proxy del livello socioeconomico dell'utente: modelli di fascia alta contro modelli economici.

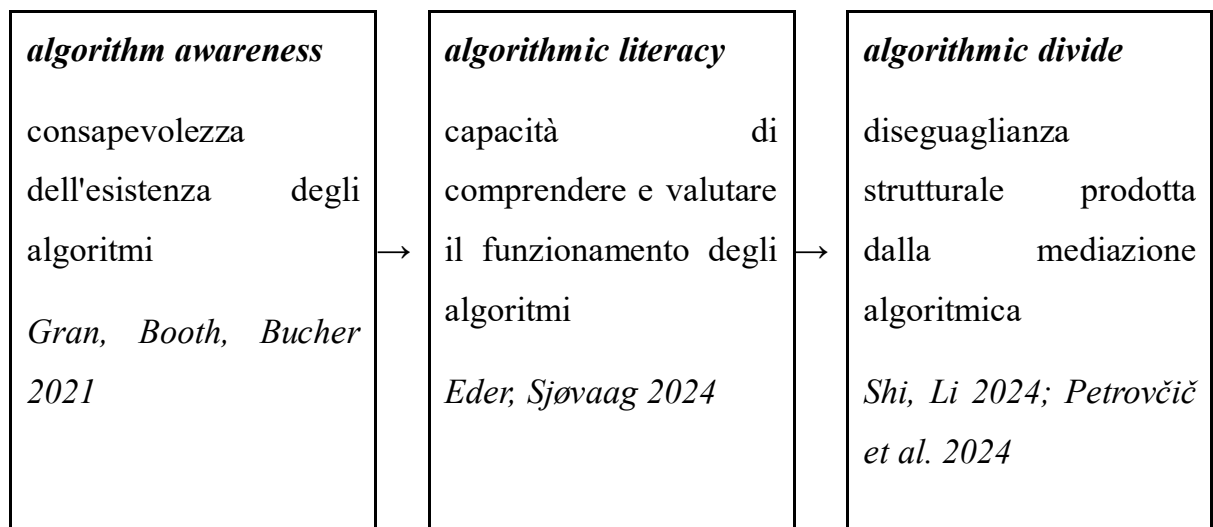
I risultati dell'esperimento sono di particolare rilievo. Anche a parità di comportamento di navigazione, con gli stessi video visualizzati, la stessa permanenza media e un pattern di interazioni equivalente, l'algoritmo finiva per raccomandare agli utenti con smartphone economici una percentuale significativamente minore di video sanitari provenienti da fonti verificate. I contenuti di qualità informativa più elevata, dunque, raggiungevano meno frequentemente gli utenti collocati nei segmenti socio-economicamente svantaggiati. Gli autori dello studio non si pronunciano sul meccanismo specifico all'origine dell'asimmetria: potrebbero essere coinvolte scelte progettuali, oppure effetti non intenzionali del processo di addestramento, o ancora pattern statistici ereditati dai dati storici. Quello che documentano in modo netto è il risultato empirico: l'algoritmo, indipendentemente dalle intenzioni di chi lo ha progettato, finisce per amplificare le disparità socioeconomiche esistenti. Il dato è coerente con la lettura teorica proposta nei paragrafi precedenti, secondo cui la mediazione tecnologica non neutralizza le asimmetrie sociali, ma tende a riprodurle e in alcuni casi a intensificarle.

Sul piano della teorizzazione, un contributo particolarmente utile alla sistemazione del campo è offerto dall'articolo pubblicato nel 2024 su «Information, Communication & Society» da Petrovčič e collaboratori, in cui viene proposto un modello esplicativo che collega la consapevolezza e la conoscenza algoritmica ai tre livelli del *digital divide* classico (Petrovčič et al., 2024). Sulla base di dati empirici raccolti in Slovenia, gli autori mostrano come le disuguaglianze nell'accesso ubiquo alla rete e nelle competenze digitali influenzino la conoscenza algoritmica, la quale a sua volta si riflette sugli usi effettivi della tecnologia e sui benefici che gli utenti ne ricavano.

La rilevanza teorica del lavoro è duplice: sul piano descrittivo, il modello fornisce un quadro integrato in cui le varie componenti del divario (accesso, competenze, benefici, consapevolezza algoritmica) non compaiono come dimensioni isolate ma si dispongono in una catena causale verificata empiricamente. Sul piano interpretativo, l'articolo suggerisce che l'*algorithmic divide* non va considerato come un fenomeno che rompe con la tradizione di ricerca avviata da van Dijk e consolidatasi nell'arco dei tre livelli, bensì come una sua estensione coerente. Le nuove competenze richieste dagli ambienti algoritmici (riconoscere gli algoritmi, comprenderne il funzionamento, valutarli criticamente) si collocano ai vertici della tassonomia delle competenze elaborata da van Dijk e van Deursen, nella zona delle competenze *content-related* e strategiche che la ricerca aveva già identificato come la più diseguale nella distribuzione sociale. Quello che cambia, in altri termini, non è il framework teorico, ma l'oggetto a cui le competenze devono essere applicate.

La progressione concettuale che si è ricostruita nel presente paragrafo, dalla consapevolezza dell'esistenza degli algoritmi alla loro alfabetizzazione critica, fino alla documentazione empirica degli esiti

differenziati che essi producono, può essere sintetizzata nello schema seguente.



Schema 4: La progressione concettuale dell'algorithmic divide

Manca però un passaggio ulteriore. Gli algoritmi di raccomandazione e filtraggio, di cui si è trattato in questo paragrafo, operano su contenuti prodotti da utenti umani o da redazioni professionali: li ordinano e li selezionano, ne regolano la visibilità, ma non li generano. Con la diffusione su larga scala dei modelli di intelligenza artificiale generativa, a partire dalla fine del 2022, si entra in un territorio in cui i sistemi non si limitano più a mediare l'accesso ai contenuti ma li producono direttamente. Questa trasformazione comporta un'ulteriore riconfigurazione delle competenze richieste agli utenti e apre la questione specifica di un AI divide distinto, pur se continuo, rispetto all'*algorithmic divide* finora discusso.

2.6 L'AI divide

La discussione condotta nel paragrafo precedente ha messo a fuoco le disuguaglianze che si generano intorno alla mediazione algoritmica, intesa come capacità dei sistemi automatici di selezionare, ordinare e

filtrare contenuti prodotti da soggetti umani. Due trasformazioni convergenti hanno modificato il panorama nell'ultimo triennio. La prima è l'arrivo dei modelli di intelligenza artificiale generativa sul mercato consumer dalla fine del 2022. La seconda è l'estensione dei sistemi di decisione automatizzata in settori a forte impatto sociale: la valutazione delle candidature lavorative, l'analisi del merito creditizio bancario, alcuni ambiti della diagnostica sanitaria, le valutazioni del rischio in ambito giudiziario sono solo gli esempi più evidenti. Il problema non riguarda più, soltanto, algoritmi che orientano l'accesso degli utenti a contenuti già esistenti. Oggi i sistemi producono direttamente contenuti nuovi, raccomandano percorsi operativi, arrivano a incidere su decisioni che riguardano i diritti delle persone. Questa transizione ha indotto una parte della letteratura più recente a parlare di un AI divide in senso proprio, distinto (pur se in continuità) rispetto *all'algorithmic divide* di cui si è trattato. Il sotto-paragrafo 2.6.1 ricostruisce in che senso l'AI divide rappresenti un fenomeno specifico, non riducibile né al *digital divide* classico né *all'algorithmic divide*. Il 2.6.2 presenta poi la definizione operativa proposta dalla letteratura empirica più recente. Il 2.6.3, infine, introduce la distinzione tra uso consapevole e uso subito, che apre il discorso sul capitale algoritmico e prepara il terreno per i capitoli successivi.

2.6.1 Continuità e discontinuità tra digital divide, algorithmic divide e AI divide

La proliferazione di etichette nella letteratura recente potrebbe far pensare a una moltiplicazione di fenomeni distinti, ciascuno con una propria logica autonoma. La ricostruzione genealogica condotta nei paragrafi precedenti suggerisce un'ipotesi diversa: digital divide, algorithmic divide e AI divide rappresentano stadi successivi di un'unica

traiettoria di stratificazione, in cui ogni nuova ondata tecnologica riconfigura le linee di frattura senza eliminare quelle precedenti. Comprendere su quali piani questi tre stadi presentino elementi di continuità e su quali presentino elementi di discontinuità è la condizione per dare conto in modo rigoroso della loro relazione reciproca, ed è la finalità di questo sotto-paragrafo.

Sul piano della continuità, tre elementi accomunano i tre stadi del divario. Il primo riguarda le variabili socio-demografiche che predicono la posizione degli individui rispetto alle tecnologie: età, livello di istruzione, reddito, collocazione geografica e, in misura minore, genere. Sono variabili che dal rapporto NTIA del 1995 fino allo studio di Wang e collaboratori del 2025 ricorrono con regolarità nei risultati empirici, indipendentemente dalla specifica tecnologia analizzata. Vi è poi il meccanismo causale che il framework di van Dijk ha messo a fuoco: le disuguaglianze tecnologiche non sono prodotte dalla tecnologia in sé, ma derivano dalla distribuzione diseguale delle risorse che precede l'introduzione della tecnologia. Quando la tecnologia si diffonde, si appoggia a una stratificazione che esiste già nella società e ne riproduce le divisioni. Un terzo elemento di continuità sta nel meccanismo di retroazione descritto dalla quinta proposizione della Resources and Appropriation Theory: la partecipazione diseguale alla società, prodotta dall'accesso diseguale alle tecnologie, finisce per rinforzare le disuguaglianze categoriali da cui era originariamente derivata. È il meccanismo cumulativo che van Dijk riconduce all'effetto Matteo formulato da Robert Merton, e che opera identicamente nei tre stadi.

Sul piano della discontinuità, ogni stadio aggiunge invece qualcosa di nuovo. Il digital divide nella sua forma classica si articola lungo tre piani: prima la disponibilità di hardware e di connessione, poi le competenze necessarie per usare gli strumenti, infine la capacità di

tradurre l'uso in benefici concreti per la propria vita. L'algoritmico divide introduce qualcosa di qualitativamente diverso, perché lo spostamento avviene dalla padronanza tecnica alla consapevolezza: la consapevolezza che la tecnologia non opera in modo neutrale, e che dietro a sistemi di selezione e raccomandazione esistono criteri ordinati in modo opaco. Per quanto riguarda l'AI divide, tre componenti specifiche non trovano corrispondenza nei divari precedenti. La prima è produttiva. Mentre gli algoritmi di raccomandazione mediano contenuti già esistenti, i modelli generativi producono direttamente testi e immagini, codice e contenuti audio: il rapporto tra utente e sistema si sposta dal consumo alla produzione, e si apre una questione nuova relativa a chi sia in grado di generare valore attraverso la tecnologia. Esiste poi una componente decisionale: i sistemi di AI stanno assumendo ruoli che eccedono la raccomandazione di contenuti, ed entrano nel dominio delle decisioni formali, quelle che incidono sull'accesso degli individui a opportunità di tipo economico, giuridico e sanitario. Va aggiunto, infine, il grado di opacità. Gli algoritmi di raccomandazione erano già opachi per la maggior parte degli utenti, ma i modelli di *deep learning* alla base dell'AI contemporanea lo sono in misura strutturalmente superiore: al punto che anche chi li progetta non è sempre in grado di ricostruire il ragionamento che ha portato a un determinato output.

Ha senso, dunque, parlare di un divide distinto e non semplicemente di una nuova fase del digital divide? La risposta affermativa si fonda sull'osservazione che i tre elementi appena enunciati (produzione, decisione, opacità) non costituiscono variazioni quantitative di tendenze precedenti, ma componenti qualitativamente nuove del rapporto tra individui e tecnologie. Generare valore attraverso un sistema che produce contenuti non è la stessa cosa che cercare informazioni in rete. Subire una decisione automatizzata su un mutuo o su una candidatura di

lavoro è ben altra cosa rispetto a ricevere un contenuto raccomandato. E interagire con un sistema opaco anche per i suoi sviluppatori comporta un grado di asimmetria che non si ritrova nel navigare una piattaforma dai criteri di ranking ignoti. Le novità non rendono però obsoleto il framework teorico precedente; ne richiedono semmai un'estensione. Le variabili predittive sono ancora le stesse, il meccanismo di retroazione opera identicamente, e anche le competenze che fanno la differenza si collocano nella zona delle *strategic skills* della tassonomia di van Dijk. Ciò che cambia è l'oggetto a cui queste competenze devono essere applicate: non più un'interfaccia da usare o un algoritmo di selezione da riconoscere, ma un sistema che produce e decide. L'AI divide è quindi, allo stesso tempo, in continuità con la tradizione di ricerca sul digital divide e in discontinuità qualitativa con essa.

2.6.2 La definizione operativa di Wang et al. (2025)

La precisazione concettuale appena delineata trova un fondamento empirico nella letteratura più recente. I tre elementi di discontinuità identificati nel sotto-paragrafo precedente non indicano una rottura rispetto al panorama del divario digitale, quanto piuttosto un'intensificazione di tendenze già in atto. L'*algorithmic divide* era costituito da asimmetrie di consapevolezza e alfabetizzazione; l'AI divide aggiunge a queste asimmetrie una componente ulteriore, legata al fatto che i sistemi con cui gli utenti interagiscono hanno acquisito una capacità di azione autonoma, nella generazione e nella decisione, che eccede la logica della mediazione. Restano validi, in questa cornice, i riferimenti teorici consolidati nei paragrafi precedenti: il framework multidimensionale di van Dijk e la stratificazione a tre livelli che ne deriva, oltre alla lettura di Castells sulle asimmetrie strutturali di potere nella *network society*. Ma il campo empirico a cui questi strumenti devono essere applicati è ridefinito.

Il contributo più sistematico alla definizione empirica dell'AI divide è, allo stato attuale della letteratura, l'articolo pubblicato nel 2025 su «New Media & Society» da Wang, Boerman, Kroon, Möller e de Vreese (Wang et al., 2025). Il titolo del contributo, *The artificial intelligence divide: Who is the most vulnerable?* indica con precisione la posta in gioco. Non si tratta di documentare l'esistenza di una nuova forma di disuguaglianza, che gli autori danno per assodata sulla base della letteratura precedente, ma di identificare quali gruppi sociali siano maggiormente esposti e attraverso quali meccanismi.

La definizione operativa che gli autori propongono articola l'AI divide lungo tre assi: le conoscenze (ciò che gli individui sanno dei sistemi di AI, del loro funzionamento, dei loro limiti), le abilità (la capacità effettiva di interagire con tali sistemi in modo efficace, di formulare richieste appropriate, di valutarne criticamente gli output), e le attitudini (il grado di fiducia, l'atteggiamento nei confronti delle potenzialità e dei rischi, la disponibilità a integrare l'AI nelle proprie pratiche quotidiane). L'indagine empirica degli autori, condotta su un campione internazionale, rileva che tutte e tre le dimensioni del divario sono correlate alle variabili socio-demografiche standard. Le più rilevanti risultano l'età, il livello di istruzione e il reddito; il genere pesa anch'esso, anche se in misura più contenuta. Il profilo tipico del soggetto più vulnerabile all'AI divide, secondo Wang e colleghi, è quello di una persona anziana che dispone di basso capitale culturale e che si colloca in segmenti socio-economici meno favoriti. Un quadro che ricalca da vicino quello emerso nei livelli precedenti.

2.6.3 Uso consapevole e uso subito: verso il capitale algoritmico

Un aspetto dell'analisi di Wang e colleghi va però messo in evidenza, perché segna una differenza di rilievo rispetto alle configurazioni classiche del divario digitale. Nei tre livelli ricostruiti nel paragrafo 2.4 la distinzione passava, con una certa nitidezza, tra chi utilizzava la tecnologia e chi ne restava escluso. Nel caso dell'AI lo scenario è cambiato: l'esclusione completa è ormai un'eccezione. La quasi totalità degli utenti connessi a Internet ha già interagito con sistemi di intelligenza artificiale, anche se in molti casi senza accorgersene. Avviene quando uno smartphone propone parole nella tastiera predittiva, quando un motore di ricerca personalizza i risultati attraverso un modello linguistico integrato, quando un sistema operativo offre funzioni generative per la sintesi di testi o per la modifica di immagini. E avviene, sempre più, attraverso interfacce conversazionali esplicite come ChatGPT. La distinzione che oggi conta non è tra chi usa l'AI e chi ne resta fuori. È, semmai, tra chi sa di interagire con un sistema artificiale e chi accetta passivamente i suoi output. È la distanza tra chi sa riconoscere la presenza del sistema e formulare richieste informate, comprendendone i limiti, e chi accetta gli output senza disporre degli strumenti per valutarli. È una distinzione che la letteratura sta iniziando a tematizzare e che ha implicazioni non marginali per il disegno della ricerca empirica a cui i prossimi capitoli fanno riferimento.

La previsione che questa evoluzione potesse condurre a una polarizzazione sociale di nuovo tipo non è, peraltro, estranea al framework classico del divario digitale. Già nell'edizione 2020 del suo volume, prima dell'esplosione mediatica dell'AI generativa, van Dijk dedica alcune pagine conclusive a una scenarizzazione prospettica che richiama in modo significativo la discussione odierna. L'autore, rielaborando una suggestione dello storico Yuval Noah Harari, prospetta la possibilità che

l'intreccio tra intelligenza artificiale e potenziamento tecnologico produca una divaricazione strutturale tra due componenti della popolazione:

«The popular historian Y. N. Harari (2016, 2018) has a vision of the future in which *super humans* are created alongside normal humans. The super humans are the information elite, with not only the best digital skills but also the power to afford, to deal with and to control advanced data worn on or as implants in the body. Normal humans would become superfluous or irrelevant once most jobs are fully or largely automated and robotized.» (van Dijk, 2020, p. 129).

Il giudizio di van Dijk rispetto a questa visione è sobrio, ma non liquidatorio: secondo l'autore, una versione attenuata della distopia descritta da Harari può effettivamente diventare realtà. Non nel senso di una separazione biologica tra super-umani e umani normali, che resta nell'ordine della speculazione narrativa, bensì nel senso di una stratificazione sociale in cui chi dispone delle competenze e delle risorse necessarie per utilizzare criticamente l'AI accumula vantaggi cumulativi, mentre chi ne è sprovvisto accumula svantaggi. È lo stesso meccanismo di retroazione che van Dijk aveva formulato nella quinta proposizione della sua teoria, applicato al nuovo contesto tecnologico.

Arrivati a questo punto del percorso, diventa possibile recuperare una categoria che è stata via via emersa nel corso della ricostruzione genealogica e che merita di essere formulata in modo esplicito. Per *capitale algoritmico* si può intendere l'insieme delle conoscenze, delle competenze e delle disposizioni che permettono a un individuo di interagire con i sistemi algoritmici e di intelligenza artificiale in modo consapevole e strategico. Non si tratta di un concetto sostitutivo rispetto al capitale digitale elaborato dalla letteratura sulle diseguaglianze digitali, ma di una sua articolazione più specifica, che mette a fuoco le risorse richieste dall'ambiente mediato dall'AI. Letta attraverso la cornice di

Castells richiamata nel paragrafo precedente, questa categoria assume una valenza ulteriore. Il capitale algoritmico è anche un fattore di posizionamento entro una configurazione di potere, oltre che una risorsa individuale: la capacità di comprendere i sistemi e di orientarne l'output diventa, in questa prospettiva, una forma di vantaggio strutturale. È peraltro un'articolazione che si lega in modo diretto al framework sociologico del Capitolo 1: dove Bourdieu, ripreso da Collins, aveva mostrato come il capitale culturale trasmesso nella famiglia di origine condizioni le traiettorie scolastiche e professionali, il capitale algoritmico costituisce oggi una forma analoga di risorsa ereditabile, distribuita in modo diseguale e capace di tradursi in vantaggi cumulativi.

La distribuzione diseguale del capitale algoritmico produce, sulla base di quanto documentato dalla letteratura richiamata nei paragrafi precedenti, conseguenze che si manifestano in almeno tre dimensioni della vita sociale contemporanea. La prima riguarda la visibilità: i sistemi algoritmici determinano in misura crescente chi appare nei flussi informativi e chi ne resta escluso, chi viene raccomandato e chi marginalizzato. C'è poi la dimensione del lavoro: l'integrazione dell'AI nei processi produttivi sta ridefinendo le competenze richieste e le traiettorie professionali, con effetti differenziati a seconda del profilo di partenza di ciascun lavoratore. La terza dimensione è quella del potere: la capacità di comprendere e utilizzare criticamente i sistemi si traduce in un'asimmetria decisionale che coinvolge non soltanto il rapporto tra utenti e piattaforme, ma più ampiamente la distribuzione dell'agency nelle relazioni sociali mediate dalla tecnologia. Sono le tre dimensioni che il Capitolo 3 affronterà in modo analitico, applicando il framework teorico ricostruito in questo capitolo alle manifestazioni concrete del divario nelle pratiche contemporanee.

2.7 Una traiettoria di lungo periodo, dall'accesso all'algoritmo

Il percorso condotto nelle pagine precedenti ha ricostruito venticinque anni di ricerca sul divario digitale in una forma genealogica. Si è partiti dal contesto in cui il termine *digital divide* nasce, negli Stati Uniti della metà degli anni Novanta, quando la frattura tra chi possedeva un computer connesso a Internet e chi non lo possedeva poteva essere descritta, senza troppi scrupoli teorici, nei termini binari di una separazione tra due popolazioni. Si è poi mostrato come quella lettura, efficace nella sua immediatezza, abbia progressivamente rivelato i propri limiti di fronte all'evoluzione del campo empirico: la critica della metafora condotta da van Dijk ha isolato quattro malintesi ricorrenti (la dicotomia, l'invalidabilità, l'assolutezza, la staticità) che una lettura matura del fenomeno deve superare.

A partire da questa critica si è presentato il modello multidimensionale elaborato dall'autore, in cui le diseguaglianze digitali vengono collegate, attraverso il concetto intermedio di risorse, alle disuguaglianze sociali preesistenti, in una catena causale circolare che produce effetti di retroazione. Su questa base è stata ripercorsa la stratificazione storica che la letteratura ha elaborato nei tre livelli oramai canonici: l'accesso fisico, le competenze, i benefici concreti. Ciascun livello non ha sostituito il precedente, ma vi si è aggiunto, configurando una stratificazione sincronica in cui forme diverse di divario coesistono nelle stesse società. A questa tripartizione consolidata si sono poi aggiunti, nell'ultimo decennio, due stadi ulteriori: l'*algorithmic divide*, generato dalla mediazione automatica dei contenuti, e l'AI divide, emerso più recentemente e legato all'interazione con sistemi generativi e decisionali dotati di margini crescenti di autonomia.

Il risultato concettuale della ricostruzione può essere formulato in modo sintetico: ogni nuova ondata tecnologica ha prodotto un

riposizionamento del divario, non il suo superamento. Le linee di frattura si sono spostate progressivamente: dalla disponibilità di hardware alla qualità della connessione, poi dalle competenze operative a quelle strategiche, infine dagli esiti nella vita offline alla consapevolezza dei sistemi automatici. In tutti i casi, e in modo più o meno riconoscibile, le variabili socio-demografiche classiche (istruzione e reddito anzitutto, e poi età e collocazione geografica) restano determinanti. Il framework di van Dijk, con le sue cinque proposizioni, offre la chiave interpretativa di questa continuità. Le disuguaglianze sociali generano distribuzioni diseguali di risorse; queste producono a loro volta accessi diseguali alle tecnologie, e l'accesso diseguale finisce per rinforzare le disuguaglianze di partenza. È un meccanismo circolare che ogni nuova fase tecnologica ripropone, modificando i contenuti ma non la logica di fondo, che si lega in modo diretto al framework sociologico ricostruito nel Capitolo 1, dove gli stessi meccanismi di riproduzione delle disuguaglianze erano stati identificati come tratto strutturale delle società moderne.

La rilevanza di questa costruzione teorica non si esaurisce sul piano concettuale. I fenomeni descritti trovano corrispondenza nei dati che le istituzioni europee raccolgono in modo sistematico. Il Digital Economy and Society Index, elaborato dalla Commissione Europea dal 2014 e dal 2023 confluito nel rapporto sullo stato del Digital Decade, fornisce una fotografia della situazione degli Stati membri sui principali indicatori di digitalizzazione: connettività, competenze digitali, integrazione delle tecnologie nelle imprese, servizi pubblici digitali (European Commission, 2023). Il dato più rilevante per la discussione appena conclusa riguarda le competenze digitali della popolazione. Secondo il rapporto 2023, soltanto il 54% dei cittadini dell'Unione Europea di età compresa tra i 16 e i 74 anni possiede competenze digitali considerate di base. Si tratta di una percentuale che, a una lettura superficiale, può suggerire che il problema

dell'accesso sia ormai marginale; a una lettura più attenta, indica invece che quasi la metà della popolazione adulta del continente non ha raggiunto la soglia minima necessaria per una partecipazione piena alla società digitale.

La posizione dell'Italia, all'interno di questo quadro, si colloca al di sotto della media europea, in una zona intermedia del continente per quanto riguarda le competenze digitali di base e in posizione più arretrata sulle competenze avanzate. Il dato assume un rilievo particolare se letto attraverso le categorie elaborate nel capitolo. Il nostro paese si trova, in altri termini, in una condizione di stratificazione simultanea dei divari: mentre segmenti consistenti della popolazione non hanno ancora consolidato le competenze di primo e secondo livello, il panorama tecnologico introduce già le sfide del terzo livello, dell'*algorithmic divide* e dell'AI divide. Si stratificano in modo sincronico, senza che gli stadi precedenti siano stati superati. È una condizione che rende il contesto italiano un campo di indagine particolarmente interessante per la ricerca sui nuovi divari, perché vi si osservano in modo contemporaneo fenomeni che in paesi più digitalizzati si sono susseguiti in fasi distinte.

La genealogia ricostruita nel presente capitolo può essere sintetizzata, in forma tabellare, nello schema seguente, che riprende e amplia la rappresentazione già proposta al termine del paragrafo 2.4 estendendola agli stadi più recenti.

Stadio	Periodo	Focus	Autori di riferimento
Primo livello	1995–2003	Accesso fisico: possesso di dispositivi e connessione	NTIA 1995; Norris 2001
Secondo livello	2004–oggi	Competenze e modalità d'uso	Hargittai 2002; Warschauer 2003; Selwyn 2004; van Deursen & van Dijk 2011
Terzo livello	2010–oggi	Benefici concreti (outcomes) nella vita offline	Robles Morales et al. 2010; Helsper & van Dijk 2012; van Deursen & Helsper 2015
<i>Algorithmic divide</i>	2020–oggi	Consapevolezza e alfabetizzazione algoritmica; esiti differenziati della mediazione automatica	Gran, Booth, Bucher 2021; Eder & Sjøvaag 2024; Shi & Li 2024; Petrovčič et al. 2024
AI divide	2023–oggi	Conoscenze, abilità e attitudini nell'interazione con l'AI generativa e decisionale	Wang et al. 2025

Schema 5: La stratificazione del divario digitale dal 1995 al 2025

Lo schema va letto, ancora una volta, nella consapevolezza che gli stadi non si escludono a vicenda, ma si stratificano in una configurazione sincronica. La popolazione contemporanea delle società avanzate sperimenta simultaneamente tutti e cinque gli stadi. Alcuni segmenti sono

ancora esclusi dall'accesso fisico, altri vi accedono ma sono privi delle competenze d'uso necessarie. Una terza fascia utilizza la tecnologia senza riuscire a tradurla in benefici concreti, mentre un'altra ancora è inconsapevole della mediazione algoritmica. Ampie aree della popolazione, infine, si trovano impreparate all'interazione con l'intelligenza artificiale. Ogni individuo occupa, in questa configurazione multidimensionale, una posizione specifica che il framework teorico presentato consente di analizzare con precisione.

Il quadro così ricostruito fornisce le coordinate teoriche entro cui la tesi si muoverà nel prossimo capitolo. Se il Capitolo 1 ha costruito la cornice sociologica generale e il presente capitolo ha articolato la genealogia del divario digitale fino all'AI divide, il Capitolo 3 affronterà le manifestazioni concrete del divario algoritmico e dell'AI divide nelle tre dimensioni della vita sociale: la visibilità nei flussi informativi mediati dalle piattaforme digitali, il lavoro nei contesti trasformati dall'integrazione dell'AI, e il potere nelle relazioni asimmetriche tra utenti, piattaforme e sistemi automatizzati. A ciascuna di queste dimensioni corrispondono manifestazioni empiriche, categorie interpretative e implicazioni critiche specifiche, che verranno esaminate applicando le categorie teoriche qui elaborate al contesto della ricerca contemporanea.

Capitolo 3. Le tre dimensioni dell'AI divide: visibilità, lavoro e potere

3.1 Introduzione

Il Capitolo 2 si è concluso con un'affermazione che vale la pena richiamare in apertura. L'AI divide non rappresenta una discontinuità teorica rispetto al divario digitale che la sociologia ha studiato a partire dagli anni Novanta, ma l'esito coerente di una traiettoria di lungo periodo che si innesta sulle dinamiche di stratificazione sociale ricostruite nel Capitolo 1. Le competenze strategiche di cui parla van Dijk, l'effetto Matteo mertoniano applicato al digitale, il divario di rendimento documentato da van Deursen e Helsper: sono tutti elementi che concorrono a definire un quadro in cui l'introduzione massiva dell'intelligenza artificiale, lungi dal cancellare le disuguaglianze preesistenti, le riarticola in forme nuove e più stratificate.

Resta tuttavia da chiarire un punto che il capitolo precedente ha solo anticipato: in cosa si articola, concretamente, l'AI divide oggi? Quali sono le sue forme di manifestazione nelle vite delle persone, e attraverso quali meccanismi opera? La tesi che orienta il presente capitolo è che l'AI divide si articoli su tre dimensioni interconnesse, che corrispondono a tre soglie di accesso e che si rafforzano reciprocamente: la visibilità, il lavoro e il potere.

La prima dimensione, quella della visibilità, ha a che fare con il modo in cui gli algoritmi di raccomandazione selezionano i contenuti che vediamo nei flussi informativi. Influenza chi viene proposto agli altri utenti e chi resta invece nell'ombra, con conseguenze dirette sulla possibilità di partecipare al dibattito pubblico e di muoversi in un universo

informativo non chiuso. Sul fronte del lavoro, le trasformazioni indotte dall'AI ridefiniscono la posizione che ciascuno occupa nei mercati occupazionali. Da questa posizione dipende, in larga parte, la capacità di mantenere o migliorare la propria condizione economica negli anni della transizione algoritmica. La terza dimensione, infine, riguarda il potere. Si tratta del controllo sui sistemi che producono decisioni automatizzate: chi possiede gli strumenti per comprenderli e governarli ha un vantaggio strutturale rispetto a chi può solo subirne le conseguenze.

Le tre dimensioni non sono giustapposte casualmente. Tra di esse esiste una relazione di rinforzo cumulativo che ricorda da vicino il meccanismo di retroazione formulato da van Dijk nella quinta proposizione della *Resources and Appropriation Theory*. Tra la visibilità e il lavoro il legame opera attraverso le competenze: chi non è visibile nello spazio informativo ha meno strumenti per partecipare al dibattito pubblico, e quindi meno occasioni di costruire le competenze che consentirebbero di leggere criticamente i contenuti algoritmici. La relazione tra lavoro e potere passa invece per le risorse: chi perde nel mercato del lavoro per via dell'automazione dispone di meno risorse economiche per acquisire le competenze necessarie a un uso consapevole dell'AI, e tende a subirla anziché governarla. Resta poi la connessione tra potere e visibilità, che è la più diretta delle tre: chi non controlla i sistemi algoritmici subisce decisioni opache che incidono direttamente sulla propria visibilità e sulle proprie opportunità lavorative. Il circuito si chiude, e produce una disuguaglianza cumulativa che è la versione contemporanea dell'effetto Matteo discusso nel Capitolo 2.

La selezione di queste tre dimensioni richiede una giustificazione esplicita. Altre dimensioni dell'AI divide sono presenti nella letteratura recente, dalla privacy all'identità digitale, dalla sicurezza dei dati alla sostenibilità ambientale dei modelli generativi. Per concentrarci proprio

su visibilità, lavoro e potere ci sono ragioni convergenti su tre piani. Sul piano bibliometrico, si tratta dei tre ambiti più studiati nella letteratura empirica sull'AI divide pubblicata tra il 2020 e il 2025, come emerge dalle *systematic review* più recenti. Sul piano concettuale, ciascuna delle tre dimensioni corrisponde a una soglia chiaramente identificabile (e quindi misurabile) della partecipazione individuale all'ambiente algoritmico. Sul piano sociologico, infine, queste sono le categorie attraverso cui la sociologia classica ha tradizionalmente articolato il discorso sulla stratificazione: applicarle al contesto dell'intelligenza artificiale permette di mantenere continuità con la tradizione, senza perdere di vista la novità del fenomeno. Le dimensioni qui non discusse, come la privacy o l'identità digitale, sono trattate trasversalmente nei tre paragrafi, là dove si intersecano con i temi principali.

Il capitolo procede in quattro tempi. Nel paragrafo 3.2 si discute la dimensione della visibilità, partendo dalla formulazione originaria della *filter bubble* proposta da Pariser fino alle *systematic review* più recenti. Il discorso viene poi ancorato al contesto italiano attraverso l'indagine dell'AGCOM sui meccanismi cognitivi che sottendono al consumo di informazione, e si chiude con un cenno alla declinazione del problema della visibilità nel regime dell'AI generativa, a partire dalla critica formulata da Bender e collaboratrici ai modelli linguistici di grandi dimensioni. Sul lavoro è dedicato il paragrafo 3.3, che presenta il quadro internazionale dell'OECD, il caso italiano documentato dai working paper dell'INAPP, e l'evidenza empirica più recente sull'impatto dell'AI generativa sul lavoro freelance ricostruita da Teutloff e colleghi. Il potere è il tema del paragrafo 3.4, che muove dalla critica della *surveillance capitalism* elaborata da Zuboff, attraversa il problema dell'opacità algoritmica nei modelli a black box, e arriva alla risposta regolatoria italiana ricostruita da Giannelli a partire dall'AGCOM. Una riflessione

sulle implicazioni del problema della trasparenza nel contesto dei sistemi generativi di nuova generazione chiude il paragrafo. Il paragrafo 3.5 conclude il capitolo, raccogliendo le tre dimensioni in una sintesi che prepara le considerazioni finali della tesi. Un'ultima considerazione di metodo. Il presente capitolo non pretende di esaurire il fenomeno dell'AI divide, che resta un campo di ricerca in trasformazione rapidissima e dai contorni ancora in via di definizione. L'obiettivo è offrire una mappa concettuale delle sue principali dimensioni, costruita su fonti recenti e calibrata sul contesto italiano dove la documentazione disponibile lo consente. Ciascuno dei tre paragrafi si chiude con un'estensione del problema al regime dell'AI generativa, che costituisce la frontiera più recente e meno consolidata del dibattito, e che merita una considerazione mirata nei punti dove il framework concettuale del paragrafo trova diretta applicazione. Le tre dimensioni vanno intese come prospettive di lettura del fenomeno, non come compartimenti stagni: chi si occupa di visibilità incontra inevitabilmente questioni di potere, chi si occupa di lavoro incontra questioni di visibilità professionale, chi si occupa di potere incontra questioni di accesso alle competenze. La separazione analitica serve a rendere il discorso più trattabile, non a suggerire che le tre dimensioni operino in modo indipendente.

3.2 Visibilità: chi appare, chi scompare nei flussi informativi

La prima dimensione dell'AI divide riguarda la visibilità informativa. Nell'ambiente digitale contemporaneo l'informazione disponibile è virtualmente infinita. L'attenzione individuale, invece, resta una risorsa scarsa. Tra queste due grandezze, sproporzionate per natura, si frappongono i sistemi automatici di raccomandazione e di filtraggio: selezionano i contenuti, li ordinano in base a un calcolo di rilevanza per ciascun utente, e infine li propongono all'attenzione. La selezione che ne

risulta non è neutra. Dipende dalle scelte progettuali delle piattaforme, dai loro modelli di business orientati al coinvolgimento dell'utente, e da una raccolta continua di dati comportamentali sull'utenza. Il risultato è che ciascun utente vive in un universo informativo personalizzato, con conseguenze rilevanti su ciò che sa, su ciò che pensa e su come partecipa al dibattito pubblico.

La rilevanza sociologica del fenomeno si coglie nel momento in cui si osserva che la personalizzazione non distribuisce in modo neutrale i suoi effetti. Per gli utenti che dispongono di una alfabetizzazione digitale avanzata e che mantengono un consumo informativo diversificato, gli effetti di chiusura risultano attenuati. Sono in grado di attivare strategie di compensazione: confrontano le fonti tra loro, cercano attivamente prospettive divergenti, e riconoscono quando il filtro algoritmico sta limitando la varietà di ciò che vedono. Negli utenti che non dispongono di queste risorse, la selezione algoritmica viene invece accettata come ambiente di default, e i suoi effetti vengono assorbiti senza possibilità di mediazione critica. È la traduzione, nello spazio della visibilità, del divario di rendimento informativo che il Capitolo 2 aveva ricostruito a partire dagli studi di van Deursen e Helsper: una stessa esposizione produce esiti diversi in funzione della posizione sociale e culturale di chi la riceve. Le pagine che seguono ricostruiscono il dibattito su questo fenomeno secondo quattro piani complementari: l'effetto algoritmico della personalizzazione (paragrafo 3.2.1), la dinamica sociale e cognitiva della chiusura informativa (paragrafo 3.2.2), i meccanismi cognitivi che ne facilitano l'efficacia nel contesto italiano (paragrafo 3.2.3), e la declinazione del problema nel regime dell'AI generativa (paragrafo 3.2.4).

3.2.1 Filtri algoritmici e personalizzazione: la filter bubble

Il termine *filter bubble* è stato introdotto nel 2011 da Eli Pariser, attivista e cofondatore di MoveOn.org, in un intervento divulgativo tenuto a TED e nel libro omonimo pubblicato lo stesso anno presso Penguin (Pariser, 2011). L'argomento centrale di Pariser muoveva da un'osservazione personale: i contenuti politicamente conservatori che l'autore seguiva regolarmente su Facebook erano progressivamente scomparsi dal suo *feed* senza alcuna comunicazione esplicita, in conseguenza del fatto che l'algoritmo di selezione aveva rilevato un'interazione più frequente con i contenuti di area liberale. La *filter bubble* veniva definita come «your own personal, unique universe of information that you live in online» (Pariser, 2011), un ambiente informativo personalizzato in cui l'utente non sceglie i criteri di selezione e non vede ciò che viene escluso.

Pariser identificava due caratteristiche dell'ambiente algoritmico che ne fanno qualcosa di diverso dai meccanismi tradizionali di selezione editoriale. La prima è l'invisibilità del filtro: il lettore di un quotidiano sa che la prima pagina è il prodotto di scelte editoriali esplicite, mentre l'utente di un servizio di raccomandazione algoritmica non sa quali criteri abbiano determinato l'ordine in cui i contenuti gli sono apparsi. La seconda è la personalizzazione individuale: la prima pagina del quotidiano è la stessa per tutti i lettori, mentre il *feed* algoritmico è strutturalmente diverso da utente a utente. Pariser sintetizzava entrambe le caratteristiche in una formula che è poi diventata canonica nella letteratura successiva: «there is no standard Google anymore» (Pariser, 2011). Il punto cruciale, che l'autore ha riproposto in più occasioni, riguarda il passaggio storico in atto, descritto come «a passing of the torch from human gatekeepers to algorithmic ones» (Pariser, 2011): la sostituzione dei tradizionali *gatekeeper* giornalistici con sistemi automatici di selezione che, secondo

Pariser, non avrebbero ancora incorporato l'etica professionale e la responsabilità civica che la cultura editoriale aveva consolidato nel corso del Novecento.

La formulazione di Pariser ha avuto un'eco straordinaria nel dibattito pubblico e nella ricerca accademica successiva, ma è stata sottoposta a una verifica empirica articolata che, a oltre un decennio di distanza, restituisce un quadro più sfumato di quello che il termine *filter bubble* sembrava promettere. La più recente *systematic review* sul tema, condotta da Hartmann, Wang, Pohlmann e Berendt (2025) sul *Journal of Computational Social Science*, ha analizzato centoventinove studi peer-reviewed pubblicati tra il 2014 e il 2023, selezionati a partire da un corpus iniziale di oltre millesettecento contributi attraverso la metodologia PRISMA 2020. Il risultato è significativo per due ragioni. La prima riguarda l'orientamento generale della letteratura: settantasei studi su centoventinove confermano l'esistenza di pattern di omofilia e di chiusura informativa nelle piattaforme social, ventisette restituiscono risultati misti e ventisei non trovano evidenza significativa del fenomeno. La maggioranza degli studi sostiene quindi l'ipotesi della *filter bubble*, ma una minoranza non trascurabile la mette in discussione. La seconda ragione è di ordine metodologico, ed è il contributo originale più importante della review: i risultati dipendono in modo sistematico dal metodo utilizzato. Gli studi basati su *trace data*, cioè sull'analisi dei dati comportamentali raccolti dalle piattaforme, e quelli che utilizzano metodi di scienza sociale computazionale tendono a confermare l'ipotesi della *filter bubble*; gli studi basati su survey e sull'analisi dell'ambiente mediatico complessivo dell'individuo tendono a smentirla (Hartmann et al., 2025).

Questo risultato metodologico è centrale per la discussione che si sta sviluppando, perché chiarisce la natura del fenomeno: la *filter bubble* esiste come pattern osservabile nelle tracce digitali degli utenti, ma quando

si chiede direttamente alle persone quali fonti consultano e con quale varietà, il quadro si articola diversamente, perché molti utenti integrano il consumo digitale con altre fonti informative (televisione, radio, conversazioni offline). Il filtro algoritmico c'è, ma agisce all'interno di un ambiente mediatico più ampio che ne attenua, o talvolta ne amplifica, gli effetti. Hartmann e collaboratori segnalano inoltre che l'esistenza e l'intensità del fenomeno variano in modo significativo a seconda della piattaforma considerata. Facebook tende a produrre comunità segregate con esposizione limitata a contenuti trasversali, secondo lo studio di Cinelli e collaboratori (2021, citato in Hartmann et al., 2025). YouTube, attraverso il sistema di raccomandazione, amplifica contenuti di estrema destra e percorsi di radicalizzazione, come documentato da Kaiser e Rauchfleisch (2020, citato in Hartmann et al., 2025). Reddit, grazie a un'architettura che lascia all'utente maggiore controllo sui criteri di selezione, mostra effetti di chiusura meno pronunciati. TikTok e Instagram, le piattaforme oggi più rilevanti per le fasce giovanili, restano significativamente sotto-ricercate, principalmente per le difficoltà di accesso alle API e per la complessità dell'analisi automatica dei contenuti video (Hartmann et al., 2025).

3.2.2 Echo chambers e polarizzazione: la dimensione sociale e cognitiva

Accanto al concetto di *filter bubble*, la letteratura ha sviluppato un secondo concetto che è spesso usato come sinonimo, ma che la ricerca più recente tende a distinguere: l'*echo chamber*. Il termine è stato introdotto da Cass Sunstein nel volume *Republic.com* del 2001, e indica un ambiente in cui le opinioni di un individuo vengono rinforzate dall'interazione ripetuta con altri individui che condividono atteggiamenti simili. La definizione operativa più diffusa nella ricerca empirica recente, ripresa

anche da Hartmann e collaboratori, descrive l'echo chamber come «environments in which the opinion, political leaning, or belief of users about a topic gets reinforced due to repeated interactions with peers or sources having similar tendencies and attitudes» (Cinelli et al., 2021, citato in Hartmann et al., 2025).

La differenza concettuale rispetto alla *filter bubble* è importante, anche se molta letteratura tende ad appiattirla. La *filter bubble* è un effetto strutturale degli algoritmi di personalizzazione: anche un utente che non interagisce attivamente con i contenuti viene collocato dal sistema in un universo informativo selezionato. L'*echo chamber* è invece un fenomeno sociale e cognitivo: emerge dalle scelte degli utenti di esporsi prevalentemente a contenuti coerenti con le proprie convinzioni preesistenti, attraverso meccanismi di omofilia (*homophily*, la tendenza ad associarsi con persone simili), esposizione selettiva (*selective exposure*, la scelta di contenuti allineati con le proprie idee) e bias di conferma (*confirmation bias*, la tendenza a interpretare le informazioni in modo da confermare le opinioni di partenza). La *systematic review* più recente sul tema, condotta da Ahmmad, Shahzad, Iqbal e Latif (2025) e specificamente focalizzata sull'impatto delle bolle informative sui giovani, mantiene questa distinzione in modo netto: «A *filter bubble* refers to the personalization effects of algorithmic curation that limit exposure to diverse viewpoints, while an echo chamber emerges through selective social interaction and confirmation bias, where users engage primarily with like-minded others» (Ahmmad et al., 2025). Nei fatti i due meccanismi tendono a sovrapporsi, perché l'algoritmo amplifica le scelte già orientate dell'utente in un circolo di rinforzo reciproco, ma mantenerli analiticamente distinti consente di cogliere il contributo specifico dell'architettura tecnica e del comportamento sociale alla produzione del fenomeno.

La review di Ahmmad e collaboratori, basata sull'analisi di trenta studi peer-reviewed pubblicati tra il 2015 e il 2025, identifica tre risultati ricorrenti nella letteratura. L'amplificazione strutturale è il primo: i sistemi di raccomandazione delle piattaforme producono in modo sistematico un'omogeneizzazione ideologica dei contenuti proposti, riducendo l'esposizione incidentale a punti di vista differenti. Per gli autori il bias algoritmico è quindi una forza strutturale, e non un dettaglio tecnico marginale: «a structural driver shaping everyday media diets» (Ahmmad et al., 2025). C'è poi la questione dell'agency degli utenti, in particolare dei giovani. La ricerca documenta che molti adolescenti e giovani adulti hanno una consapevolezza parziale ma reale del funzionamento degli algoritmi, e attivano strategie di diversificazione del proprio consumo informativo: seguono account di posizioni divergenti, utilizzano piattaforme multiple. Questa agency è però limitata dall'opacità delle architetture algoritmiche e da una *digital literacy* diseguale, che si distribuisce secondo le variabili socio-demografiche classiche già discusse nel Capitolo 2. Va segnalato infine un risultato che dal punto di vista sociologico è il più interessante, perché complica una lettura puramente negativa del fenomeno: oltre a funzionare come dispositivi di chiusura ideologica, le *echo chamber* agiscono anche come spazi di rinforzo identitario e di appartenenza culturale, dove si costruiscono lessici condivisi e forme di socialità subculturale. La polarizzazione politica si accompagna quindi a una funzione di costruzione identitaria, che rende queste comunità attraenti per i loro membri al di là dei contenuti politici veicolati (Ahmmad et al., 2025).

Per quanto riguarda gli esiti delle *echo chamber*, la review di Hartmann e collaboratori discute tre effetti documentati dalla letteratura. La polarizzazione ideologica è il primo: gli utenti collocati in ambienti omogenei tendono a radicalizzare progressivamente le proprie posizioni,

e questo processo è stato documentato in modo particolarmente solido nelle reti politiche statunitensi durante eventi elettorali. C'è poi la questione della *misinformation*, perché le comunità chiuse tendono a far circolare informazioni false con maggiore velocità e con minore esposizione a meccanismi di correzione. Va segnalata anche l'amplificazione dell'estremismo, in particolare nei sistemi di raccomandazione di YouTube, dove a partire da contenuti moderati tendono a essere proposti contenuti progressivamente più radicali: il fenomeno è stato documentato dallo studio di Kaiser e Rauchfleisch (2020) ripreso nella review (Hartmann et al., 2025). I tre effetti non operano in modo indipendente, ma si rafforzano a vicenda, e hanno alimentato negli ultimi anni un'attenzione crescente da parte dei regolatori europei sul tema dei rischi sistemici delle piattaforme di grandi dimensioni.

La connessione con il framework sociologico ricostruito nei capitoli precedenti è diretta. Chi è più esposto agli effetti delle *echo chamber*? La letteratura documenta una distribuzione diseguale che ricalca i tratti già osservati per il divario digitale di terzo livello. La diversità del consumo informativo correla positivamente con l'istruzione e con l'impegno politico, secondo lo studio di Dubois e Blank (2018) ripreso da Hartmann; gli utenti con maggiore capitale culturale attivano strategie di compensazione efficaci, gli altri assorbono in misura più piena gli effetti della chiusura. Si configura quindi una nuova manifestazione del divario di rendimento informativo: una stessa esposizione mediale produce esiti diversi in funzione delle competenze strategiche disponibili. Le competenze *content-related* della tassonomia di van Dijk, e in particolare quelle *strategic*, risultano qui decisive non solo per estrarre benefici concreti dalla rete, ma anche per non subirne gli effetti di chiusura. La visibilità informativa diventa, in questo senso, una funzione delle

competenze culturali e cognitive di partenza, e quindi della posizione sociale di chi accede all'ambiente algoritmico.

3.2.3 Il caso italiano: dispercezioni e meccanismi cognitivi

Il quadro internazionale ricostruito nei sotto-paragrafi precedenti trova nel contesto italiano una specifica articolazione, documentata in modo sistematico dall'AGCOM, l'Autorità per le Garanzie nelle Comunicazioni, nell'indagine conoscitiva «Piattaforme digitali e sistema dell'informazione» avviata nel 2016 con la delibera 309/16/CONS e conclusa nel 2020 con il rapporto «Percezioni e disinformazione. Molto razionali o troppo pigri?» (AGCOM, 2020). L'indagine si distingue per l'approccio metodologico, che combina lo strumento della survey rappresentativa con un disegno sperimentale di psicologia cognitiva: il questionario, somministrato da SWG a un campione rappresentativo di 1.358 individui italiani di età compresa tra i 14 e i 74 anni, abbina rilevazioni socio-demografiche tradizionali a test sulla conoscenza dei fenomeni sociali ed economici e sulla capacità di riconoscere notizie di qualità diversa (AGCOM, 2020, p. 10).

La scelta di rivolgere l'attenzione al lato della domanda, e cioè ai meccanismi cognitivi degli utenti, anziché soltanto al lato dell'offerta degli algoritmi, è significativa per il discorso che si sta sviluppando. Il rapporto si fonda su una constatazione che ha radici nella ricerca psicologica e nell'economia comportamentale, e che AGCOM riconduce esplicitamente alla teoria dualistica dei processi mentali elaborata da Daniel Kahneman in *Pensieri lenti e veloci* (2017). Le scelte di consumo informativo sono mediate da due sistemi cognitivi distinti. Il sistema 1 è quello intuitivo e veloce, basato su routine associative che non richiedono attenzione deliberata. Il sistema 2, di contro, è analitico e lento, e poggia su un

ragionamento controllato. Il dibattito scientifico recente ha proposto due interpretazioni complementari del modo in cui la disinformazione opera. La prima attribuisce l'efficacia dei contenuti falsi alla «pigrizia» cognitiva, cioè al funzionamento prevalente del sistema 1, che non attiva i processi controllati di verifica. La seconda chiama in causa il *motivated reasoning*, una strategia cognitiva in cui il sistema 2 viene attivato selettivamente per «avvalorare le proprie convinzioni e proteggere la propria ideologia politica» (AGCOM, 2020, p. 2). AGCOM non sceglie tra le due interpretazioni, ma le considera complementari: la disinformazione fa leva sia sulla pigrizia sia sul ragionamento motivato, e produce di conseguenza una vulnerabilità che si manifesta lungo tutto l'arco delle competenze cognitive.

Il dato empirico più rilevante che emerge dal rapporto, e che documenta la specificità del contesto italiano, riguarda l'esistenza di una dispercezione sistematica dei fenomeni sociali ed economici. Riprendendo gli studi condotti da IPSOS a partire dal 2012 attraverso la rilevazione *The Perils of Perceptions*, AGCOM osserva che in Italia esiste una distanza notevole tra l'entità reale di alcuni fenomeni (la criminalità, l'immigrazione, lo stato dell'economia, le condizioni sanitarie) e l'impressione che i cittadini hanno della stessa. Si tratta di una distanza che si è dimostrata persistente nel tempo e che mostra una direzione sistematica: «una sistematica propensione a dipingere la realtà in modo peggiore di quanto non sia», che il rapporto definisce «dispercezione negativa» (AGCOM, 2020, p. 9). I cittadini italiani tendono cioè a sovrastimare la portata dei fenomeni problematici e a sottostimare gli indicatori positivi, in un pattern stabile che precede ed eccede il consumo di disinformazione online.

Il rilievo di questo dato per il discorso sulla visibilità algoritmica è duplice. Da un lato, mostra che le bolle informative algoritmiche non

operano nel vuoto: si innestano su un terreno cognitivo già predisposto, in cui le dispercezioni preesistenti orientano la ricerca selettiva di contenuti coerenti con la propria visione del mondo, e quindi facilitano la formazione delle *echo chamber* dal lato della domanda. Dall'altro lato, mostra che le dispercezioni rendono più efficace la disinformazione, perché «rendono meno riconoscibili i fenomeni di disinformazione» e perché le stesse strategie di manipolazione possono sfruttare le false percezioni dei cittadini «indirizzandosi su di esse alimentandole ulteriormente» (AGCOM, 2020, p. 3). Si configura un circuito cumulativo in cui dispercezione cognitiva, selezione algoritmica e disinformazione strategica si rafforzano a vicenda, producendo una vulnerabilità della sfera pubblica che il rapporto definisce esplicitamente in termini di rischio democratico: una società in cui la disinformazione si innesta su false percezioni della realtà «rischia di diventare più vulnerabile e gli effetti della disinformazione rischiano di amplificarsi» (AGCOM, 2020, p. 3).

La rilevanza di questa analisi per la tesi sull'AI divide è di ordine strutturale, non soltanto contestuale. Il rapporto AGCOM dimostra che la visibilità informativa nelle piattaforme italiane non si distribuisce in modo neutro tra i cittadini, ma riproduce e amplifica le diseguaglianze cognitive di partenza, che a loro volta correlano con le variabili socio-demografiche classiche già osservate nel Capitolo 2. Chi parte da una posizione cognitiva di vulnerabilità, definita dalla combinazione di dispercezioni negative, scarsa attivazione del sistema 2 e ragionamento motivato in direzione delle proprie convinzioni preesistenti, è anche chi subisce in misura maggiore gli effetti di chiusura prodotti dagli algoritmi di raccomandazione. La distribuzione del rischio non è uniforme: si stratifica seguendo le linee di frattura che il framework sociologico ricostruito nei capitoli precedenti ha imparato a riconoscere. La dimensione della visibilità non è soltanto una questione di accesso ai contenuti, ma una questione di rendimento

cognitivo dell'accesso, che si lega in modo diretto al divario di terzo livello formulato da van Deursen e Helsper e all'effetto Matteo mertoniano applicato al consumo informativo.

3.2.4 Visibilità nell'AI generativa: dati di training e amplificazione dei bias

I tre sotto-paragrafi precedenti hanno discusso la visibilità nei sistemi di raccomandazione algoritmica, che operano selezionando contenuti già esistenti e proponendoli all'utente in ordine personalizzato. L'avvento dell'AI generativa, con la diffusione pubblica dei modelli linguistici di grandi dimensioni a partire dal rilascio di GPT-3 nel 2020, inizialmente accessibile via API agli sviluppatori, e soprattutto dall'apertura al grande pubblico di ChatGPT nel novembre 2022, introduce un regime ulteriore: i contenuti non vengono soltanto selezionati, ma *prodotti automaticamente*. La domanda sulla visibilità si sposta di conseguenza: non riguarda più soltanto cosa l'utente vede tra i contenuti preesistenti, ma anche quali voci, quali culture, quali prospettive entrano nei testi che la macchina genera per lui.

Il riferimento canonico per questo passaggio è l'articolo di Bender, Gebru, McMillan-Major e Shmitchell (2021) pubblicato negli atti della conferenza ACM FAccT 2021 con il titolo *On the Dangers of Stochastic Parrots*. Il contributo identifica una serie di rischi strutturali dei *Large Language Models* (LLM), di cui il più rilevante per il discorso sulla visibilità riguarda la composizione dei dati di addestramento. Le autrici documentano che i grandi *dataset* utilizzati per addestrare i modelli linguistici, ottenuti per la maggior parte da raccolte testuali estratte dal web, «overrepresent hegemonic viewpoints and encode biases potentially damaging to marginalized populations» (Bender et al., 2021, p. 610). Le

ragioni di questa sovrarappresentazione sono di tre ordini. In primo luogo, l'accesso a internet non è distribuito uniformemente nella popolazione mondiale, e la composizione dei contenuti riflette demografie specifiche. In secondo luogo, le piattaforme da cui i dati vengono raccolti hanno profili sbilanciati: gli utenti statunitensi di Reddit, sorgente di buona parte dei testi usati per addestrare GPT-2 e GPT-3, sono per il 67% uomini e per il 64% sotto i trent'anni, secondo i dati Pew Research citati nello studio (Bender et al., 2021, p. 614). In terzo luogo, i filtri applicati per pulire i dataset eliminano sistematicamente termini propri delle comunità marginalizzate, riducendo ulteriormente la loro rappresentazione nei modelli.

Il termine *stochastic parrot*, coniato dalle autrici per descrivere gli LLM, è diventato categoria standard nella letteratura critica sull'AI generativa. La metafora descrive i modelli linguistici come sistemi che producono linguaggio per ricombinazione probabilistica senza comprensione semantica: «an LM is a system for haphazardly stitching together sequences of linguistic forms it has observed in its vast training data, according to probabilistic information about how they combine, but without any reference to meaning» (Bender et al., 2021, p. 617). La rilevanza per il discorso sulla visibilità è duplice. Da una parte gli LLM non si limitano a riflettere i bias dei dati di addestramento, ma li amplificano e li propagano: «LMs producing text will reproduce and even amplify the biases in their input» (Bender et al., 2021, p. 618). Dall'altra, gli output dei modelli generativi, una volta diffusi nello spazio informativo pubblico, diventano a loro volta dati di addestramento per i modelli della generazione successiva, in un circuito di rinforzo che rischia di consolidare le distorsioni invece di correggerle. È una versione algoritmica e accelerata del meccanismo cumulativo che il Capitolo 2 aveva ricostruito a partire dall'effetto Matteo mertoniano.

Un secondo punto sollevato da Bender e collaboratrici riguarda il modo in cui gli esseri umani interagiscono con gli output degli LLM. La fluidità del linguaggio prodotto induce il lettore a percepirvi un'intenzionalità comunicativa che il sistema non possiede, in quello che le autrici sintetizzano con una formula efficace: «coherence is in fact in the eye of the beholder» (Bender et al., 2021, p. 616). Il rischio è che gli utenti attribuiscano agli output dei modelli una credibilità non giustificata, e che questo porti a una propagazione non controllata delle distorsioni nei testi che la macchina produce. Il problema si lega in modo diretto al discorso sulla competenza algoritmica già discusso nel Capitolo 2: chi non possiede gli strumenti per riconoscere i limiti epistemici degli LLM tende ad accettare i loro output come affidabili. La distribuzione di questa competenza non è neutra. Come per i sistemi di raccomandazione, correla con il livello di istruzione, con l'età, con la familiarità tecnologica, e quindi riproduce le linee di stratificazione sociale che il framework della tesi ha imparato a riconoscere.

Il quadro che emerge dai quattro sotto-paragrafi del 3.2 può essere sintetizzato in due osservazioni conclusive. La prima è che il problema della visibilità asimmetrica si ripresenta a livelli diversi e con meccanismi diversi a ogni nuova ondata tecnologica: dai gatekeeper editoriali dei media broadcast ai sistemi di raccomandazione delle piattaforme social, dalle dispercezioni cognitive degli utenti alle distorsioni dei modelli generativi. Ciascun livello aggiunge una propria specificità al problema senza eliminare quelli precedenti, in un meccanismo di stratificazione che richiama da vicino la genealogia ricostruita nel Capitolo 2. La seconda osservazione è che la distribuzione di questa asimmetria non è neutra: ricalca le linee di stratificazione sociale già discusse nei capitoli precedenti, e tende ad accumularsi in un effetto Matteo che amplifica le disuguaglianze di partenza invece di ridurle. Le successive dimensioni

dell'AI divide, il lavoro e il potere, declinano la stessa logica su piani diversi della vita sociale.

3.3 Lavoro: chi è esposto, chi guadagna, chi perde nei mercati algoritmici

Il paragrafo precedente ha mostrato come l'AI ridisegni l'asimmetria della visibilità informativa, distribuendo in modo non neutrale l'accesso ai contenuti e la rappresentazione delle voci nei flussi mediati dagli algoritmi. La seconda dimensione dell'AI divide riguarda il lavoro. La premessa empirica da cui muovere è una considerazione che converge nella letteratura più recente: l'AI non distrugge l'occupazione in aggregato, perlomeno nel breve periodo, ma ne ridistribuisce le opportunità lungo linee di stratificazione che ricalcano e talvolta ribaltano quelle preesistenti. La differenza strutturale rispetto alle ondate tecnologiche precedenti è che l'AI ha cominciato a investire le occupazioni cognitive ad alta qualifica, fino a poco tempo fa ritenute al riparo dall'automazione. Questo cambiamento di scenario sposta il baricentro del discorso: non più una questione di sostituzione dei lavori manuali ripetitivi, ma una questione di redistribuzione del valore tra profili professionali, livelli di istruzione, posizioni nella distribuzione salariale.

Il dibattito su queste trasformazioni viene ricostruito nelle pagine che seguono su quattro piani complementari. Il primo è il quadro internazionale offerto dal rapporto OECD 2023, oggetto del paragrafo 3.3.1. Il paragrafo 3.3.2 si concentra invece sul caso italiano, documentando l'esposizione delle professioni all'AI grazie allo studio INAPP di Ferri, Porcelli e Fenoaltea. La distribuzione dell'impatto sui salari italiani è oggetto del paragrafo 3.3.3, che riprende il lavoro di Brunetti e Mussida. Chiude il paragrafo 3.3.4 con la declinazione del

problema nel regime dell'AI generativa, attraverso l'evidenza recente di Teutloff e collaboratori. Il filo conduttore consiste nel leggere ciascuno di questi piani come declinazione lavorativa di concetti già introdotti nel Capitolo 2: il divario di rendimento informativo formulato da van Deursen e Helsper, l'effetto Matteo mertoniano nella sua applicazione alla distribuzione delle opportunità, la distinzione tra uso consapevole e uso subito delle tecnologie digitali.

3.3.1 Il quadro internazionale: l'AI come frontiera della trasformazione del lavoro

La fonte istituzionale di riferimento per il dibattito contemporaneo sulle relazioni tra AI e lavoro è il rapporto *OECD Employment Outlook 2023: Artificial Intelligence and the Labour Market*, curato da Andrea Bassanini e Stijn Broecke. Si tratta della prima edizione della pubblicazione annuale dell'OECD interamente dedicata all'AI, e in particolare ai suoi effetti sull'occupazione, sulla qualità del lavoro e sull'inclusione. Il capitolo terzo del volume, di Andrew Green, affronta esplicitamente la questione dell'impatto dell'AI sull'occupazione, mentre il capitolo quarto, di Green con Salvi Del Pero e Verhagen, ne discute gli effetti sulla qualità del lavoro. L'apertura del rapporto colloca l'AI nel quadro delle innovazioni tecnologiche di lungo periodo: «AI is the current technological innovation sparking hopes of rapid productivity gains and stirring fears of job loss» (OECD, 2023, p. 103). La formula è equilibrata e dichiara apertamente che il rapporto non si pronuncia univocamente in senso utopico o distopico, ma costruisce le condizioni di un'analisi empirica.

La novità strutturale dell'AI rispetto alle precedenti ondate di automazione, secondo il rapporto, riguarda il tipo di compiti che la

tecnologia è capace di svolgere. Il rapporto osserva che «AI has made the most progress in its ability to perform non-routine, cognitive tasks, such as information ordering, memorisation and perceptual speed» (OECD, 2023, p. 103). La differenza con il passato è importante: le precedenti rivoluzioni tecnologiche avevano automatizzato compiti routinari, colpendo perciò soprattutto le occupazioni a basse competenze (operai di catena di montaggio, sportellisti bancari, addetti alle catene di produzione). L'AI sta invece automatizzando compiti cognitivi non routinari, e quindi colpisce per la prima volta in maniera significativa le occupazioni ad alte competenze. Il rapporto sintetizza la conseguenza in modo netto: «high-skilled occupations are those most likely to involve non-routine cognitive tasks, and they have therefore been the most exposed to progress in AI. Examples of such occupations include business professionals, managers, chief executives and science and engineering professionals» (OECD, 2023, p. 103). È un ribaltamento del quadro tradizionale sulla relazione tra tecnologia e disuguaglianza, ed è un punto da cui il discorso non può prescindere.

La domanda fondamentale resta quella sull'impatto netto: l'AI produrrà in aggregato più lavoro o meno lavoro? La risposta dell'OECD è di sospensione metodica del giudizio: «Theoretically, the net impact of AI on employment is ambiguous» (OECD, 2023, p. 103). Il rapporto identifica tre meccanismi che operano contemporaneamente e in direzioni opposte: il *displacement effect* (sostituzione di lavoro umano), il *productivity effect* (aumento della domanda di lavoro tramite i guadagni di produttività che riducono i prezzi e quindi aumentano la domanda di beni e servizi), e il *reinstatement effect* (creazione di nuovi compiti e di nuove occupazioni complementari all'AI). Quale dei tre meccanismi prevarrà nel lungo periodo dipende da fattori che la letteratura economica non è oggi in grado di prevedere con precisione. L'evidenza empirica disponibile al

momento della pubblicazione del rapporto suggerisce che, almeno nel breve termine, l'effetto netto dell'AI sull'occupazione aggregata nei paesi OECD è stato limitato: l'AI non è ancora un fattore di distruzione massiccia del lavoro.

Anche se l'effetto aggregato è contenuto, esiste una concentrazione del rischio in alcune categorie occupazionali. Il dato sintetico più significativo del rapporto è il seguente: «On average across OECD countries in the sample, the occupations at the highest risk of automation account for 27% of employment» (OECD, 2023, p. 117). Circa un lavoratore su quattro nei paesi OECD è impiegato in occupazioni esposte ad alto rischio di automazione. Il dato non significa che un quarto dei lavoratori perderà il proprio impiego, ma che un quarto opera in mansioni in cui i compiti svolti hanno un alto grado di sovrapposizione con quelli che l'AI è capace di replicare. Su questa esposizione si innesta un paradosso che è importante per il framework della tesi: Tra i lavoratori più qualificati, pur essendo tecnicamente i più esposti dal punto di vista dell'AI, si registra anche la maggiore capacità di adattamento. Per i lavoratori con basse competenze il rapporto si rovescia: l'esposizione tecnica è minore, ma quando si concretizza la mobilità verso altre posizioni risulta più difficile: «low-skill workers were often at the greatest disadvantage because when their tasks would be automated, they were often the most difficult to retrain or move to another position within the firm» (OECD, 2023, p. 118). Il rischio si distribuisce quindi in modo asimmetrico lungo l'asse delle competenze. Una parte considerevole dell'esposizione effettiva ricade sui lavoratori in posizione intermedia, che combinano un'esposizione operativa elevata con una bassa mobilità.

Sul versante salariale, il capitolo quarto del rapporto documenta un effetto già operante: chi possiede competenze AI riceve un premio sostanziale rispetto a chi non ne possiede: «After controlling for skills

demanded and the local labour market, they find a wage premium of 11% for job postings that require AI skills within the same firm, and 5% within the same firm and job title» (OECD, 2023, p. 131). Il premio salariale per le competenze AI è dell'undici per cento per posti di lavoro analoghi nella stessa azienda, e del cinque per cento per posti dello stesso ruolo nella stessa azienda. Il rapporto sottolinea che questo premio «is higher than the premiums associated with other skills commonly demanded»: l'AI non è una skill come le altre, è una skill che genera un differenziale salariale particolarmente accentuato. È un primo segnale empirico che la disuguaglianza salariale legata all'AI non è una previsione teorica, è un fenomeno già in atto. Resta da chiedersi come questa disuguaglianza si declini nel contesto italiano specifico, dove la trasformazione tecnologica incontra stratificazioni sociali, territoriali e settoriali peculiari. È la domanda del paragrafo successivo.

3.3.2 L'esposizione italiana all'AI: tra rischio di sostituzione e potenziale di complementarità

Il quadro internazionale dell'OECD viene declinato sul caso italiano nel Working Paper n. 125 dell'INAPP pubblicato a settembre 2024 a firma di Valentina Ferri, Rita Porcelli ed Enrico Maria Fenoaltea, con il titolo *Lavoro e Intelligenza artificiale in Italia: tra opportunità e rischio di sostituzione*. L'INAPP è l'ente pubblico di ricerca italiano specializzato nell'analisi delle politiche del lavoro, vigilato dal Ministero del Lavoro; il CREF (Centro Ricerche Enrico Fermi) di Roma è centro di ricerca interdisciplinare a cui afferisce Fenoaltea. L'obiettivo dichiarato dal paper è duplice: «determinare l'esposizione all'Intelligenza artificiale (IA) e l'importanza della complementarità dell'IA nelle attività lavorative quotidiane» (Ferri et al., 2024, p. 4). L'innovazione metodologica consiste nell'adattamento al contesto italiano di due indicatori sviluppati nella

letteratura internazionale: l'indice AIOE (*Ability level AI Occupational Exposure*), costruito da Felten, Raj e Seamans (2021) sul sistema classificatorio statunitense O*NET, e l'indice C-AIOE (AIOE corretto per la complementarità), sviluppato da Pizzinelli e collaboratori (2023) nel Working Paper n. 216 dell'International Monetary Fund. Le 52 attitudini dell'Indagine Campionaria sulle Professioni 2013 di INAPP e ISTAT corrispondono alle 52 abilities di O*NET; l'esposizione calcolata su questa base viene applicata alla Rilevazione Continua sulle Forze di Lavoro 2022 di ISTAT a livello del quinto digit della classificazione delle professioni.

L'applicazione dell'indice AIOE al panorama professionale italiano produce risultati qualitativamente coerenti con quelli emersi dalla letteratura internazionale, ma calibrati sulle specifiche del mercato del lavoro nazionale. Le professioni con esposizione più bassa appartengono al lavoro manuale, fisico e artistico: manovali e personale non qualificato dell'edilizia civile, intonacatori, atleti, conduttori di veicoli a trazione animale, ballerini, acrobati e artisti circensi compaiono nelle prime venti posizioni della classifica delle meno esposte (Ferri et al., 2024, Tabella 2, p. 11). Le professioni con esposizione più alta sono invece concentrate nel lavoro amministrativo e di gestione informativa: addetti al protocollo e allo smistamento di documenti, addetti alle buste paga, uscieri, direttori generali delle amministrazioni dello Stato, tecnici dell'organizzazione del traffico ferroviario, magistrati (Ferri et al., 2024, Tabella 3, p. 12). Il quadro qualitativo conferma quanto già osservato nel rapporto OECD: l'AI italiana attacca i compiti cognitivi non routinari di natura amministrativa e gestionale, lasciando relativamente intatto il lavoro fisico e quello artistico-creativo.

La distribuzione dell'esposizione sul territorio italiano segue le linee della stratificazione storica del paese. Sul fronte regionale i valori massimi di esposizione media si registrano in Lombardia (0,523), seguita

da Lazio (0,519) ed Emilia-Romagna (0,514). All'estremo opposto della classifica si collocano Calabria (0,425), Puglia (0,431) e Sicilia (0,440) (Ferri et al., 2024, p. 13). Anche a livello provinciale la geografia ricalca quella nazionale: Milano (0,568) è in testa, seguita da Bologna (0,559) e Roma (0,543), mentre in fondo alla classifica si trovano Ragusa, Catanzaro e Vibo Valentia. Settorialmente, le attività finanziarie e assicurative (0,83) e i servizi di informazione e comunicazione (0,79) presentano la massima esposizione; agricoltura (0,11) e costruzioni (0,19) la minima (Ferri et al., 2024, Tabella 4, p. 15). L'AI divide italiano si sovrappone quindi a una geografia economica già nota, dove le aree e i settori più digitalizzati e finanziarizzati sono anche quelli più investiti dalla nuova ondata tecnologica. Il quadro che emerge è coerente con quanto discusso nel Capitolo 1 sulle disuguaglianze territoriali del digitale: l'AI non livella, amplifica le asimmetrie preesistenti.

Su un punto specifico Ferri e collaboratori forniscono un'evidenza che merita di essere sottolineata, perché ribalta uno schema interpretativo classico della letteratura sul digital divide. L'esposizione all'AI cresce con il livello di istruzione: 0,698 per chi possiede ITS, laurea o dottorato; 0,548 per il diploma di maturità; 0,209 per chi non ha alcun titolo; 0,190 per chi possiede solo la licenza elementare (Ferri et al., 2024, Tabella 7, p. 16). La regressione OLS conferma il risultato: «i laureati hanno un'esposizione molto alta rispetto a chi non ha conseguito alcun titolo. Questo suggerisce che le competenze acquisite attraverso l'istruzione avanzata sono quelle più esposte all'automazione tramite IA» (Ferri et al., 2024, p. 17). Si tratta di un'inversione rispetto al digital divide tradizionale, che colpiva soprattutto i lavoratori meno istruiti. L'AI non penalizza chi ha meno strumenti culturali, almeno non in via diretta: penalizza chi ha competenze cognitive sofisticate. Anche due altri dati socio-anagrafici meritano attenzione: le donne sono mediamente più esposte degli uomini (0,548

contro 0,450; Tabella 5, p. 15) e i giovani tra i quindici e i ventiquattro anni sono il gruppo meno esposto in assoluto (0,414; Tabella 6, p. 15). Ne risulta una fotografia in cui l'AI investe il lavoro qualificato, femminile e in età lavorativa centrale.

La distinzione concettualmente decisiva del paper non riguarda però l'esposizione in quanto tale, ma il suo incrocio con la complementarità. L'indice C-AIOE corretto per il fattore *theta* permette di distinguere le occupazioni esposte secondo due assi: quelle ad alta esposizione e bassa complementarità (l'AI sostituisce) e quelle ad alta esposizione e alta complementarità (l'AI potenzia). Il dato sintetico finale del paper, ricavato dall'intersezione di questi due indici, è particolarmente potente: «il 23% dei lavoratori in Italia è a rischio di sostituzione [...] poco più di un lavoratore su quattro beneficeranno dell'IA. Complessivamente, circa metà della forza lavoro italiana è coinvolta in qualche misura dall'impatto dell'Intelligenza artificiale, sia in termini positivi che negativi» (Ferri et al., 2024, p. 22). Le cinque professioni più a rischio secondo l'indice corretto sono addetti al protocollo e allo smistamento di documenti, uscieri, addetti alle buste paga, addetti all'informazione nei call center senza funzioni di vendita, e centralinisti (Ferri et al., 2024, p. 26). I lavoratori più qualificati, in particolare i laureati, tendono invece a collocarsi nel quadrante della complementarità: alta esposizione e alto potenziale di potenziamento. Si delinea così una versione italiana dell'effetto Matteo applicato al lavoro: chi è già qualificato sarà ulteriormente potenziato dall'AI, mentre chi possiede competenze intermedie nelle fasce a rischio di sostituzione subirà l'impatto più diretto. Il paper si chiude con un'indicazione di policy: «per i lavoratori a maggior rischio, soprattutto in età più giovane, occorrerebbe pianificare sin da subito strumenti di reskilling che permettano di assicurare un futuro adeguato» (Ferri et al., 2024, p. 26). Le politiche di riqualificazione sono

però soltanto una parte della questione: dall'altra parte resta da capire come l'esposizione all'AI si traduca concretamente in differenziali salariali, e quindi in disuguaglianze economiche misurabili. È la questione del paragrafo successivo.

3.3.3 L'impatto distributivo sui salari italiani: la polarizzazione del mercato del lavoro

Lo studio di Ferri e collaboratori risponde alla domanda su chi sia esposto all'AI nel mercato del lavoro italiano. Resta da affrontare una domanda complementare e altrettanto rilevante per il discorso sulle disuguaglianze: chi guadagna economicamente da questa esposizione, e con quale distribuzione? A questa domanda risponde il lavoro di Irene Brunetti (INAPP) e Chiara Mussida (Università Cattolica del Sacro Cuore di Piacenza), pubblicato nel 2025 come INAPP Working Paper n. 135 con il titolo *AI occupational exposure and wage distribution: The case of Italy*. Il paper si colloca nella collana INAPP successiva a quella di Ferri e ne costituisce uno sviluppo logico, applicando lo strumento concettuale dell'esposizione all'AI a una domanda di natura propriamente distributiva.

L'obiettivo dichiarato dello studio è: «esaminare la relazione tra l'esposizione all'intelligenza artificiale (IA) delle occupazioni e i salari guardando all'intera distribuzione salariale» (Brunetti & Mussida, 2025, abstract). Il punto metodologico decisivo è proprio quest'ultimo. La maggior parte degli studi precedenti sul rapporto tra AI e salari aveva osservato gli effetti sui valori medi della retribuzione: chi è più esposto all'AI guadagna in media di più o di meno? Brunetti e Mussida si chiedono invece se l'effetto sia uniforme lungo tutta la distribuzione salariale, o se vari a seconda della posizione del lavoratore nella scala delle retribuzioni. Per rispondere a questa domanda utilizzano un dataset employer-employee

italiano per il periodo 2011-2019, una regressione quantilica non condizionata (UQR, *Unconditional Quantile Regression*) e una decomposizione Oaxaca-Blinder per scomporre il differenziale salariale in parte spiegata dalle caratteristiche osservabili e parte non spiegata. L'indicatore di esposizione utilizzato è costruito in modo da catturare soprattutto la dimensione di complementarità, in linea con la letteratura più recente sul tema.

Il risultato centrale dello studio è formulato in un passaggio chiave dell'abstract: «i risultati evidenziano l'esistenza di un premio salariale associato alle occupazioni altamente esposte/complementari all'AI, soprattutto per quei lavoratori che si trovano nella parte alta della distribuzione, mentre si riduce per quei lavoratori che si trovano nella parte centrale suggerendo che l'intelligenza artificiale potrebbe aggravare la disuguaglianza salariale tra lavoratori a bassa e ad alta retribuzione, contribuendo così al perdurare di un mercato del lavoro polarizzato» (Brunetti & Mussida, 2025, abstract). Il premio salariale per l'esposizione all'AI esiste, ma non è distribuito in modo uniforme lungo la scala delle retribuzioni: chi si trova nei decili più alti della distribuzione guadagna proporzionalmente di più dall'esposizione, mentre chi si trova nei decili centrali ne ricava un beneficio salariale minore. Il risultato netto è che l'AI, nel mercato del lavoro italiano, tende ad ampliare il divario tra chi sta sopra e chi sta in mezzo, contribuendo al fenomeno noto in letteratura economica come polarizzazione del mercato del lavoro.

La connessione di questa evidenza con il framework della tesi è diretta. Si tratta della versione empirica italiana di quello che il rapporto OECD aveva descritto in chiave aggregata con il premio salariale dell'undici per cento per le skill AI, e di una traduzione concreta sul piano salariale del meccanismo che il Capitolo 2 aveva ricostruito come effetto Matteo. Chi parte da una posizione di vantaggio nella distribuzione

salariale viene ulteriormente avvantaggiato dall'esposizione all'AI: la disuguaglianza si riproduce e si amplifica. Il caso italiano confermato da Brunetti e Mussida è particolarmente significativo perché si colloca in un contesto storico di polarizzazione del lavoro già documentata nella letteratura sulla deindustrializzazione e sulla terziarizzazione degli ultimi decenni. L'AI non inaugura la polarizzazione, ma vi si innesta come fattore di ulteriore accelerazione. Va segnalato un limite temporale dello studio: i dati coprono il periodo 2011-2019 e precedono completamente l'esplosione dell'AI generativa pubblica dopo il rilascio di ChatGPT nel novembre 2022. Per cogliere gli effetti specifici dell'AI generativa sul mercato del lavoro serve un'evidenza più recente, ed è il tema dell'ultimo paragrafo.

3.3.4 Lavoro nell'AI generativa: vincitori, perdenti e shift in expertise

I paragrafi precedenti hanno discusso l'impatto dell'AI sul lavoro nei suoi termini consolidati: automazione di compiti cognitivi non routinari, esposizione differenziata delle professioni, distribuzione asimmetrica dei premi salariali. Tutti questi dati condividono un limite temporale: i periodi di osservazione precedono o appena lambiscono il rilascio pubblico di ChatGPT nel novembre 2022. L'AI generativa ha introdotto un regime ulteriore di trasformazione del lavoro, i cui meccanismi specifici sono parzialmente diversi da quelli dell'AI tradizionale. Le prime evidenze empiriche di larga scala sull'impatto dei modelli generativi nei mercati del lavoro sono emerse soltanto negli ultimi due anni, e quella su cui ci si soffermerà ora è oggi il riferimento più solido.

Lo studio di Teutloff e collaboratori (2025), pubblicato sul *Journal of Economic Behavior & Organization* con il titolo *Winners and losers of*

generative AI: Early evidence of shifts in freelancer demand, è il risultato di una collaborazione multi-istituzionale che riunisce ricercatori del Copenhagen Center for Social Data Science, dell'Oxford Internet Institute, dell'University College London, di OpenAI, dell'ETLA di Helsinki e del Complexity Science Hub di Vienna. L'oggetto dell'analisi è esplicitamente la trasformazione della domanda di lavoro freelance dopo il rilascio pubblico di ChatGPT: «we examine how ChatGPT has changed the demand for freelancers in jobs where generative AI tools can act as substitutes or complements to human labor». L'ambizione metodologica è significativa: lo studio analizza oltre tre milioni di job postings da una delle principali piattaforme globali di freelancing, raccolti tra gennaio 2021 e settembre 2023, ovvero circa un anno prima e un anno dopo il rilascio di ChatGPT. La metodologia computazionale impiega *BERTopic* per partizionare i job postings in 116 cluster di skill, e GPT-4o per classificare ciascun cluster come sostituibile, complementare o non interessato dall'AI generativa.

Il risultato centrale dello studio è formulato nel passaggio chiave dell'abstract: «labor demand increased after the launch of ChatGPT, but only in skill clusters that were complementary to or unaffected by the AI tool. In contrast, demand for substitutable skills, such as writing and translation, decreased by 20–50% relative to the counterfactual trend» (Teutloff et al., 2025). La domanda di freelance per skill sostituibili, come la scrittura di contenuti e la traduzione, è diminuita tra il venti e il cinquanta per cento rispetto al trend atteso in assenza di ChatGPT, con la riduzione più acuta sui lavori brevi della durata di una-tre settimane. La domanda per skill complementari è invece cresciuta: in particolare, secondo i dati riportati nel comunicato ufficiale dell'Oxford Internet Institute, la domanda di sviluppo chatbot e di *natural language processing* è quasi triplicata, e quella di programmazione machine learning è cresciuta

del ventiquattro per cento. L'aggregato netto della domanda di lavoro freelance è aumentato, non diminuito: la conferma empirica del fatto che l'AI generativa non sia un job-killer aggregato è coerente con quanto già documentato dall'OECD per l'AI tradizionale. Quello che cambia non è il volume del lavoro, ma la sua composizione.

Un risultato del paper merita di essere segnalato per la sua rilevanza teorica. La distribuzione dell'esperienza richiesta dai committenti si è modificata in direzioni opposte nei due tipi di skill. Per le skill sostituibili, la riduzione di domanda colpisce in misura maggiore i lavoratori esperti, perché il premio salariale legato all'esperienza non si giustifica più se l'AI generativa può svolgere lo stesso lavoro a costo inferiore. Per le skill complementari, al contrario, la riduzione colpisce i novizi, perché le aziende preferiscono professionisti capaci di sfruttare l'AI come strumento di amplificazione del proprio lavoro. La sintesi di questo doppio movimento è offerta dall'autrice corrispondente dello studio, R. Maria del Rio-Chanona, nel comunicato ufficiale dell'Oxford Internet Institute: «for substitutable skills, demand shifts away from expert workers, while for complementary skills, demand may shift toward greater expertise». Si delinea così quello che gli autori chiamano uno *shift in expertise*: l'AI generativa non riduce in aggregato la domanda di lavoro, ma ridistribuisce la domanda di esperienza in direzioni opposte a seconda del tipo di skill posseduta. Il punto è teoricamente decisivo perché lega in modo diretto l'AI generativa al divario di rendimento ricostruito nel Capitolo 2: a parità di esposizione formale, il rendimento dell'esperienza dipende ora dal tipo di skill che la accompagna. Una stessa quantità di anni di lavoro vale di più o di meno a seconda della posizione rispetto all'AI generativa.

Il quadro che emerge dai quattro sotto-paragrafi del 3.3 può essere ricondotto a due osservazioni conclusive. La prima è che l'AI, nelle sue forme sia tradizionali sia generative, non è un job-killer aggregato: tutti i

dati esaminati convergono nel mostrare che la domanda totale di lavoro non si riduce, e in molti casi cresce. La seconda osservazione è che la stabilità aggregata nasconde una redistribuzione interna delle opportunità lungo nuove linee di stratificazione. Le nuove linee di stratificazione sono multiple. Sul livello di istruzione, l'evidenza italiana di Ferri e collaboratori mostra in modo paradossale che i laureati sono più esposti all'AI rispetto a chi possiede titoli inferiori. Sulla posizione nella distribuzione salariale, lo studio di Brunetti e Mussida documenta che il guadagno proporzionale dall'esposizione cresce con la posizione di partenza. Quanto al tipo di skill, dal lavoro di Teutloff e colleghi emerge una divaricazione tra chi possiede skill complementari all'AI generativa e chi possiede skill sostituibili: i primi vedono crescere la propria domanda di lavoro, i secondi la vedono contrarsi. Sul livello di esperienza, infine, la stessa ricerca rivela un ribaltamento del rendimento dell'esperienza a seconda della skill posseduta. Tutti questi shift ricalcano una stessa logica strutturale: l'effetto Matteo applicato al lavoro, dove la posizione iniziale viene amplificata dalla esposizione all'AI, e dove chi parte in vantaggio amplifica il vantaggio mentre chi parte nello svantaggio amplifica lo svantaggio. La terza dimensione dell'AI divide, il potere, declina la stessa logica al livello della governance dei sistemi algoritmici, ed è il tema del paragrafo successivo.

3.4 Potere: chi decide, chi è deciso, chi è regolato dagli algoritmi

I due paragrafi precedenti hanno mostrato come l'AI ridisegni l'asimmetria della visibilità informativa e delle opportunità lavorative. Questo paragrafo si occupa della terza dimensione dell'AI divide, quella sotto certi aspetti più radicale: il potere. La premessa concettuale del discorso è che i sistemi algoritmici non si limitano a osservare e selezionare contenuti o a mediare l'incontro tra domanda e offerta di

lavoro, ma producono attivamente decisioni che incidono sui diritti fondamentali delle persone. Stabiliscono chi viene assunto in un processo di selezione e chi non passa al colloquio successivo, a chi viene concesso il credito e con quali condizioni, quali soggetti sono ritenuti a rischio nelle valutazioni di sicurezza pubblica, in che modo vengono assegnati o negati specifici servizi sociali. La domanda che il paragrafo affronta diventa allora più articolata: chi ha il potere di decidere come queste decisioni sono prese? Chi ne è responsabile? E con quali strumenti l'ordinamento può intervenire per regolare il fenomeno?

Il discorso si articola in quattro piani complementari. Il primo introduce la formulazione fondativa del concetto di *surveillance capitalism* elaborata da Shoshana Zuboff nel 2015 (paragrafo 3.4.1), che fornisce il vocabolario teorico con cui ancora oggi si discute il potere algoritmico delle Big Tech. Il secondo riprende il framework analitico sulla discriminazione algoritmica e sui diritti umani sviluppato da Evgeni Aizenberg e Jeroen van den Hoven (paragrafo 3.4.2). Il terzo cala il discorso nel contesto italiano, ricostruendo la prassi regolatoria di AGCOM attraverso lo studio di Matteo Giannelli del 2021 (paragrafo 3.4.3). Il quarto chiude il paragrafo con un'estensione del framework di Aizenberg e van den Hoven al regime dell'AI generativa (paragrafo 3.4.4). Il filo conduttore connette il discorso ai concetti del Capitolo 2 ricostruiti nel framework della tesi: il divario di terzo livello come differenziale di potere decisionale, l'effetto Matteo applicato alla distribuzione del controllo algoritmico, la distinzione tra uso consapevole e uso subito come asimmetria di potere nella relazione tra utenti e sistemi.

3.4.1 Surveillance capitalism e Big Other: la formulazione fondativa di Zuboff

La fonte teorica di riferimento del dibattito contemporaneo sul potere algoritmico è l'articolo di Shoshana Zuboff *Big other: surveillance capitalism and the prospects of an information civilization*, pubblicato nel 2015 sul *Journal of Information Technology*. L'autrice, professoressa emerita ad Harvard Business School e affiliata al Berkman Center for Internet and Society dell'Università di Harvard, ha aperto con questo paper un campo di studi e ha fornito il vocabolario con cui ancora oggi si discute il potere delle grandi piattaforme tecnologiche. Il libro successivo dell'autrice, *The Age of Surveillance Capitalism* del 2019, ha esteso e raffinato l'analisi del paper del 2015, ma le coordinate teoriche fondamentali sono già qui.

Il paper di Zuboff descrive «an emergent logic of accumulation in the networked sphere, 'surveillance capitalism'» (Zuboff, 2015, p. 75). La lente d'analisi non è teorica ma empirica: l'autrice parte dall'osservazione delle pratiche di Google, e in particolare da due articoli scientifici dell'allora chief economist dell'azienda Hal Varian, che identificano quattro usi tipici delle transazioni mediate dal computer. I quattro usi sono: estrazione e analisi dei dati, nuove forme contrattuali rese possibili dal monitoraggio capillare, personalizzazione dei servizi, ed esperimenti continui sugli utenti. Questi quattro pilastri tecnici costituiscono, secondo Zuboff, l'architettura della *surveillance capitalism*: una nuova logica di accumulazione che non si fonda sulla vendita di beni o servizi tradizionali, ma sull'estrazione sistematica di dati comportamentali, sulla loro mercificazione e sul loro utilizzo per prevedere e modificare il comportamento futuro degli utenti.

Sul piano del potere, l'innovazione concettuale del paper consiste nell'aver dato un nome alla nuova forma di potere che questa logica di

accumulazione produce. Zuboff lo chiama *Big Other*, in esplicita contrapposizione al *Big Brother* orwelliano. La definizione canonica del concetto, ormai entrata nel vocabolario standard del dibattito sociologico e regolatorio, è la seguente: «Big Other is a ubiquitous networked institutional regime that records, modifies, and commodifies everyday experience from toasters to bodies, communication to thought, all with a view to establishing new pathways to monetization and profit» (Zuboff, 2015, p. 81). Il punto teorico della contrapposizione tra le due figure va sottolineato: Il *Big Brother* orwelliano era un potere centralizzato, incarnato dallo Stato di sorveglianza: si manifestava in modo visibile, agiva in chiave repressiva e si reggeva sulla coercizione esterna ai propri sottoposti. Il *Big Other* di Zuboff opera in modo diametralmente opposto. Una pluralità di attori privati lo distribuisce sul territorio digitale, in luogo della centralizzazione statale. La sua presenza si confonde con l'esperienza ordinaria della tecnologia, dove il *Big Brother* era riconoscibile come autorità repressiva. E al posto della coercizione esterna, si regge sulla cooperazione apparentemente volontaria degli utenti, che producono in continuazione i dati attraverso cui essi stessi vengono poi sorvegliati.

La descrizione tecnica del funzionamento del *Big Other* si articola in tre meccanismi che operano contemporaneamente. Sempre dal paper: il sistema è costituito da «mechanisms of extraction, commodification, and control that effectively exile persons from their own behavior while producing new markets of behavioral prediction and modification» (Zuboff, 2015). Il punto sociologicamente più potente è la metafora dell'esilio. Gli individui sono allontanati dal proprio stesso comportamento perché i dati che producono attraverso le proprie azioni quotidiane, dagli acquisti alle ricerche online, dagli spostamenti alle conversazioni, diventano materia prima per nuovi mercati a cui essi stessi non hanno accesso e di cui non controllano il funzionamento. Si crea un nuovo tipo

di asimmetria informativa e di potere: il valore economico viene prodotto dagli individui, ma se ne appropriano le piattaforme; nella relazione informativa, gli utenti sono trasparenti per il sistema, mentre il funzionamento del sistema rimane opaco per gli utenti.

Il messaggio politico del paper di Zuboff è formulato in modo netto: «surveillance capitalism challenges democratic norms and departs in keyways from the centuries-long evolution of market capitalism» (Zuboff, 2015). La *surveillance capitalism* non è semplicemente un nuovo modello di business, è una rottura con il capitalismo di mercato classico, fondato sulla reciprocità tra produttori e consumatori. Rappresenta una sfida alle norme democratiche perché introduce asimmetrie strutturali di informazione e di potere che le istituzioni democratiche tradizionali faticano a regolare. La connessione con il framework della tesi è diretta: la *surveillance capitalism* è il meccanismo storico-economico che converte l'uso quotidiano della tecnologia digitale in rendimento asimmetrico per le piattaforme, in una versione capitalistica e algoritmica dell'effetto Matteo discusso nel Capitolo 2. La questione che resta aperta è come questo nuovo tipo di potere possa essere regolato. Sul piano analitico, una delle articolazioni più rilevanti del problema è quella offerta dal framework che si discute nel paragrafo successivo.

3.4.2 Discriminazione algoritmica e diritti umani: il framework di Aizenberg e van den Hoven

Sei anni dopo l'articolo fondativo di Zuboff, il dibattito sul potere algoritmico ha conosciuto una significativa articolazione tecnica. Una delle formulazioni più rilevanti è quella di Evgeni Aizenberg e Jeroen van den Hoven, ricercatori alla Delft University of Technology. Van den Hoven, in particolare, è figura di riferimento internazionale per il

framework *Design for Values*, oggi adottato anche dalle istituzioni europee come riferimento metodologico per la *Trustworthy AI*. Il loro articolo *Designing for human rights in AI*, pubblicato in open access nel 2020 sulla rivista *Big Data & Society*, fornisce una griglia analitica per affrontare il problema della trasparenza, dell'accountability e della discriminazione algoritmica. È utile riprendere l'articolo nel quadro del pilastro Potere della tesi perché ne risulta una sintesi operativa, non solo critica, dei problemi sollevati da Zuboff cinque anni prima.

Il punto di partenza degli autori è una constatazione empirica sulla pervasività delle decisioni algoritmiche nella vita sociale contemporanea: «in the age of Big Data, companies and governments are increasingly using algorithms to inform hiring decisions, employee management, policing, credit scoring, insurance pricing, and many more aspects of our lives» (Aizenberg & van den Hoven, 2020). L'argomento è importante. Gli algoritmi non operano più in nicchie tecniche periferiche: sono entrati negli ambiti più sensibili della vita sociale. Vengono utilizzati nei processi di assunzione e nella gestione del personale, nelle decisioni di polizia e nella valutazione del merito creditizio, nel calcolo dei prezzi assicurativi. Si tratta di ambiti in cui le decisioni incidono direttamente sui diritti delle persone, e in cui l'introduzione di sistemi automatici produce conseguenze nuove e spesso poco prevedibili.

Il rischio specifico che gli autori individuano nei sistemi AI in questi ambiti è formulato con precisione: «AI systems can help us make evidence-driven, efficient decisions, but can also confront us with unjustified, discriminatory decisions wrongly assumed to be accurate because they are made automatically and quantitatively» (Aizenberg & van den Hoven, 2020). Il punto centrale è il paradosso epistemico delle decisioni algoritmiche, ovvero la legittimazione tecnica della discriminazione. La forma quantitativa e automatica del giudizio induce a

percepirlo come oggettivo e accurato, anche quando in realtà riproduce o amplifica discriminazioni esistenti nei dati di addestramento. Ciò che la società ha imparato a riconoscere come discriminazione quando è prodotta da un decisore umano, e quindi a contestare per via giuridica o politica, viene presentato come neutralità statistica quando è prodotto da un sistema algoritmico, e di conseguenza appare più difficile da contestare. L'argomento si lega in modo diretto al discorso sui *bias* dei dati di training degli LLM sviluppato nel paragrafo 3.2.4 a partire da Bender e collaboratrici, mostrando come il problema sia strutturale ai sistemi AI e non limitato all'AI generativa.

Il nodo tecnico che rende difficile contestare le decisioni algoritmiche è approfondito nella rassegna di Ben Chester Cheong (2024), pubblicata su *Frontiers in Human Dynamics*, che dedica al problema della trasparenza e dell'accountability dei sistemi AI una review della letteratura legale ed etica. Il punto di partenza è il cosiddetto problema della scatola nera: i sistemi di AI fondati su reti neurali profonde sono «'black boxes' whose inner workings are inscrutable to humans» (Cheong, 2024). Nei modelli più complessi l'opacità è una caratteristica strutturale, non un difetto accidentale, e pone una difficoltà diretta per la possibilità di rendere conto delle decisioni prese. A questa difficoltà il dibattito giuridico ha risposto con strumenti come il diritto alla spiegazione introdotto dal Regolamento generale sulla protezione dei dati europeo, che secondo Cheong configura «a '*right to explanation*' requiring disclosure of the logic behind solely automated decisions» (Cheong, 2024). Lo stesso autore però segnala il limite di questo approccio: «transparency alone is insufficient for accountability, as explanations of AI systems can still be highly technical» (Cheong, 2024). La sola trasparenza, in altre parole, lascia aperto il problema della comprensibilità effettiva delle decisioni per chi le subisce, perché la spiegazione tecnica di un modello complesso può

restare inaccessibile al cittadino comune e perfino al regolatore. Per il discorso sul potere che attraversa questo paragrafo, la conclusione di Cheong è particolarmente pertinente: l'accountability effettiva richiede di affrontare gli squilibri di potere tra chi sviluppa i sistemi e chi ne è soggetto, perché «meaningful AI accountability requires grappling with power imbalances between AI developers and those affected by their systems» (Cheong, 2024). Per riequilibrare l'asimmetria di potere descritta da Zuboff, la trasparenza non basta da sola. Va affiancata da meccanismi istituzionali: il controllo esterno, l'audit indipendente, e il coinvolgimento attivo dei soggetti più esposti alle decisioni automatizzate.

L'inquadramento offerto dagli autori va oltre la dimensione etica generica e tocca il piano dei diritti fondamentali: «these technological developments are consequential to people's fundamental human rights» (Aizenberg & van den Hoven, 2020). Il riferimento esplicito ai diritti umani sposta il discorso da un piano etico a un piano costituzionale e internazionale. Ciò che è in gioco non sono valori generici, ma diritti riconosciuti a livello costituzionale nazionale e nelle convenzioni internazionali. Su questa base, gli autori muovono una critica articolata a due derive complementari del dibattito contemporaneo. La prima è il *solutionism trap* (Morozov, 2013), ossia l'idea che ogni problema sociale abbia una soluzione tecnica, e che l'AI sia di volta in volta la soluzione disponibile. La seconda è il limite degli approcci etici generici: «calls for more ethically- and socially-aware AI often fail to provide answers for how to proceed beyond stressing the importance of transparency, explainability, and fairness» (Aizenberg & van den Hoven, 2020). Il rischio è che termini come trasparenza, spiegabilità ed equità diventino parole d'ordine retoriche senza un'effettiva traduzione operativa nei sistemi.

La proposta degli autori per superare entrambe le derive si chiama *Design for Values*: un framework metodologico che integra il *Value Sensitive Design* di Friedman con il *Participatory Design* della tradizione scandinava per tradurre i diritti umani in requisiti progettuali concreti, attraverso il coinvolgimento attivo degli stakeholder sociali fin dalle prime fasi della progettazione di un sistema. Sul piano teorico, il concetto più potente del paper è quello di *statistical dehumanization*, la disumanizzazione statistica intesa come riduzione di un individuo a una rappresentazione numerica, con perdita della specificità contestuale e dell'autonomia di autodefinirsi. La connessione con il discorso di Zuboff è ravvicinata: la *statistical dehumanization* di Aizenberg e van den Hoven è il meccanismo che converte la persona umana in un punto in uno spazio statistico, in una nuova forma di alienazione algoritmica che richiama da vicino l'esilio dal proprio comportamento descritto da Zuboff. Resta aperta la domanda su come gli ordinamenti nazionali e sovranazionali possano operativamente regolare questi sistemi.

3.4.3 La risposta regolatoria italiana: la prassi AGCOM ricostruita da Giannelli

Il discorso teorico di Zuboff e quello analitico di Aizenberg e van den Hoven inquadrano il problema del potere algoritmico in termini generali. Resta da capire come gli ordinamenti nazionali abbiano risposto al problema sul piano regolatorio concreto. Per il caso italiano, la fonte di riferimento è l'articolo di Matteo Giannelli *Poteri dell'AGCOM e uso degli algoritmi. Tra innovazioni tecnologiche ed evoluzioni del quadro regolamentare*, pubblicato nel 2021 sulla rivista *Osservatorio sulle fonti*, classificata in classe A ANVUR per le scienze giuridiche, in un numero monografico interamente dedicato al rapporto tra autorità amministrative indipendenti e regolazione delle decisioni algoritmiche. L'obiettivo

dichiarato del saggio è di «analyse and describe the practice of decisions, studies and any other documentation that AGCOM has published on the subject of algorithmic decisions in recent years» (Giannelli, 2021). Si tratta quindi di una ricostruzione sistematica e ragionata della prassi regolatoria, non di una riflessione meramente teorica.

AGCOM (Autorità per le Garanzie nelle Comunicazioni) è l'autorità amministrativa indipendente italiana competente sui media e sulle comunicazioni, istituita nel 1997, ed è la stessa autorità il cui rapporto del 2020 sulle dispercezioni cognitive italiane è stato discusso nel paragrafo 3.2.3 sulla visibilità. La sua competenza si è progressivamente estesa dalle materie tradizionali, ovvero radio, televisione e stampa, alle piattaforme digitali e ai servizi di intermediazione algoritmica, in linea con l'evoluzione del quadro regolatorio europeo. Tra le iniziative ricostruite da Giannelli, particolare rilievo ha l'indagine conoscitiva congiunta condotta da AGCOM con l'AGCM (Autorità Garante della Concorrenza e del Mercato) e con il Garante per la protezione dei dati personali sui *Big Data*, avviata nel 2017 e conclusa nel 2020 con la pubblicazione di linee guida e raccomandazioni di policy. L'indagine ha identificato problemi convergenti di privacy, di concorrenza e di libertà di informazione legati all'uso di grandi masse di dati da parte delle piattaforme.

Un nodo centrale del paper riguarda il problema della trasparenza nella selezione dei contenuti da parte delle piattaforme e i limiti delle forme di auto-regolazione adottate da queste ultime. AGCOM ha tentato in più occasioni di richiamare le piattaforme al rispetto del codice di autoregolamentazione su una pluralità di temi. Tra i casi documentati da Giannelli emergono in particolare due interventi: l'applicazione del divieto di pubblicità di giochi e scommesse introdotto dal cosiddetto decreto Dignità del 2018, con i suoi effetti sui servizi di indicizzazione e di promozione algoritmica dei siti web, e l'atto di richiamo del 2020 alle

principali piattaforme di condivisione contenente l'invito al rispetto dei principi a tutela della correttezza dell'informazione, nel contesto della pandemia di Covid-19. I casi ricostruiti mostrano un quadro ricorrente nella prassi regolatoria italiana del periodo. L'auto-regolazione delle piattaforme appare insufficiente perché il loro modello di business non è coerente con gli obiettivi di pluralismo informativo, di correttezza dell'informazione e di tutela degli utenti che la regolazione pubblica intende perseguire. La conseguenza, ricavata da Giannelli, è la necessità di forme più strutturate di co-regolazione o di regolazione vera e propria, in cui l'autorità pubblica definisca vincoli effettivi e ne verifichi il rispetto.

Dal 2021, anno di pubblicazione del paper di Giannelli, il quadro regolatorio italiano si è ulteriormente trasformato per via dell'esplosione dell'AI generativa nel grande pubblico, conseguente al rilascio di ChatGPT nel novembre 2022. Già nel marzo 2023, il Garante per la protezione dei dati personali italiano ha temporaneamente sospeso ChatGPT nel territorio nazionale per presunte violazioni del Regolamento generale sulla protezione dei dati. Come riportato dal rapporto OECD del 2023, «during April 2023, Italy's data protection agency temporarily banned ChatGPT from processing personal data of Italian data subjects due to several alleged violations of the GDPR» (OECD, 2023, p. 200). Si è trattato di uno dei primi grandi interventi regolatori al mondo sull'AI generativa di *general purpose*, e ha aperto un precedente discusso in numerosi report internazionali. Il caso italiano del Garante mostra come la prassi regolatoria descritta da Giannelli per il periodo precedente sia continuata, e si sia rapidamente estesa al nuovo regime tecnologico. Resta da capire come i problemi di potere, di trasparenza e di accountability discussi nei paragrafi precedenti si declinino specificamente nel regime generativo, ed è il tema dell'ultimo paragrafo del pilastro Potere.

3.4.4 Potere nell'AI generativa: dal framework di Aizenberg al regime post-ChatGPT

I paragrafi precedenti hanno discusso il potere algoritmico nei suoi termini consolidati: la *surveillance capitalism* come logica di accumulazione dei dati comportamentali, la discriminazione algoritmica come problema di diritti umani fondamentali, la risposta regolatoria italiana di AGCOM e del Garante per la protezione dei dati personali. Tutti questi piani condividono un riferimento implicito a sistemi di AI che selezionano, classificano e decidono. L'AI generativa, con la diffusione pubblica dei modelli linguistici di grandi dimensioni a partire dal 2022, introduce un regime ulteriore: i sistemi non solo decidono, ma producono. La domanda sul potere si articola di conseguenza in modo nuovo. Non si tratta più soltanto di chi decide come gli algoritmi classificano le persone, ma di chi decide quali testi, quali immagini, quali rappresentazioni del mondo entrano nello spazio pubblico mediato dall'AI.

Il framework di Aizenberg e van den Hoven, discusso nel paragrafo 3.4.2, si presta in modo naturale a essere esteso al regime generativo, perché i problemi che identifica sono strutturali al rapporto tra sistemi AI e diritti umani fondamentali. La discriminazione algoritmica mascherata da accuratezza quantitativa, criticata dagli autori per i sistemi decisionali tradizionali, si ripresenta negli LLM nella forma del *bias* dei dati di training discusso nel paragrafo 3.2.4 a partire da Bender e collaboratrici. I modelli generativi sovrarappresentano le visioni egemoniche e marginalizzano le voci minoritarie, e questa distorsione viene amplificata dalla forma fluente e autorevole della loro produzione testuale, che induce nel lettore una percezione di credibilità sproporzionata rispetto al rigore epistemico effettivo del contenuto. Il problema della scatola nera segnalato da Cheong (2024) si acuisce nel regime generativo, perché l'opacità non riguarda più soltanto il criterio di una singola decisione, ma l'intero

processo di produzione di un testo, di cui l'utente non può ricostruire né le fonti né i passaggi inferenziali. Il concetto di *statistical dehumanization* trova nell'AI generativa una nuova declinazione. La persona umana non solo viene ridotta a un punto nello spazio statistico, come avveniva nei sistemi decisionali tradizionali, ma può diventare oggetto di descrizioni testuali, narrazioni e ritratti prodotti automaticamente, con tutto il peso di apparente autorevolezza che la forma scritta porta con sé.

Una parte della riflessione di Aizenberg e van den Hoven che si applica con particolare acutezza al regime generativo è la critica al *solutionism trap*. Nel paper del 2020 gli autori scrivono che «we should not fall into the solutionism trap, assuming that some form of AI is always the solution to the problem or goal at hand» (Aizenberg & van den Hoven, 2020). Il monito ha una rilevanza specifica nell'epoca attuale, caratterizzata dalla corsa all'integrazione degli LLM in ogni servizio pubblico e privato, dagli *assistant* per l'assistenza clienti agli strumenti per la pubblica amministrazione, dai dispositivi medici ai sistemi educativi. Il rischio di assumere che l'AI generativa sia sempre la soluzione, anche quando il problema da risolvere non si presta a una formulazione algoritmica, è oggi particolarmente acuto. Il framework *Design for Values* fornisce uno strumento metodologico per identificare i casi in cui l'AI generativa non dovrebbe essere usata, e ciò richiede il coinvolgimento esplicito degli stakeholder sociali nel processo di progettazione, secondo le metodologie partecipative discusse nel paragrafo precedente.

Riprendendo i tre meccanismi di *Big Other* identificati da Zuboff, ossia estrazione, mercificazione e controllo, è possibile osservare come ciascuno di essi assuma nel regime generativo una forma intensificata. L'estrazione di dati non si limita più alla cattura del comportamento di navigazione dell'utente, ma include i contenuti testuali e visivi che gli utenti producono nel dialogare con i modelli, contenuti che diventano essi

stessi materia prima per ulteriori cicli di addestramento. La mercificazione non riguarda più soltanto le previsioni di comportamento, ma anche le capacità generative del modello, vendute come servizi a imprese e istituzioni che le integrano nei propri prodotti. Il controllo, infine, opera su un piano nuovo: l'AI generativa non condiziona soltanto le decisioni che vengono prese sulle persone, ma plasma il modo stesso in cui le persone formulano le proprie idee, scrivono i propri testi e accedono al sapere disponibile. Si delinea così una stratificazione di asimmetrie tra chi controlla l'infrastruttura computazionale dei grandi modelli, accessibile soltanto a un numero ridotto di attori economici globali, e chi ne usa gli output, ovvero la stragrande maggioranza degli utenti finali.

Il framework di Aizenberg e van den Hoven, integrato con il vocabolario teorico di Zuboff, fornisce una griglia analitica robusta per leggere il regime generativo. La discriminazione mascherata da accuratezza quantitativa, la *statistical dehumanization*, il rischio del *solutionism trap* e l'inquadramento nel piano dei diritti umani fondamentali sono tutti elementi che mantengono la loro pertinenza nel passaggio dall'AI tradizionale all'AI generativa, e che richiedono semmai un grado di articolazione maggiore alla luce della specifica capacità produttiva dei modelli. Il discorso non si esaurisce qui, e una sintesi più ampia delle tre dimensioni dell'AI divide ricostruite nel capitolo è demandata alle considerazioni conclusive del paragrafo finale.

3.5 Tre dimensioni, un solo meccanismo

Il percorso condotto nelle pagine precedenti ha articolato l'AI divide su tre dimensioni interconnesse, secondo la tesi enunciata in apertura del capitolo. La visibilità ha a che fare con la presenza nei flussi informativi mediati dagli algoritmi di raccomandazione. Sul fronte del

lavoro è la posizione di ciascuno nei mercati occupazionali trasformati dall'intelligenza artificiale a essere in gioco. Sul potere, infine, conta il controllo sui sistemi algoritmici che decidono, selezionano e oggi producono contenuti. Lungo i tre pilastri è emerso con regolarità un esito comune: l'AI divide non distribuisce i propri effetti in modo neutro. Ricalca e amplifica le linee di stratificazione sociale e territoriale già esistenti, oltre a quelle costruite sulle competenze, secondo un meccanismo che la tesi ha ricondotto in più punti all'effetto Matteo ricostruito nel Capitolo 2. La transizione algoritmica riproduce e intensifica le asimmetrie dell'ordine sociale precedente, ne modifica i contenuti senza alterarne la logica di fondo.

Nel paragrafo 3.2, dedicato alla visibilità, è emerso come la selezione algoritmica dei contenuti non sia un processo tecnico neutrale, ma un meccanismo che distribuisce in modo diseguale la possibilità di accedere a un universo informativo plurale. La rilettura della *filter bubble* e dell'*echo chamber* alla luce delle *systematic review* più recenti ha permesso di superare le formulazioni più allarmistiche e di precisare i meccanismi effettivi della chiusura informativa, mentre l'indagine condotta da AGCOM ha calato il problema nel contesto italiano, mostrando come la dispercezione della realtà sociale segua le linee della stratificazione cognitiva. Il punto teorico decisivo emerso dal paragrafo è che la visibilità non è una proprietà del contenuto, ma un esito della distribuzione diseguale del capitale algoritmico tra gli utenti. Chi dispone di alfabetizzazione avanzata può attivare strategie di compensazione e riconoscere il filtro; chi non ne dispone subisce la selezione come ambiente di default. Nel regime dell'AI generativa, come mostrato a partire dal lavoro di Bender e collaboratrici, la distorsione si sposta dai contenuti selezionati ai contenuti prodotti, con un'amplificazione delle visioni egemoniche e una marginalizzazione delle voci minoritarie.

Sul lavoro (paragrafo 3.3) il quadro che emerge è in apparenza rassicurante e in realtà problematico. L'evidenza disponibile, dal rapporto OECD agli studi italiani di Ferri e collaboratori e di Brunetti e Mussida, converge nel mostrare che l'AI non è un fattore di distruzione aggregata dell'occupazione. La domanda totale di lavoro non si riduce, e in molti casi cresce. Questa stabilità aggregata, però, nasconde una redistribuzione interna delle opportunità lungo nuove linee di stratificazione. L'esposizione all'AI cresce con il livello di istruzione, in un'inversione paradossale del digital divide classico; il premio salariale dell'esposizione si concentra nella parte alta della distribuzione, aggravando la polarizzazione del mercato del lavoro italiano; lo *shift in expertise* documentato da Teutloff e collaboratori nel regime dell'AI generativa ridistribuisce la domanda di esperienza in direzioni opposte a seconda del tipo di competenza posseduta. Anche qui la logica sottostante è quella dell'effetto Matteo applicato al lavoro: la posizione di partenza viene amplificata dalla transizione algoritmica, a vantaggio di chi parte avvantaggiato e a svantaggio di chi parte in posizione intermedia o debole.

Quanto al potere, il paragrafo 3.4 ha ricostruito la genealogia di una stratificazione di forme di potere algoritmico che si sono progressivamente accumulate. Il potere di estrazione e mercificazione dei dati comportamentali, formulato da Zuboff attraverso i concetti di *surveillance capitalism* e di *Big Other*; il potere decisionale automatico nei domini sensibili dei diritti, analizzato da Aizenberg e van den Hoven attraverso le nozioni di discriminazione algoritmica e di *statistical dehumanization*; il problema dell'opacità dei sistemi, ricostruito da Cheong attraverso la nozione di scatola nera e il limite del solo principio di trasparenza; il potere di selezione e regolazione della sfera pubblica informativa, ricostruito nel caso italiano attraverso la prassi di AGCOM studiata da Giannelli e l'intervento del Garante per la protezione dei dati personali sul caso

ChatGPT. Il filo che tiene insieme questa genealogia è che il potere algoritmico non si limita a registrare la realtà sociale, ma la modifica. Le decisioni prese o assistite dai sistemi automatici incidono direttamente sui diritti fondamentali delle persone, e nel regime dell'AI generativa questo potere assume una forma ulteriore, perché i sistemi non solo decidono sulle persone, ma producono i testi, le immagini e le rappresentazioni del mondo entro cui le persone formano le proprie idee.

Le tre dimensioni non vanno lette come compartimenti separati, ma come parti di un unico meccanismo che si rafforza in modo circolare. Tra di esse opera una relazione di rinforzo cumulativo che richiama da vicino la quinta proposizione della *Resources and Appropriation Theory* di van Dijk, dove le disuguaglianze digitali retroagiscono sulle disuguaglianze sociali di partenza. Chi è marginalizzato nella visibilità informativa dispone di meno strumenti per costruire le competenze che gli consentirebbero di collocarsi favorevolmente nei mercati del lavoro trasformati dall'AI; chi occupa una posizione debole nel lavoro algoritmico ha meno potere di comprendere e contestare i sistemi che decidono sulla sua vita; chi ha meno potere decisionale è a sua volta più esposto alla marginalizzazione nei flussi informativi, in un circolo che tende a chiudersi su se stesso. Il regime dell'AI generativa, che attraversa come filo conduttore tutti e tre i pilastri del capitolo, non si limita a intensificare ciascuna dimensione presa singolarmente, ma rafforza i legami che le uniscono. La stessa infrastruttura computazionale che produce i contenuti che plasmano la visibilità è quella che ridefinisce le competenze richieste dal mercato del lavoro ed è quella che concentra il potere generativo in un numero ridotto di attori economici globali. Le tre dimensioni dell'AI divide non sono tre problemi distinti, ma tre facce di un'unica trasformazione, e chiedono di essere affrontate come tali.

Il quadro ricostruito in questo capitolo conferma e specifica la tesi di fondo del lavoro. L'AI divide è l'esito coerente della genealogia ricostruita nel Capitolo 2, che aveva mostrato come ogni ondata tecnologica produca un riposizionamento del divario digitale senza il suo superamento, e si innesta sulle dinamiche di riproduzione delle disuguaglianze sociali analizzate nel Capitolo 1. Le categorie elaborate nei capitoli precedenti, dall'effetto Matteo mertoniano al divario di terzo livello di van Deursen e Helsper, dalla distinzione tra uso consapevole e uso subito alla catena causale circolare tra risorse e disuguaglianze, hanno trovato nelle tre dimensioni della visibilità, del lavoro e del potere una declinazione concreta e contemporanea. Il caso italiano, in particolare, si è rivelato un laboratorio di osservazione privilegiato, perché vi si manifesta in forma simultanea quella stratificazione dei divari che in altri contesti si è dispiegata in fasi distinte: segmenti consistenti della popolazione non hanno ancora consolidato le competenze di base, mentre il panorama tecnologico introduce già le sfide poste dall'AI divide nelle sue tre dimensioni. Si tratta, in altre parole, di una stratificazione simultanea dei divari, dove ogni ondata si sovrappone alle precedenti invece di superarle. Resta da raccogliere le fila dell'intero percorso e da trarne le implicazioni complessive, sul piano teorico e su quello delle possibili risposte: è il compito delle considerazioni conclusive che chiudono il presente lavoro.

Capitolo 4. Dalla diagnosi alla proposta: l'AI divide nel contesto italiano ed europeo

4.1 Introduzione e metodo

I capitoli precedenti hanno ricostruito l'AI divide sul piano teorico. È emerso un quadro in cui l'intelligenza artificiale, lungi dall'introdurre una frattura del tutto inedita, si innesta sulle diseguaglianze esistenti lungo i tre assi della visibilità, del lavoro e del potere. Questo capitolo verifica quel quadro sul terreno del contesto italiano ed europeo, mettendo a confronto due ordini di evidenze: i dati statistici più recenti sulla diffusione dell'intelligenza artificiale e sulla percezione del rischio a essa associato, e le testimonianze di chi, per ruolo professionale, osserva il fenomeno da vicino. L'intento di queste pagine non è quello di misurare l'AI divide in termini quantitativi esaustivi. Si tratta piuttosto di comprenderne le forme concrete e le interpretazioni, affiancando ai dati lo sguardo di chi opera sul campo.

Il disegno della ricerca è di tipo prevalentemente qualitativo, con un sostegno quantitativo. L'asse centrale è costituito da un ciclo di interviste a esperti, scelti in modo da rappresentare prospettive diverse e complementari sul fenomeno. A queste si affiancano, come base oggettiva di riferimento, i dati della rilevazione ISTAT sull'adozione delle tecnologie digitali nelle imprese italiane e quelli dell'indagine europea *European Social Survey* sulla percezione del rischio occupazionale legato all'intelligenza artificiale. La combinazione dei due piani risponde a una logica precisa: i dati documentano l'esistenza e l'ampiezza del divario, mentre le voci degli esperti ne restituiscono le ragioni e i meccanismi, permettendo di interpretare ciò che le statistiche si limitano a registrare.

Le interviste sono state condotte in forma semi-strutturata, sulla base di una traccia comune di quattro domande aperte, uguali per tutti gli interlocutori e adattate solo nel riferimento iniziale al profilo professionale di ciascuno. La prima domanda invitava l'intervistato a descrivere se e in quali forme l'AI divide si manifesti concretamente nel proprio ambito di osservazione. Si passava poi al nodo di chi rischi oggi di restare indietro e chi invece tragga vantaggio dalla diffusione dell'intelligenza artificiale. La terza domanda spostava l'attenzione sulle risposte possibili, interrogando gli intervistati sull'adeguatezza della regolazione europea (e in particolare dell'AI Act), e chiedendo quali strumenti oggi mancherebbero. La traccia si chiudeva infine con un invito a indicare la preoccupazione principale per gli anni a venire e una priorità per chi deve decidere le politiche pubbliche. Una scelta metodologica ha guidato la costruzione della traccia: le tre dimensioni della ricerca, visibilità, lavoro e potere, non sono mai state nominate nelle domande, in modo da lasciarle emergere spontaneamente dalle risposte ed evitare di orientare gli intervistati verso le categorie della tesi.

Gli esperti coinvolti provengono da ambiti deliberatamente eterogenei, così da illuminare il fenomeno da angolazioni diverse. Iolanda Castellana, marketing project manager della Scuola di Scienze Aziendali e Tecnologie Industriali Piero Baldesi, porta il punto di vista della formazione manageriale e del rapporto tra competenze e impresa. C'è poi Alessandro Ferrari, fondatore e amministratore delegato di ARGO Vision e docente di intelligenza artificiale all'Università di Pavia, che unisce la prospettiva tecnica dello sviluppatore a quella dell'imprenditore in un mercato regolamentato. Da Paolo Perulli, sociologo dell'economia, viene l'inquadramento teorico e di lungo periodo del fenomeno. Antonio Romeo, vice president presales con esperienze in Microsoft e Dell, osserva la trasformazione dall'interno del settore tecnologico e della consulenza alle

imprese. Lo sguardo di chi opera quotidianamente nei mercati digitali arriva infine da Massimo Ciotta, CEO e Advertising specialist di Piano MAKE. La diversità dei loro osservatori è una risorsa più che un limite, perché consente di verificare la convergenza (o la divergenza) di letture provenienti da mondi distanti sui medesimi nodi.

Accanto alle interviste, il disegno della ricerca ha previsto uno strumento quantitativo di natura esplorativa: un questionario anonimo, costruito su dodici domande a risposta chiusa. Lo strumento è stato diffuso attraverso la newsletter di un incubatore di startup, un canale che ha raggiunto un pubblico eterogeneo per profilo professionale, settore ed età, per quanto non rappresentativo della popolazione generale. Le domande, formulate in coerenza con i temi affrontati nei colloqui, indagavano l'uso reale dell'intelligenza artificiale nel proprio lavoro, la percezione del cambiamento in atto e atteso, la lettura del divario e di chi rischia di restare indietro, il grado di preparazione personale e la consapevolezza del quadro normativo. La rilevazione, condotta su base volontaria e nel rispetto delle norme sul trattamento dei dati, ha raccolto quaranta risposte. Si tratta di un numero troppo limitato per fondare un'analisi statisticamente rappresentativa, e per questa ragione il questionario non viene qui trattato come una fonte di evidenza al pari dei dati ISTAT ed ESS, ma come un pretest: una prima ricognizione che serve a saggiare lo strumento e a far emergere alcune tendenze, da verificare in una ricerca futura di più ampio respiro. I suoi risultati sono ripresi, in questa funzione, nel paragrafo conclusivo del capitolo.

Come ciascuno di questi strumenti, anche l'impianto della ricerca nel suo insieme presenta dei limiti, che è opportuno dichiarare fin da subito. Le interviste raccolgono il punto di vista di un numero ristretto di testimoni privilegiati, e non costituiscono un campione statisticamente rappresentativo. Il loro valore risiede piuttosto nella competenza e nella

posizione di osservazione degli intervistati, più che nella generalizzabilità numerica delle loro affermazioni. Allo stesso modo, i dati statistici utilizzati provengono da rilevazioni concepite per altri scopi, e vengono qui impiegati come sfondo di contesto, con le cautele che il paragrafo seguente espliciterà. Sono limiti che definiscono la natura dell'analisi senza indebolirla: si tratta di cogliere in profondità le dinamiche dell'AI divide così come si presentano a chi le vive, lasciando ad altri studi la misura definitiva. È in questa cornice che vanno letti i risultati esposti nei paragrafi che seguono, a partire dall'evidenza dei dati.

4.2 Il divario esiste e cresce: l'evidenza dei dati

Il terzo capitolo ha ricostruito sul piano teorico le tre dimensioni lungo cui l'intelligenza artificiale ridisegna le diseguaglianze, distinguendo i terreni della visibilità, del lavoro e del potere. Quel quadro trova riscontro nella realtà italiana ed europea? I dati più recenti permettono di rispondere in modo abbastanza preciso. La rilevazione ISTAT sull'uso delle tecnologie dell'informazione e della comunicazione nelle imprese, che ogni anno fotografa la diffusione del digitale nel tessuto produttivo nazionale, documenta che il divario nell'adozione dell'intelligenza artificiale esiste, e che la sua tendenza è quella ad allargarsi (ISTAT, 2025).

Nel 2025 la quota di imprese italiane con almeno dieci addetti che utilizza almeno una tecnologia di intelligenza artificiale raggiunge il 16,4%, raddoppiando l'8,2% dell'anno precedente. Due anni prima, nel 2023, ferma al 5,0%. La crescita è quindi rapida, anche se la diffusione resta minoritaria: oltre quattro imprese su cinque non hanno ancora adottato alcuna tecnologia di questo tipo. Il dato aggregato nasconde però una frattura interna sostanziale. Le piccole e medie imprese si fermano al

15,7%; le grandi imprese arrivano invece al 53,1%, e per la prima volta superano la soglia di una su due. Quello che è davvero nuovo è la dinamica nel tempo della distanza tra i due gruppi: la forbice nell'intensità di utilizzo tra grandi e PMI valeva circa venti punti percentuali nel 2023, è cresciuta fino a venticinque nel 2024, e ha raggiunto i trentasette nel 2025.

Questo andamento assume un significato particolare se collocato nel quadro più ampio della digitalizzazione. Per la maggior parte degli indicatori digitali i divari dimensionali tra le imprese tendono oggi a ridursi: l'adozione dell'intelligenza artificiale fa eccezione, e la forbice tra grandi e piccole imprese continua ad allargarsi. È la traduzione empirica, sul piano delle imprese, di quel meccanismo di vantaggio cumulativo che il capitolo precedente ha riconosciuto tra i lavoratori. È l'effetto Matteo: chi parte avanti accelera ulteriormente, mentre chi parte indietro fatica a recuperare il terreno, e la tecnologia più avanzata invece di chiudere le distanze preesistenti ne apre di nuove e più profonde. Iolanda Castellana, nell'intervista, traduce questo dato in un'espressione efficace. Avverte il rischio di una «polarizzazione del tessuto produttivo» tra «poche grandi aziende e startup native digitali capaci di correre a velocità esponenziale» e «la stragrande maggioranza delle nostre piccole e medie imprese tradizionali che rischiano l'estraniamento competitivo per mancanza di competenze interne».

Le radici di questo divario sono antiche. Per adottare l'intelligenza artificiale serve una dotazione digitale di base che le imprese italiane possiedono in misura diseguale. L'analisi dei dati, per esempio, è praticata dal 41,9% delle piccole e medie imprese: nelle grandi imprese la quota sale all'83,6%. Distanze analoghe si riscontrano nell'uso dei software gestionali. Anche la dimensione territoriale conferma la concentrazione del fenomeno: le imprese del Nord-ovest registrano la crescita più marcata nell'adozione, che è passata dall'8,9% del 2024 al 19,3% del 2025. Chi non

dispone delle competenze e delle infrastrutture digitali necessarie non può accedere all'intelligenza artificiale, e in questo modo il nuovo divario si appoggia su quelli già esistenti, rafforzandoli.

Sul versante oggettivo i dati ISTAT misurano l'adozione della tecnologia. Sul versante percettivo i dati di un'indagine europea restituiscono un quadro complementare. La terza ondata dell'indagine europea *CRONOS*, condotta nell'ambito *dell'European Social Survey*, ha chiesto ai cittadini di undici paesi europei se ritengano che l'intelligenza artificiale rappresenti un rischio per i lavori del proprio campo (ESS ERIC, 2025). Il 33,8% degli intervistati percepisce un rischio, quasi sempre nella forma più misurata: il 27,9% ritiene che ne siano a rischio *alcuni*, e solo il 5,9% pensa che lo siano tutti. Il 43,4% degli intervistati, che resta la fascia più ampia, non vede alcun rischio. Quello che emerge è dunque la percezione di una trasformazione selettiva che colpisce solo alcune posizioni lavorative, più che un timore di sostituzione generalizzata.

Un dato in particolare richiama il ribaltamento già osservato nel terzo capitolo. La percezione del rischio non è distribuita in modo uniforme tra i livelli di istruzione: raggiunge il valore più alto proprio fra i più istruiti. Tra chi possiede un titolo di livello terziario avanzato, equivalente alla laurea magistrale o al dottorato, la quota di chi avverte un rischio sale al 42,0%, ben sopra la media del 33,8%. Sul piano soggettivo europeo, in altre parole, si conferma ciò che le ricerche citate nel terzo capitolo mostravano su quello oggettivo italiano: a sentirsi più esposti all'intelligenza artificiale sono anche, e soprattutto, coloro che hanno investito di più in formazione. Va precisato che l'Italia non rientra tra i paesi coperti dall'indagine: il dato offre quindi uno sfondo comparativo europeo, mentre per il caso nazionale restano centrali i dati ISTAT. Resta comunque significativo che la stessa direzione emerga da fonti diverse e con metodi diversi, a conferma della solidità del fenomeno.

4.3 Chi avanza e chi resta indietro: la lettura degli esperti

I dati del paragrafo precedente fotografano un divario che cresce. Quello che non spiegano, però, è perché si produca e chi colpisca. Per questo serve la voce di chi il fenomeno lo osserva da vicino. Le cinque interviste raccolte per questa ricerca coinvolgono testimoni provenienti da ambiti molto distanti tra loro: la formazione manageriale e la ricerca sociologica, l'impresa tecnologica, la consulenza digitale, l'imprenditoria nell'e-commerce. Pur muovendo da posizioni così diverse, le loro analisi finiscono per incontrarsi su un punto. Con l'intelligenza artificiale, la linea di separazione tra chi avanza e chi resta indietro non coincide più con quella del vecchio divario digitale: si ridisegna secondo criteri nuovi.

Il primo punto su cui le voci si ritrovano è lo spostamento del problema dall'accesso all'uso. L'intelligenza artificiale, a differenza delle tecnologie che l'hanno preceduta, ha una soglia d'ingresso bassissima. Antonio Romeo lo dice con chiarezza: «rispetto a internet, rispetto al computer, l'AI è più democratica, ha un gradino di accesso più basso», perché con i modelli conversazionali «il salto cognitivo per usarlo era a zero». Il divario, semmai, si apre dopo. Sta nel saper trasformare quello strumento in produttività, e qui «l'utilizzo banale è alla portata di tutti; portarlo ad aumentare la propria produttività, sapendo come usarla al lavoro, le limitazioni, come scambiare informazioni, ha un gap cognitivo più alto». Castellana individua la stessa difficoltà: non sta nell'usare lo strumento, ma nel «capire come, quando e perché integrarlo nei processi aziendali». E Massimo Ciotta, che osserva il fenomeno dal mondo dell'e-commerce, lo traduce nel linguaggio delle barriere d'ingresso. L'intelligenza artificiale avvantaggia «tutti quei settori dove le barriere d'ingresso si abbassano», ma mette in difficoltà «tutta quella fascia di mercato scarsamente specializzata», i professionisti senza una conoscenza

profonda della propria materia, i più facili da sostituire. Non conta più avere accesso alla tecnologia. Conta saperla usare per produrre valore.

C'è però un'osservazione che ricorre nelle interviste più di ogni altra, ed è anche la più controintuitiva. Riguarda chi, oggi, rischia davvero. Tre interlocutori che non si sono mai parlati, e che guardano il fenomeno da mondi opposti, arrivano alla stessa conclusione: a restare indietro sono proprio i profili qualificati, e non più i meno istruiti come avveniva con le precedenti tecnologie. Castellana sposta il discrimine dal titolo di studio alle attitudini. Rispetto al vecchio divario digitale, «dove il discrimine era il titolo di studio o la disponibilità economica, oggi le linee di confine si ridisegnano sulla base della flessibilità cognitiva e della capacità di problem solving». A faticare sono «i profili professionali, soprattutto quelli con più anni di esperienza lavorativa, più rigidi e legati a mansionari standardizzati», quelli per cui «chi è abituato a eseguire senza interrogarsi sul processo fa molta fatica». Alessandro Ferrari ci arriva dal versante tecnico e nota un rovesciamento rispetto al passato: «noi siamo abituati spesso che il progresso colpisce in maniera più dura le classi meno preparate, meno colte, che hanno studiato di meno, e in realtà questa rivoluzione potrebbe cambiare gli equilibri». Le più esposte, secondo lui, sono «le classi un pochino più abituate a stare bene, come gli sviluppatori, gli informatici, chi lavora nel marketing, chi lavora nelle risorse umane», mentre i lavori manuali restano «più safe, più salve da questa accelerazione». La formulazione più netta è quella di Paolo Perulli, che legge il fenomeno nel tempo lungo e ne misura la portata: «per la prima volta non sono a rischio le fasce di qualificazione medio-bassa, ma sono più a rischio le fasce di qualificazione medio-alta, e questo cambia tutto». Se l'intelligenza artificiale arrivasse a permeare tutte le professioni qualificate, avverte, «si apre una vera e propria voragine dal punto di vista del mercato del lavoro qualificato». E non si tratta di una sua impressione

isolata, perché su questo, aggiunge, esiste «una convergenza da parte di quasi tutti gli studiosi».

Che tre osservatori così lontani arrivino alla stessa conclusione è il dato più rilevante. Si tratta di un'esperta di formazione che ragiona sulle competenze, di un imprenditore tecnologico che osserva il mercato del lavoro dall'interno, e di uno studioso di scienze sociali che lo legge in prospettiva storica. Quella stessa conclusione trova conferma sia nei numeri sia negli studi già citati. La percezione del rischio rilevata in Europa, che il paragrafo precedente ha mostrato più alta proprio tra i più istruiti, punta nella stessa direzione. Lo confermano anche le ricerche sul caso italiano richiamate nel terzo capitolo, che collocavano i lavoratori più qualificati nell'area di maggiore esposizione all'automazione. Quello che gli esperti raccontano dall'interno del mondo del lavoro, in altre parole, i dati e la letteratura lo registrano dall'esterno. L'intelligenza artificiale colpisce per prime le competenze di concetto, e non più solo le mansioni esecutive come accadeva in passato.

Sarebbe però un errore leggere tutto questo come l'annuncio di licenziamenti di massa imminenti. Lo stesso Romeo invita alla prudenza, e distingue ciò che si racconta da ciò che accade davvero: «tutti i licenziamenti annunciati a causa dell'AI sono, tra virgolette, l'AI come scusa: le aziende stanno tagliando personale per risparmiare soldi». Arriva a dire di non conoscere ancora «una singola persona licenziata perché il suo lavoro da domani lo fa l'AI». L'impatto, nella sua lettura, prende un'altra strada. È competitivo, non diretto. Tra due imprese rivali, quella che usa l'intelligenza artificiale per diventare più produttiva finisce per spingere l'altra fuori dal mercato, qualunque cosa decida sul personale. È lo stesso processo che i dati ISTAT del paragrafo precedente fotografano nella forbice tra grandi imprese e piccole e medie imprese: chi non adotta la tecnologia non viene licenziato dall'intelligenza artificiale, ma rischia

di soccombere a chi la adotta. Romeo lo condensa in un'immagine: «se un lavoratore senza l'AI incontra un lavoratore con l'AI, il lavoratore senza l'AI è un disoccupato». Il problema, allora, non è tanto il posto di lavoro di oggi. Sono le competenze di domani.

Sul fronte del lavoro, dunque, le interviste convergono. Toccano molto meno il tema della visibilità, che il terzo capitolo ha indicato come uno dei tre assi dell'AI divide: segno che l'attenzione di chi opera nel settore è oggi puntata soprattutto sugli effetti occupazionali. Un'eccezione viene da Ciotta, che segnala un rischio nuovo nel rapporto tra intelligenza artificiale e reputazione: «chiunque può costruire affermazioni o fatti completamente falsi ma apparentemente veritieri sul conto di un'altra persona». I modelli generativi sanno produrre contenuti credibili, e ridisegnano anche qui le posizioni di chi resta esposto rispetto a chi riesce a controllare la propria immagine, lungo le linee già analizzate nel capitolo precedente. Resta che il nodo più sentito, per chi lavora nel settore, riguarda il lavoro e le competenze. È lì che si gioca, nella lettura degli esperti, la partita più importante dell'equità. E come affrontarla è la domanda da cui prende le mosse il paragrafo successivo.

4.4 Cosa si può fare: il dibattito sulle risposte

Una volta accertato il divario e capito chi colpisce, resta la domanda più difficile: cosa fare. Anche qui i dati offrono un punto di partenza. Tra le imprese italiane che hanno preso in considerazione l'intelligenza artificiale senza poi adottarla, il 47,3% indica come ostacolo la mancanza di chiarezza sulle conseguenze legali, mentre la carenza di competenze adeguate è segnalata dal 58,6% (ISTAT, 2025). Sono due numeri che spostano il baricentro del discorso. L'incertezza sulle regole è essa stessa un freno all'adozione, oltre che un problema di diritti e di garanzie, e

finisce per agire come un motore del divario. Da qui parte il confronto con gli esperti sul terreno delle risposte, e il primo banco di prova è il principale strumento che l'Europa si è data, il regolamento sull'intelligenza artificiale noto come AI Act. La domanda che attraversa tutte le interviste è una sola: basta una buona regolazione?

La difesa più netta dell'AI Act viene da Alessandro Ferrari, e arriva in controtendenza rispetto a un sentire diffuso. «È abbastanza condiviso il malcontento verso l'AI Act, ma io mi sento di essere un pochino in disaccordo», premette. Per quanto migliorabile e vago su alcuni punti, a suo giudizio il regolamento resta «un segnale chiaro da parte della comunità europea di normare qualcosa che ha la potenzialità di una bomba atomica». Ferrari non nasconde il limite strutturale di ogni regolazione in questo campo, e cioè che il legislatore arriva sempre dopo: «la velocità con la quale insegue la tecnologia è sempre stata inferiore, e oggi siamo su ordini di grandezza diversi, perché l'AI di tre settimane fa è già completamente diversa dall'AI di oggi». Eppure, proprio dalla sua esperienza di imprenditore in un mercato regolamentato, quello farmaceutico, ricava l'argomento più originale a favore delle regole. L'AI Act, sostiene, «non va visto solo come un meccanismo che possa frenare l'imprenditorialità europea», perché nei settori dove esistono già vincoli stringenti «paradossalmente diventa un vantaggio competitivo per le aziende che governano questi processi in conformità». Chi è abituato a lavorare dentro un quadro di norme, nella finanza come nel medicale, parte avvantaggiato. La regola, da ostacolo, diventa competenza.

Non tutti, però, condividono questa fiducia. Da più parti arriva un'obiezione di fondo: la regolazione, da sola, non risolve il problema. Iolanda Castellana lo dice in modo diretto. L'AI Act è «un passo fondamentale e necessario dal punto di vista etico e della sicurezza», ma «regolamentare i rischi non significa automaticamente colmare i divari di

competenza». Il motivo è che una norma agisce su un piano diverso da quello dove si gioca davvero il divario: «una norma per sua natura fissa dei limiti, ma non crea cultura d'impresa né capacità di innovazione», e quello dell'AI è prima di tutto «un problema intrinsecamente sociale ed economico, legato alla transizione del mercato del lavoro». Antonio Romeo sposta l'obiezione sul terreno economico. Per lui l'AI Act è «una mossa nella direzione giusta», ma rischia di restare un guscio vuoto se l'Europa non affronta il nodo degli investimenti: «il rischio è che noi abbiamo una legge ricchissima, ma chi sviluppa i modelli se ne frega, perché è dall'altra parte». I progressi nell'intelligenza artificiale, ricorda, dipendono da potenza di calcolo e dati, e quindi da capitali che oggi si concentrano altrove, lasciando l'Europa «totalmente nelle mani degli americani». Paolo Perulli tiene insieme le due cose. Da un lato attribuisce all'AI Act un valore che va oltre il suo contenuto tecnico, perché «l'importanza dell'atto europeo è dimostrata dalle reazioni che ha provocato»: l'Europa che regola è una potenza che frena, e questo «dà molto fastidio agli attori chiave dell'intelligenza artificiale» interessati a combattere ogni vincolo. Dall'altro avverte che il ruolo di regolatore non basta a se stesso: «l'Europa può essere un efficace attore regolativo, ma non deve essere solo questo. Deve anche dire la sua, avere una propria voce».

Se la regolazione non basta, su cosa si deve agire allora? Dalle interviste emergono due direzioni complementari. La prima è quella delle competenze, ed è la più ricorrente. È la stessa che i dati segnalano come primo ostacolo all'adozione, ed è il terreno su cui Castellana costruisce la sua proposta più concreta. Non un'altra legge, dunque, ma «un Piano Nazionale di Alfabetizzazione Funzionale all'AI per le PMI e i giovani professionisti». Quello che manca, nella sua lettura, è «un "ponte" istituzionale che traduca la tecnologia in competenze pratiche per i singoli

settori professionali», costruito appoggiandosi alla rete di incubatori, istituti tecnici superiori e formazione professionale già presente sul territorio. Romeo introduce una seconda direzione, più inattesa, che riguarda i benefici da distribuire e non solo i rischi da contenere. I modelli di intelligenza artificiale, osserva, sono addestrati su «l'intelligenza diffusa degli umani», sul sapere prodotto da tutti, e questo pone un problema che va oltre il diritto d'autore: «se c'è un aumento di produttività grazie all'AI, perché deve andare tutto ad Anthropic, OpenAI, Google? Almeno dovremmo dividerlo». La redistribuzione, precisa, potrebbe tradursi in tempo di lavoro più che in denaro, nel fatto che «non sia più necessario lavorare tutti otto ore». È un terreno, ammette, che la politica non ha ancora toccato. Sono due piste diverse, le competenze e la redistribuzione, ma muovono dalla stessa convinzione: il divario non si chiude con un divieto, e le risposte vanno cercate accanto alla regolazione, oltre che in essa. Su questa convinzione condivisa si fonda la sintesi proposta nel paragrafo conclusivo.

4.5 Una direzione possibile: dalle competenze al territorio

I paragrafi precedenti hanno seguito un percorso preciso. Sul piano dei dati è emerso un divario reale e in crescita, che separa le grandi imprese dalle piccole e attraversa il territorio e le competenze. La natura del divario è stata poi chiarita dalle voci degli esperti, che hanno rivelato un rovesciamento: a restare indietro è oggi chi non sa integrare l'intelligenza artificiale nel proprio lavoro, e non più, come per le tecnologie precedenti, i meno istruiti. E di fronte alla domanda su cosa fare, gli stessi esperti hanno convenuto su un punto, pur partendo da posizioni diverse. La regolazione è necessaria, ma non sufficiente da sola. È da questa convergenza che conviene partire per provare a indicare una direzione.

La leva su cui agire per prima è quella delle competenze. È un'indicazione che viene sia dai dati sui primi ostacoli all'adozione, sia dal rovesciamento descritto nel paragrafo precedente: il confine fra chi avanza nella transizione algoritmica e chi vi resta indietro coincide ormai con la capacità di usare la tecnologia in modo consapevole. Sul terreno delle competenze, la proposta più concreta emersa dalle interviste è quella di Iolanda Castellana, che immagina «un Piano Nazionale di Alfabetizzazione Funzionale all'AI per le PMI e i giovani professionisti», capace di colmare l'assenza di «un "ponte" istituzionale che traduca la tecnologia in competenze pratiche per i singoli settori professionali». È una direzione che questo lavoro fa propria, perché risponde alla radice del problema. Le competenze, in fondo, sono la condizione che rende efficaci tutte le altre risposte. Una norma ben scritta o un incentivo generoso restano lettera morta senza chi sappia applicarli sul campo, e una tecnologia accessibile a tutti, come si è visto, non riduce le distanze se solo alcuni sanno trasformarla in valore. Investire sulla formazione, dalla scuola fino alla riqualificazione di chi già lavora, è la premessa di ogni altra politica sull'intelligenza artificiale.

Accanto a questa leva, già disponibile, c'è una frontiera più avanzata e ancora quasi inesplorata, introdotta nelle interviste da Antonio Romeo. È il tema della redistribuzione. Se l'intelligenza artificiale viene addestrata sul sapere prodotto dall'umanità intera, sostiene, i benefici che ne derivano non dovrebbero concentrarsi per intero in poche mani: «se c'è un aumento di produttività grazie all'AI, perché deve andare tutto a chi sviluppa i modelli? Almeno dovremmo dividerlo». È un'idea che sposta il discorso su un piano nuovo. Oltre a contenere i rischi e diffondere le competenze, si tratta di ripensare il modo in cui i guadagni di produttività generati dall'intelligenza artificiale vengono distribuiti nella società. È il terreno più difficile, quello che la politica, come ammette lo stesso Romeo,

non ha ancora davvero affrontato, e proprio per questo merita di essere indicato come una delle direzioni più promettenti per il futuro. Questo lavoro si limita qui ad accennarne la rilevanza, riservando alle conclusioni una riflessione più distesa sulle forme che una simile redistribuzione potrebbe assumere, all'interno di una cornice regolativa attenta non solo ai rischi ma anche all'equità.

C'è infine una terza direzione, che riguarda non solo l'uso dell'intelligenza artificiale ma anche la sua produzione. È il nodo sollevato da Paolo Perulli quando osserva che l'Europa, e l'Italia con essa, «usa solo le tecnologie degli Stati americani» ed è «totalmente dipendente», priva di una strategia che non sia quella di regolare ciò che altri sviluppano. Perulli indica però anche una via d'uscita possibile, legata ai settori in cui il paese eccelle: il «made in Italy dell'AI» potrebbe servire «per essere fra i primi a sperimentare quello che succede in settori in cui siamo leader», dettando standard applicativi avanzati e non limitandosi a importare soluzioni altrui. Vale la pena estendere questa intuizione oltre i comparti tradizionali a cui l'intervistato la riferiva. Le eccellenze del territorio italiano non si fermano alla moda o all'alimentare. Attraversano anche l'agricoltura di qualità, la ricerca medica e farmaceutica, la difesa, e il settore della tutela e valorizzazione del patrimonio culturale. In ciascuno di questi ambiti l'intelligenza artificiale potrebbe essere sviluppata a partire dal sapere che già esiste e che il paese padroneggia, anziché adottata in forme pensate altrove. È la differenza fra il subire la tecnologia e l'orientarla. Una strategia che investisse contemporaneamente sulle competenze e sulla valorizzazione di questi giacimenti di conoscenza permetterebbe all'Italia di muoversi, almeno in alcuni settori, dalla posizione di utilizzatore dipendente verso quella di produttore capace di imporre i propri criteri.

Alcune di queste direzioni trovano un primo, parziale riscontro nei risultati del questionario esplorativo richiamato in apertura. Pur con tutte

le cautele imposte da un campione ridotto, due indicazioni meritano di essere segnalate. La prima conferma, dal basso, il rovesciamento descritto in questo capitolo: alla domanda su chi rischi di restare indietro, la risposta più frequente non è stata chi possiede un titolo di studio più basso, ma chi dispone di minori competenze digitali a prescindere dal livello di istruzione. La percezione dei rispondenti coincide, su questo punto, con la lettura degli esperti. La seconda indicazione è più ambivalente, e proprio per questo istruttiva. Richiesti di indicare l'intervento prioritario per ridurre le diseguaglianze legate all'intelligenza artificiale, i partecipanti hanno scelto più spesso nuove regole che non investimenti nella formazione. Il dato sembrerebbe contraddire la priorità qui attribuita alle competenze, ma va letto accanto a un altro risultato: la grande maggioranza dei rispondenti si dichiara poco o per nulla informata sulle regole già esistenti. La domanda di nuove norme nasce, almeno in parte, da una conoscenza insufficiente di quelle in vigore, e finisce così per confermare, indirettamente, che il nodo prioritario resta quello della consapevolezza e delle competenze. Più che per i suoi numeri, comunque troppo esigui per trarne conclusioni, questa ricognizione vale come prova di uno strumento: una volta affinato e somministrato a un campione ampio e mirato, un questionario di questo tipo potrebbe affiancarsi all'indagine sulle imprese delineata più avanti, contribuendo a fondare, quella sì, raccomandazioni di policy solide. I dati raccolti finora offrono soltanto un'idea preliminare delle tendenze in gioco; è la ricerca futura, costruita su basi più ampie, a poter trasformare quelle tendenze in indicazioni operative.

Le considerazioni svolte in questo capitolo aprono, infine, alcune piste di ricerca che esulano dai suoi confini ma che da esso prendono le mosse. La prima pista riguarda le imprese. I dati ISTAT documentano l'ampiezza del divario nell'adozione, ma un'indagine sistematica su un

campione ampio e rappresentativo di imprese italiane, articolato per dimensione e settore, permetterebbe di osservarne da vicino le cause e i meccanismi: è un campo dove gli studi dedicati al caso nazionale restano ancora pochi. Una seconda direzione di approfondimento è la dimensione della visibilità, toccata solo in margine dalle interviste e meritevole di un'analisi a sé. Il modo in cui i sistemi generativi ridisegnano la reputazione delle persone, oltre alla loro rappresentazione e al loro accesso all'attenzione, è un terreno ancora poco esplorato e destinato a crescere d'importanza. Vale infine la pena raccogliere l'invito implicito di Perulli, e verificare sul campo se (e in quali settori d'eccellenza) il sistema produttivo italiano stia effettivamente sviluppando applicazioni dell'intelligenza artificiale capaci di dettare standard, oppure se, come l'intervistato temeva, questa possibilità rimanga per ora soltanto un'occasione mancata. Sono domande che questo lavoro non poteva esaurire, ma che indicano la rotta lungo cui proseguire la riflessione.

Conclusioni

Questo lavoro ha preso le mosse da un interrogativo che attraversa il dibattito pubblico sull'intelligenza artificiale, dove la discussione viene spesso ridotta a slogan contrapposti. Per i sostenitori più ottimisti, la tecnologia, abbassando le barriere di accesso al sapere e agli strumenti, avrebbe democratizzato opportunità un tempo riservate a pochi. Per chi è più preoccupato, invece, la stessa tecnologia rischia di scavare divari più profondi anziché livellare il terreno. La tesi ha scelto di non accontentarsi di nessuna delle due narrazioni, e di interrogare il fenomeno con gli strumenti delle scienze sociali. Il percorso si è mosso in modo concentrico, attraversando prima la cornice sociologica sulle disuguaglianze, poi la genealogia del divario digitale fino all'emergere dell'AI divide, e quindi le tre dimensioni lungo cui quel divario si manifesta, fino ad arrivare alla verifica empirica nel contesto italiano ed europeo, attraverso dati e testimonianze dirette.

Quattro domande hanno fatto da guida al lavoro. Sulla natura del fenomeno verteva la prima: *l'AI divide costituisce una frattura inedita rispetto alle disuguaglianze del passato, oppure rappresenta la continuazione, in forma nuova, di dinamiche di esclusione già note?* C'era poi una seconda domanda, sulle manifestazioni concrete del fenomeno: *attraverso quali dimensioni della vita sociale si manifesta l'AI divide?* E sul nodo distributivo verteva la terza: *nella transizione algoritmica, chi avanza e chi resta indietro, e con quali criteri si ridisegnano le posizioni nella società?* Lo sguardo si spostava infine dalla diagnosi alle risposte, con un interrogativo specifico sulle politiche: *di fronte all'AI divide, quali interventi sono possibili sul piano delle politiche pubbliche, e la sola regolazione è sufficiente a contrastarlo?* È alla luce dell'intero percorso che queste domande trovano ora una risposta.

La prima domanda riceve dalla ricerca una risposta netta: l'AI divide è l'ultima forma di una storia molto più lunga, e non una frattura del tutto inedita. Nel secondo capitolo è stata ricostruita la genealogia del fenomeno. Ogni ondata tecnologica, a partire dalla diffusione di massa del personal computer e arrivando fino ai sistemi generativi attuali, ha prodotto un riposizionamento del divario digitale anziché il suo superamento. Lo spostamento è avvenuto per stadi successivi. La prima soglia riguardava la disponibilità di hardware; poi è venuto il momento della qualità della connessione; più tardi è stato il turno delle competenze, prima quelle operative e poi quelle strategiche; e ci si è alla fine concentrati sugli esiti effettivi nella vita offline degli utenti e sul loro grado di consapevolezza dei sistemi automatici. In ogni passaggio, però, le variabili socio-demografiche classiche sono rimaste determinanti: in particolare istruzione e reddito. La chiave di questa continuità sta nella catena causale circolare formulata da van Dijk. Le diseguaglianze sociali generano una distribuzione diseguale delle risorse; questa produce a sua volta accessi diseguali alle tecnologie, e l'accesso diseguale finisce per retroagire sulle stesse diseguaglianze sociali di partenza, rinforzandole. Si tratta del meccanismo che Merton aveva chiamato effetto Matteo, e che attraversa l'intera tesi come filo conduttore. L'intelligenza artificiale eredita questo circolo e ne intensifica gli effetti, anziché spezzarlo. Cambiano i contenuti del divario. La sua logica di fondo continua a operare nello stesso modo. In questo senso il caso italiano si è rivelato un laboratorio privilegiato. Vi convivono in forma simultanea divari che altrove si erano succeduti per fasi distinte: ampie porzioni della popolazione non hanno ancora consolidato le competenze digitali di base, e il sistema deve già affrontare nello stesso tempo le sfide poste dall'AI generativa.

La seconda domanda trova risposta nelle tre dimensioni isolate dal terzo capitolo. Sulla visibilità, è emerso come la selezione algoritmica dei

contenuti rappresenti tutt'altro che un processo tecnico neutrale: si tratta di un meccanismo che distribuisce in modo diseguale la possibilità di accedere a un universo informativo plurale. La visibilità, in altre parole, è l'esito della distribuzione diseguale del capitale algoritmico tra gli utenti, più che una proprietà intrinseca del contenuto. Con l'AI generativa la distorsione si sposta inoltre dai contenuti selezionati a quelli prodotti, amplificando le visioni egemoniche e marginalizzando, nello stesso tempo, le voci minoritarie. Sul lavoro, di cui si dirà tra un momento, è la posizione di ciascuno nei mercati occupazionali trasformati dall'automazione a essere in gioco. Quanto al potere, la genealogia ricostruita ha mostrato un'accumulazione progressiva di forme di controllo algoritmico: il capitalismo della sorveglianza teorizzato da Zuboff, la disumanizzazione statistica delle decisioni automatiche analizzata da Aizenberg e van den Hoven, l'opacità della scatola nera, e infine il potere di selezione della sfera pubblica esercitato sui contenuti. Si tratta di un potere che agisce sulla realtà sociale modificandola attivamente, attraverso le decisioni che produce. Il risultato più importante del terzo capitolo, tuttavia, riguarda non tanto le singole dimensioni del fenomeno, quanto il loro intreccio reciproco. Visibilità, lavoro e potere sono tre facce di un unico meccanismo che si rinforza in modo circolare. Non vanno intese come compartimenti separati. Chi è marginale nei flussi informativi dispone di meno strumenti per costruire le competenze che il lavoro algoritmico oggi richiede. Per chi occupa una posizione debole nel lavoro è poi più difficile comprendere e contestare i sistemi che decidono sulla sua vita. E un soggetto con meno potere è a sua volta più esposto alla marginalizzazione informativa. Il divario tende a chiudersi su se stesso, e la stessa infrastruttura computazionale che produce i contenuti plasmandone la visibilità ridefinisce anche le competenze richieste dal lavoro, e concentra il potere generativo in pochi attori economici globali.

La terza domanda ha riservato la risposta più controintuitiva, ed è forse il contributo più significativo del lavoro. Lo schema classico del divario digitale lasciava prevedere che a soffrire l'arrivo dell'intelligenza artificiale fossero, come sempre, i soggetti meno dotati di strumenti culturali ed economici. I dati raccontano una storia diversa. L'esposizione all'automazione cresce con il livello di istruzione: lo studio italiano di Ferri e collaboratori misura un'esposizione di gran lunga superiore per chi possiede una laurea o un dottorato rispetto a chi possiede titoli inferiori, in un'inversione paradossale del digital divide tradizionale, che storicamente colpiva i meno istruiti. Sarebbe però un errore concludere che siano i più qualificati a perdere. Qui sta una distinzione concettualmente decisiva. Essere esposti a un'innovazione tecnologica e risultarne vulnerabili sono due cose diverse. L'incrocio fra esposizione e complementarità mostra come i profili qualificati siano insieme i più esposti all'AI e i più capaci di adattarvisi: per loro la tecnologia funziona spesso da strumento di potenziamento, mentre la sostituzione tout court rimane un'ipotesi residuale. Il rischio concreto di sostituzione, che riguarda circa un lavoratore italiano su quattro, si concentra invece sulle posizioni intermedie e di routine cognitiva. I meno qualificati restano poco esposti, ma sono anche poco mobili, e per questo risultano fragili nel momento in cui la tecnologia li raggiunge. A questa asimmetria si aggiunge quella salariale documentata da Brunetti e Mussida. Il premio dell'esposizione si concentra nella parte alta della distribuzione mentre si riduce nella zona centrale, ed è questo a determinare l'aggravarsi della polarizzazione nel mercato del lavoro. A ridisegnare le posizioni non è più, semplicemente, il titolo di studio. Conta soprattutto il rapporto che la singola persona ha con la tecnologia: la qualità delle sue competenze digitali, la capacità di usare l'AI in modo complementare al lavoro che svolge, e la collocazione di partenza nella distribuzione dei redditi. Sulla stessa lettura convergono anche le voci degli esperti raccolte nell'ultimo capitolo, lette da osservatori

diversi. E dalla ricognizione esplorativa condotta con il questionario è emerso lo stesso quadro: i partecipanti hanno indicato il discrimine decisivo nelle competenze digitali, con il titolo di studio relegato a un ruolo secondario. Conta piuttosto il rapporto con la tecnologia: il tipo di competenza posseduta, la capacità di usare l'AI come strumento complementare al proprio lavoro, e la collocazione di partenza nella distribuzione dei redditi. Le voci degli esperti raccolte nell'ultimo capitolo hanno confermato questa lettura da osservatori diversi, e anche la ricognizione esplorativa condotta con il questionario ha visto i partecipanti indicare il discrimine decisivo nelle competenze digitali, più che nel titolo di studio. Anche qui opera l'effetto Matteo: a essere premiate o penalizzate sono le posizioni di partenza, che la transizione algoritmica amplifica in entrambe le direzioni, più che le persone in sé.

Resta la quarta domanda, quella sulle risposte, su cui le evidenze dei dati e le testimonianze degli esperti convergono senza eccezioni: la regolazione è una condizione necessaria, anche se non sufficiente da sola. È la conclusione da cui prende le mosse la proposta che chiude questo lavoro, e che conviene articolare con la precisione che il tema richiede.

Il primo livello di ogni risposta è quello regolativo, e il suo strumento principale è l'AI Act: primo tentativo al mondo di disciplinare l'intelligenza artificiale in modo organico attraverso un approccio graduato sul livello di rischio. Si tratta di un risultato di valore. È un segnale politico che afferma il primato dei diritti fondamentali sulla logica dell'innovazione a ogni costo, al punto da provocare la reazione ostile dei grandi operatori del settore. Una reazione che Paolo Perulli legge come la prova migliore della sua efficacia: l'Europa che regola è una potenza che frena, e dà fastidio a chi vorrebbe muoversi senza vincoli. In alcuni comparti già abituati a operare entro quadri normativi stringenti, come il farmaceutico o il finanziario, la conformità può perfino tradursi in un

vantaggio competitivo, come ha osservato Alessandro Ferrari dalla sua esperienza diretta. L'AI Act, però, porta con sé due limiti che vanno detti con chiarezza. Un primo limite, quello strutturale, sfugge a qualsiasi tentativo di perfezionamento normativo. La velocità con cui si muove la tecnologia non è compatibile con quella del diritto, e una norma che richiede anni a essere concepita e applicata si troverà sempre a regolare sistemi nel frattempo cambiati. Lo ricorda Ferrari nell'intervista, quando osserva che l'intelligenza artificiale di tre settimane prima è già diversa da quella di oggi. Il secondo limite è invece sostanziale, e riguarda l'oggetto stesso della regolazione. Una norma fissa limiti e protegge diritti, ma non è in grado di produrre da sé né competenze né capacità produttiva. La regolazione, in altre parole, costituisce la cornice indispensabile della risposta più che la risposta stessa. Affidare a una legge il compito di chiudere il divario significa confondere la premessa con la soluzione.

Sul secondo livello, quello delle competenze, si gioca la vera priorità. Si tratta anche del livello più trascurato nel dibattito pubblico, dove un divieto o una sanzione fanno facilmente più notizia di un investimento formativo. Se il confine fra chi avanza nella transizione e chi vi resta indietro coincide con la capacità di usare l'intelligenza artificiale in modo consapevole, allora la formazione diventa la condizione che rende efficace ogni altra misura, più che una misura fra le altre. La proposta più concreta emersa dalla ricerca è il piano di alfabetizzazione funzionale all'intelligenza artificiale immaginato da Iolanda Castellana per le piccole imprese e i giovani professionisti. Indica con esattezza la direzione: si tratta di costruire un ponte istituzionale che traduca la tecnologia in competenze pratiche settore per settore, appoggiandosi alla rete già presente sul territorio di formazione professionale, di istituti tecnici superiori e di incubatori. Non un intervento calato dall'alto, dunque. È una proposta che questo lavoro fa propria, e che porta alle sue conseguenze.

Per essere all'altezza del problema, un investimento sulle competenze non può limitarsi all'aggiornamento di chi già lavora. Per essere efficace, deve estendersi all'intero ciclo formativo: già dalla scuola dell'obbligo, dove si formano le capacità cognitive di base che permettono ogni apprendimento successivo, e poi lungo tutta la vita lavorativa attraverso forme di riqualificazione continua. In assenza di un'infrastruttura formativa di questo tipo, diffusa sul territorio e capillare nella sua presenza, ogni politica sull'intelligenza artificiale rischia di restare lettera morta. Il divario, in quel caso, continuerà a riprodursi automaticamente lungo le linee delle diseguaglianze già esistenti. A quel punto saranno sempre i contesti già avvantaggiati ad accedere per primi alle competenze che contano: le famiglie più istruite, i territori più ricchi, le imprese che hanno già investito nel digitale.

Il terzo livello è il più difficile, e il capitolo precedente ha potuto soltanto introdurlo: si tratta della distribuzione dei benefici. Il punto sollevato da Antonio Romeo è più radicale di quanto possa apparire a prima vista. I modelli di intelligenza artificiale vengono addestrati su ciò che lui chiama l'intelligenza diffusa degli umani, ossia sui testi e sulle immagini prodotte dall'intera collettività, comprese le conversazioni che ci scambiamo online. Eppure i guadagni di produttività che ne derivano si concentrano nelle mani di un numero ristretto di imprese. È una contraddizione di fronte a cui il diritto d'autore mostra i propri limiti, dal momento che è stato pensato per tutelare singole opere identificabili. Qui in gioco c'è invece un patrimonio cognitivo collettivo, formato grazie al contributo di tutti, e su cui nessuno ha titolo per rivendicare diritti pieni. Affrontare la questione seriamente significa aprire un campo che la politica, per ammissione dello stesso Romeo, non ha ancora osato toccare davvero. Si tratta di ripensare il modo in cui il valore generato dall'automazione tecnologica viene poi distribuito nella società. Le forme

possibili sono diverse, e non sono alternative tra loro. Sul fronte fiscale, una tassazione mirata sui profitti generati dall'intelligenza artificiale potrebbe finanziare proprio quegli investimenti in competenze di cui si è detto. Si creerebbe così un circuito virtuoso tra chi guadagna dalla nuova tecnologia e chi corre il rischio di essere escluso dai suoi benefici. C'è poi la possibilità di una redistribuzione sul fronte del lavoro, che assuma la forma del tempo piuttosto che del denaro: gli incrementi di produttività potrebbero tradursi in una riduzione dell'orario a parità di retribuzione, e l'efficienza guadagnata dall'automazione tornerebbe in questo modo agli individui sotto forma di tempo liberato. È il terreno più impervio, perché tocca gli equilibri di potere economico più consolidati, ma è anche quello su cui si misurerà, nei prossimi anni, una differenza significativa. Un'intelligenza artificiale che contribuisce a ridurre le diseguaglianze esistenti, o un'intelligenza artificiale che le amplifica.

C'è infine una quarta direzione, che risponde al nodo sollevato con forza da Perulli e riguarda non solo l'uso dell'intelligenza artificiale ma anche la sua produzione. La dipendenza quasi totale dell'Europa dalle tecnologie sviluppate oltreoceano è una debolezza strutturale che riguarda l'Italia non meno degli altri paesi europei. Esiste però un margine d'azione contro questa dipendenza, legato ai settori in cui il paese conserva un'eccellenza riconosciuta a livello mondiale. Perulli stesso lo intuisce parlando di un *made in Italy* dell'intelligenza artificiale capace di dettare standard nei comparti in cui siamo già leader, e l'intuizione merita di essere estesa oltre i settori tradizionali a cui di solito la si associa. Le eccellenze del territorio italiano riguardano la moda e l'alimentare, certo, ma il quadro è più ampio. Includono l'agricoltura di qualità; comprendono la ricerca medica e farmaceutica; coinvolgono il settore della difesa e quelli legati alla cura e alla valorizzazione del patrimonio culturale. Sono ambiti in cui l'Italia possiede un sapere specialistico riconosciuto a livello

internazionale, e su cui l'intelligenza artificiale potrebbe essere sviluppata internamente, sulla base di quel sapere, evitando l'importazione di forme pensate altrove e per altri contesti. È la differenza fra il subire una tecnologia di provenienza esterna e l'orientarla in proprio. Più nel concreto, è la differenza tra il rimanere utilizzatori in posizione di dipendenza tecnologica e il diventare produttori capaci di imporre i propri criteri, almeno nei domini in cui il sistema-paese vanta riconosciuti vantaggi. Le quattro direzioni indicate (regolare, formare, redistribuire, produrre) non rappresentano opzioni alternative fra cui scegliere. Sono i pilastri di un'unica strategia complessiva, e ciascuno è fragile in assenza degli altri. La regolazione, in assenza di competenze, è lettera morta; le competenze, senza una qualche forma di redistribuzione, non garantiscono equità sociale; e l'intera costruzione poggia sul vuoto se l'Europa rinuncia a produrre in proprio. È questa interdipendenza fra i quattro pilastri, più di ogni singola misura presa isolatamente, la principale indicazione che emerge dal lavoro.

Resta da dichiarare ciò che questo lavoro non ha potuto fare. Le testimonianze raccolte, per quanto autorevoli, provengono da un numero ristretto di osservatori privilegiati, e non costituiscono un campione rappresentativo. Il questionario, da parte sua, ha avuto il valore di una prova esplorativa, e ben difficilmente potrebbe essere considerato una rilevazione conclusiva. C'è poi un terzo limite, che riguarda l'ampiezza del confronto: lo sguardo si è concentrato soprattutto sul caso italiano, mentre il confronto sistematico con altri contesti nazionali è rimasto sullo sfondo. A questi limiti se ne aggiunge uno che nessuna ricerca su questo tema può evitare, la rapidità con cui il fenomeno si trasforma, e che rende ogni fotografia provvisoria nel momento stesso in cui viene scattata. Sono limiti che indicano la strada a chi vorrà proseguire, piuttosto che indebolire le conclusioni di questo lavoro. La prima pista futura riguarda un'indagine

sistematica sulle imprese italiane, articolata per dimensione e settore e condotta su un campione ampio, che permetterebbe di misurare il divario di adozione con una precisione che qui non era possibile. Il questionario presentato in queste pagine, una volta affinato, potrebbe esserne il primo strumento. C'è poi la dimensione della visibilità, che le interviste hanno toccato solo in margine, e che meriterebbe un'analisi dedicata. La questione diventa ancora più rilevante nel regime generativo, dove i sistemi sono ormai in grado di produrre contenuti in autonomia, e questo aggiunge un livello ulteriore alle dinamiche di selezione algoritmica già discusse. Un ultimo filone di ricerca potrebbe verificare empiricamente sul campo l'ipotesi di un made in Italy dell'intelligenza artificiale, capace di dettare standard nei settori in cui l'eccellenza italiana è già riconosciuta a livello internazionale.

Se un filo unisce l'intero percorso, è questo. L'intelligenza artificiale non distribuisce da sé vantaggi e svantaggi secondo una logica propria e ineluttabile. La tecnologia, da sola, non decide chi resta visibile e chi finisce ai margini dei flussi informativi. Né decide quale lavoro venga potenziato dall'AI e quale invece risulti svalutato, o quali soggetti governino i sistemi e quali altri ne siano governati. A decidere sono piuttosto le strutture sociali su cui la tecnologia si appoggia, e soprattutto le scelte collettive che vengono compiute (o che vengono rimandate). L'AI divide non è scritto in alcun codice. È il risultato, ogni volta rinnovabile, di ciò che le società decidono di fare e di ciò che decidono di non fare. Comprenderlo è il primo passo. Quanto a ciò che viene dopo, la responsabilità non appartiene alle macchine, è nostra.

Appendice delle interviste

Intervistato 1

Founder e CEO di ARGO Vision, docente di Intelligenza Artificiale all'Università di Pavia

Chiamata telefonica, 23/06/2026

L'intervista è stata condotta e registrata con il consenso dell'intervistato all'uso del nome, del ruolo e dei contenuti ai fini del presente lavoro.

Intervistatore: La ringrazio. Prima di iniziare le riporto un paragone emerso in un'altra intervista, fatto da una marketing manager: è come se da un giorno all'altro ci fossero state date le chiavi di una supercar, mentre fino al giorno prima guidavamo una Panda. Non si sapeva come usarla, bisognava imparare. La prima domanda è questa: si inizia a parlare di una nuova frattura legata all'intelligenza artificiale, l'AI divide, che non riguarda più tanto l'accesso alla tecnologia quanto la capacità di trarne un beneficio. Dal punto di vista del suo lavoro, è un fenomeno che incontra concretamente, e in che forme lo vede manifestarsi?

Intervistato: Io vengo da un mondo prettamente informatico e tecnologico; quindi, sono molto ben predisposto ad assorbire ciò che è migliorativo dal punto di vista tecnico. Per noi che veniamo da questo settore è molto più semplice trarne un beneficio tangibile. La velocità con cui queste tecnologie stanno impattando le abitudini e i processi lavorativi è però qualcosa di complesso, che non necessariamente si traduce in un beneficio immediato o nel medio periodo, perché i processi a cui siamo abituati non sono pronti ad accogliere questo tipo di superpotere. La tecnologia dell'intelligenza artificiale è incredibile, ma non è così scontato che il beneficio sia tangibile o duraturo in tutti i settori in cui andremo ad adottarla. È un tema di processi, di regolamentazione, di abitudini, e anche della capacità dell'essere umano di comprendere che è uno strumento

talmente dirompente da dover cambiare radicalmente anche i processi che gli girano intorno.

Intervistatore: Quindi l'AI non sembra colpire tutti allo stesso modo, e può creare divari rispetto a situazioni preesistenti: pensando ad aziende che non avevano in precedenza un livello tecnologico avanzato, o che sono solo in fase di avvio.

Intervistato: Negli ultimi vent'anni abbiamo assistito a due momenti. Lo dico per l'Italia, ma il paragone è probabilmente estendibile anche ad altri paesi. Circa trent'anni fa è iniziato in maniera drastica il processo della digitalizzazione, che però è stato un'occasione persa, perché è stato un processo molto lento e non sempre ha portato a risultati eccezionali. Oggi invece l'AI è diventata agli occhi di quasi tutti il sacro graal della produttività, e questo ha avviato un processo di vera digitalizzazione dell'azienda: ci si è resi conto che per adottare l'AI serve uno step precedente, cioè la disponibilità e l'organizzazione dei dati, il cambiamento dei processi. Da questo punto di vista l'AI è diventata un volano per un progresso non necessariamente legato all'uso di modelli di AI, ma che ha reso le aziende consapevoli che la digitalizzazione non è più rimandabile. Per questo credo che il beneficio sarà abbastanza condiviso. È chiaro che i divari ci sono sempre stati nella storia dell'umanità e continueranno a esserci, ma l'AI rimescola un po' le carte. Noi siamo abituati a pensare che il progresso colpisca in maniera più dura le classi meno preparate, meno colte, che hanno studiato di meno; in realtà questa rivoluzione potrebbe cambiare gli equilibri, toccando classi più abituate a stare bene, come gli sviluppatori, gli informatici, chi lavora nel marketing o nelle risorse umane, cioè tutti quei processi che oggi cerchiamo di accelerare o di demandare all'intelligenza artificiale. Altre classi, più legate a lavori manuali o meno di concetto, sembrano essere un po' più al sicuro da questa accelerazione improvvisa. Il divario, quindi, è

sempre esistito e continuerà a esistere, ma non è detto che le componenti che lo caratterizzano siano quelle a cui siamo stati abituati.

Intervistatore: Ha già risposto in parte alla seconda domanda. Sul piano normativo, noi come Unione Europea abbiamo l'AI Act, uno strumento che prova a regolare i sistemi di intelligenza artificiale in base al loro rischio. Secondo lei è sufficiente a ridurre l'AI divide, oppure bisogna introdurre altri strumenti, normativi o più concreti, che aiutino le aziende in questa fase di introduzione dell'AI?

Intervistato: Premesso che il malcontento verso l'AI Act è abbastanza condiviso nell'opinione pubblica, io sinceramente mi sento un po' in disaccordo con questo malcontento. L'AI Act, per quanto migliorabile, per quanto vago su alcuni aspetti e difficile da recepire su altri, è comunque un tentativo, «un segnale chiaro da parte della comunità europea di normare qualcosa che ha la potenzialità di una bomba atomica». Nel mondo della tecnologia è sempre stato evidente che il legislatore non ha tempi compatibili con lo sviluppo tecnologico: la velocità con cui insegue la tecnologia è sempre stata inferiore, e oggi siamo su ordini di grandezza diversi, perché l'AI di tre settimane fa è probabilmente già completamente diversa da quella di oggi. I tempi della genesi di un framework normativo non sono paragonabili a quelli della nascita di modelli nuovi che possono rivoluzionare il settore. L'AI Act, d'altra parte, si concentra su aspetti fondanti e non su aspetti pratici o specifici: è molto focalizzato sulla libertà, sulle discriminazioni, per evitare che certe derive avvengano. Quindi, pur con difetti migliorabili, è un tentativo notevole di porre un freno a una deriva rischiosa, che vediamo nel mondo dell'informazione oltreoceano e nei sistemi di videosorveglianza di massa del mondo cinese, dove l'utilizzo dell'AI è fuori controllo da tanti punti di vista. Questa è la mia valutazione, da europeista convinto. Dall'altra parte, l'AI Act non va visto solo come un meccanismo che frena l'imprenditorialità europea: ci

sono settori dove è esattamente il contrario. Nei mercati regolamentati come quello in cui lavoro con la mia azienda, il farmaceutico, l'AI Act «paradossalmente diventa un vantaggio competitivo per le aziende che governano questi processi in conformità» ai framework che normano l'intelligenza artificiale. Chi lavora nei mercati normati, come il finanziario, il medicale, il farmaceutico, è già abituato ad avere a che fare con framework legali che governano certi processi. Da questo punto di vista può essere un valore aggiunto, un vantaggio competitivo, e non solo un ostacolo.

Intervistatore: Secondo lei, nei prossimi anni, il ruolo dell'AI in Italia ridistribuirà nuove opportunità o alimenterà nuove diseguaglianze sociali? E c'è una priorità che metterebbe sul tavolo del governo?

Intervistato: Io faccio un altro mestiere, non sono un politico, quindi è difficile per me indicare una priorità. Quello che posso dire è che sarà molto complesso e delicato gestire l'ipertrofia che l'AI genererà nel libero mercato, perché l'AI è un abilitatore di processi esponenziale e riduce spesso la componente di lavoro umano necessaria a ottenere lo stesso risultato. La scelta sarà quindi tra aumentare la produttività, il fatturato e la ricchezza mantenendo costante l'occupazione, oppure ottimizzare i costi aumentando gli EBITDA a discapito della forza lavoro umana. Questo è il rischio che esiste oggi. Per me è indubbio che saranno necessarie politiche di contenimento delle insicurezze sociali, perché il rischio lo vediamo già negli Stati Uniti, dove il mercato del lavoro è più liquido e le grandi aziende non si pongono problemi a lasciare a casa il venti, venticinque, trenta per cento della forza lavoro in un giorno, perché l'AI sostituisce le persone nei processi. Poi, come è sempre accaduto, nel medio periodo svilupperemo gli anticorpi per trasformare questo strumento in una parte del nostro ecosistema culturale, sociale ed economico. L'unico rischio vero, oggi, ed è un mio parere personale, è una

cosa che secondo me non è mai successa nella storia dell'umanità: è una tecnologia che ha una viralità e una velocità di adozione che non abbiamo mai visto. Se pensi al fuoco, all'aviazione, all'automobile, alle medicine, ai vaccini, a tutte le invenzioni che l'essere umano ha governato, credo di non sbagliare dicendo che nessuna ha avuto l'esponenzialità nell'adozione dell'intelligenza artificiale. Ricorda un po' il Covid: non è che l'AI non esistesse fino a ieri, esiste da ottant'anni in forme molto diverse, ma a un certo punto si arriva a un punto di rottura in cui, come nel Covid, si passa da mille contagiati a cento milioni nel giro di qualche settimana. Allo stesso modo, grazie a ChatGPT, l'AI è passata dall'essere una tecnologia di nicchia a essere estremamente popolare. Io ci lavoro da circa vent'anni e, pur avendo l'esperienza per fare previsioni, faccio molta fatica a dire cosa succederà tra sei mesi. Questo ci racconta di una tecnologia intrinsecamente esponenziale, che come tutte le tecnologie dirompenti, con risvolti militari, sociali ed economici, sarà in qualche modo governata dalle istituzioni, statali o internazionali. L'unico rischio, ripeto, è nel breve periodo non riuscire a governare bene il processo, perché è molto più veloce di quanto siamo abituati, e il Covid ci ha insegnato che quando si prendono decisioni rilevanti per l'umanità in pochi giorni non sempre si fanno le scelte giuste.

Intervistatore: Al momento non è neanche semplice prevedere quali siano le decisioni più giuste per questa tecnologia, perché saranno scelte politiche.

Intervistato: Derivanti da quali sono le priorità. In Europa la priorità è difendere la privacy, la salute, la libertà di pensiero; in altri contesti le priorità sono la sicurezza nazionale o il prodotto interno lordo. Dipende realmente dalle priorità che gli stati o le organizzazioni interstatali si daranno su questo tema.

Intervistatore: È sostanzialmente quello che diceva lei. Stavo leggendo un libro di un sociologo dell'economia, un professore dell'Università del Piemonte Orientale cheavrò il piacere di intervistare, che parla di Peter Thiel e di come la sua visione sull'AI sia molto improntata a un mindset americano, orientato anche al controllo e all'utilizzo militare. Gli altri stati hanno priorità diverse rispetto all'Unione Europea, e Thiel vede la regolazione europea come una limitazione sostanziale degli utilizzi possibili dell'AI.

Intervistato: Non è necessariamente un'accezione negativa che gli altri stati abbiano altre priorità. Ognuno vive in un contesto sociale diverso. Noi siamo tendenzialmente occidentalocentrici e crediamo che solo il nostro orticello sia quello giusto, ma ci sono paesi nel mondo dove la democrazia non è vista come un valore, o dove l'uguaglianza tra uomo e donna non è considerata un valore necessario. Io so dove voglio vivere, e vivo qui in Europa, ne sono convinto; però quando si analizza un fenomeno globale di questa portata bisogna avere almeno il buon gusto di porsi il problema che non tutti la pensano come noi. L'errore di credere che il mondo sia eurocentrico lo stiamo già facendo da diversi anni, e ci stiamo accorgendo che il mondo va in un'altra direzione, con i BRICS, con la potenza della Cina, con gli Stati Uniti. L'AI è una delle tecnologie strategiche, come il nucleare, come la filiera farmaceutica, come l'approvvigionamento energetico. Ognuno ha le sue priorità, e noi dobbiamo essere consapevoli che è un fenomeno da governare a livello planetario, comprendendo anche le sensibilità diverse degli altri paesi.

Intervistatore: La ringrazio, professore. È stato un apporto importante per la tesi, perché il suo è il punto di vista di una persona che lavora costantemente con l'AI e nel settore farmaceutico. La ringrazio per la disponibilità.

Intervistato 2

*Vice President, Presales | SaaS & Cloud Strategist | Driving Decision Excellence
| AI & Cybersecurity Advocate*

Chiamata online, 23 giugno 2026

L'intervista è stata condotta e registrata con il consenso dell'intervistato all'uso del nome, del ruolo e dei contenuti ai fini del presente lavoro.

Intervistatore: La prima domanda si basa proprio su questo AI divide, questa nuova frattura legata all'intelligenza artificiale che non riguarda più tanto l'accesso alla tecnologia, quanto la capacità di trarne un beneficio. Dal punto di vista del suo lavoro, è un fenomeno che incontra concretamente, e in che termini lo vede manifestarsi?

Intervistato: Io lavoro con aziende che, tra virgolette, sono tra i casi più avanzati, perché non hanno il grosso problema dei costi. Però volevo capire meglio la domanda: quando parliamo di divide, intende la difficoltà di accedere per ragioni di cultura, di skill, di preparazione, come vogliamo chiamarla, piuttosto che la possibilità di accesso dal punto di vista tecnologico o economico? Per esempio, c'è stato il taglio dell'accesso ai sistemi cloud più avanzati da parte degli Stati Uniti. Su quale aspetto si sta concentrando di più?

Intervistatore: Si sta concentrando più sul primo aspetto. Siamo partiti banalmente da quello che poteva essere un primo divario negli anni Novanta, tra chi aveva — in quel caso era proprio economico — l'accesso a comprare un computer e chi no. Poi negli anni Duemila si è rimodulato: quasi tutti possono acquistare un computer, però c'è chi lo sa usare e chi no. L'evoluzione che stiamo studiando è questa: come questo divario, che in realtà è sempre stato lineare nel tempo, si sta sviluppando ora con l'AI, e se conseguentemente possa portare a — mi passi il termine — una

distruzione di posti di lavoro, una perdita di posti di lavoro, ma magari allo stesso tempo alla creazione di posti nuovi.

Intervistato: Sicuramente, secondo me, da un certo punto di vista, rispetto a internet e rispetto al computer — per tornare alle rivoluzioni che le menzionavo — l'AI è più democratica, nel senso che ha un gradino di accesso più basso. Io vedo utilizzare l'AI da chiunque: dalla massaia che chiede la ricetta al grande manager che cerca di capire l'utilizzo migliore per tirare fuori delle informazioni. Quindi, dal punto di vista del puro accesso, al di là del costo... I computer prima, internet poi, hanno sempre avuto un alto gradino di apprendimento richiesto. Internet si è diffuso, poi adesso con l'arrivo dei telefonini e delle app molte cose sono diventate comuni; però la verità è che con la cosiddetta digitalizzazione di tutti i servizi, pubblici e non, ci stiamo lasciando indietro delle persone che non hanno le conoscenze sufficienti. Sono gli anziani che devono fare un certificato e gli dici «vai online», e sono obbligati a chiamare il nipote o l'amico per farlo, perché adesso non si riesce più a parlare con un umano. E questo mentre, da una parte, per me utente avanzato — tra virgolette — preferisco magari interagire prima con il sito web; ma questo non è necessariamente vero per tutti. Da questo punto di vista, secondo me, qui ancora non è chiara l'evoluzione, perché, a differenza dei salti tecnologici precedenti, chiunque può accedere, e anche il costo per il momento non è proibitivo: anche con le AI gratuite si riescono a fare delle cose ragionevoli, come sempre a condizione di essere disposti a cedere i propri dati a chi che sia. Però dal punto di vista economico, ma soprattutto dal punto di vista delle skill, secondo me qui è il grosso differenziale rispetto alle rivoluzioni tecnologiche precedenti — il pc, internet, il cellulare, se vogliamo fare i tre macro-trend. Non so se li riassumo bene.

Intervistatore: Certo, assolutamente.

Intervistato: Questa è stata infatti la cosa grande di ChatGPT nel giorno del lancio: nella prima settimana ha avuto il più alto numero di iscritti di qualunque altra tecnologia. Anche Facebook, mi pare, ci ha messo otto-nove mesi ad arrivare a cento milioni di iscritti — se adesso non ricordo male i numeri — invece ChatGPT ci è arrivato in una settimana, perché il salto cognitivo per usarlo era a zero: potevi finalmente parlare il linguaggio naturale, che è sempre stato un problema delle tecnologie.

Intervistatore: Certo, quindi l'accesso in questo caso è stato più agevolato rispetto...

Intervistato: Alla portata di tutti, sì. Ma poi c'è chi non lo vuole usare, chi non si fida, chi non vuole neanche accendere il computer o il telefonino per accedere all'intelligenza artificiale; questo vale per qualsiasi tecnologia. Però secondo me questo è un grosso differenziale. Se devo pensare all'AI divide, io penso che sia più nella capacità di trarre direttamente vantaggio, che richiede un po' più di conoscenza. Di sicuro, se io continuo a fare il mio lavoro, qualunque esso sia, senza saper utilizzare le AI, tra qualche anno ci sarà un'ecatombe di posti di lavoro. Anzi, diciamo così: come in un vecchio film western, dove se un uomo con la pistola incontra un uomo con il fucile, l'uomo con la pistola è un uomo morto; allo stesso modo, se un lavoratore senza le AI incontra un lavoratore con le AI, il lavoratore senza le AI è un disoccupato. È importante che si creino dei programmi, altrimenti questo avverrà. Mentre l'utilizzo banale — «voglio la ricetta, fai la domanda», certi GPT che ormai hanno tutti — portare questo ad aumentare la propria produttività, quindi sapere come usarla al lavoro, sapere le limitazioni, come scambiare informazioni, nell'utilizzo produttivo richiede un gap cognitivo un po' più alto. Per cui potrebbe creare dei settori in cui io vado al lavoro e non so che farci con le AI. E poi lì dipende da vari fattori, tra cui il primo è il livello di intelligenza — tra virgolette — delle AI stesse, che devo dire è

incrementato tantissimo negli ultimi tre anni, ma a ogni iterazione l'incremento si fa meno importante, se vogliamo. Io al momento, devo dire, ho l'impressione — nell'industria, non parlo direttamente — che tutti questi pretesi licenziamenti a causa delle AI... Io sono stato in Microsoft, sono stato in Dell, conosco un po' l'industria in generale, le grosse società, ed è come quelle cose, mi ha detto mio cugino: le aziende stanno tagliando personale per risparmiare soldi, da investire poi magari nelle AI — questo sì — ma il taglio del personale... le AI sono una scusa perfetta. Io al momento non conosco una persona licenziata, una singola persona licenziata perché «il mio lavoro da domani lo fa l'AI».

Intervistatore: Ma questo è importante. Perché appunto è una credenza: sentendo parlare tante persone, sembra quasi che l'intelligenza artificiale sia arrivata per chiudere posti di lavoro e basta. La ricerca che stiamo conducendo è proprio questo, cioè cercare di capire se in realtà questa cosa possa avere delle fondamenta, e quanto — molto probabilmente, come è successo per tantissimi lavori.

Intervistato: In questo momento, secondo me, ancora no. Ripeto: tutti i licenziamenti annunciati a causa delle AI sono, tra virgolette, le AI usate come scusa — «non devo discutere con i sindacati», perché se dici che licenzi a causa delle AI... purtroppo, se non lo faccio io lo fanno i miei competitor, e sono fuori: questa è la scusa, tra virgolette. Di certo, però, c'è una cosa: è chiaro che se in generale aumenta la produttività — quelle cose che prima richiedevano venti persone adesso le fanno in cinque — un impatto ci sarà. La domanda è... penso, per esempio, alla programmazione. Nel campo della programmazione ci sono molti dibattiti, però oggi un programmatore con le AI scrive codice a una velocità impensabile. Certo, ha altri problemi: quel codice bisogna testarlo, ma [...]. Al di là di quello, è vero; però è vero anche un altro aspetto. Diciamo che io e lei siamo due competitor nel campo del software.

Io ho dieci dipendenti che improvvisamente scrivono codice con le AI, e lei ha dieci dipendenti che scrivono codice con le AI. Il problema è che, anche se con le AI i miei dieci diventano molto più produttivi e io ne licenzio cinque, lei ha dieci dipendenti che diventano molto più produttivi e non ne licenzia nessuno: io fra due anni sono fuori mercato, perché lei tirerà fuori funzionalità a una velocità quadrupla rispetto a me. Poi è chiaro: se fossimo in altri settori... Microsoft potrebbe licenziare metà degli sviluppatori di Windows, perché Microsoft non ha competitor nel mercato del sistema operativo; però se Apple mantiene tutti i suoi sviluppatori e dopo due o tre anni il sistema operativo di Apple diventa significativamente migliore di quello di Microsoft, Microsoft va fuori mercato. Quindi secondo me c'è anche questo: un conto è l'aumento della produttività, un conto sono alcune funzioni — tra virgolette — di bassa intelligenza. Se devo fare la classificazione delle fatture e metterle in ordine in un folder a mano, con le AI invece di dieci dipendenti me ne bastano cinque; ma lì stiamo parlando di task ripetitivi che comunque non portano molto valore. Dove c'è il valore, invece, è vero che aumenta la mia produttività: ci sono molte cose che adesso posso sviluppare a casa da solo e che prima dovevo per forza far fare a una software house, o fare a mano, o non fare. Poi ci sono sempre i casi limite, per carità.

Intervistatore: Assolutamente. Lei ha già risposto a due delle quattro domande, quindi anche alla seconda. Passerei all'ambito un po' normativo. Presuppongo che lei lo conosca anche per il suo lavoro, quindi immagino abbia diverse conoscenze in merito. Ho avuto altre interviste che mi hanno dato pareri concordanti su alcune cose e discordanti su altre. Secondo lei una risposta normativa come quella attuale è sufficiente per cercare di ridurre questo gap, in questo caso proprio di conoscenze? Oppure, come un po' mi aveva già accennato, è un problema sociale ed economico e

quindi la norma serve a poco? E, a livello normativo europeo — poi parleremo dell'italiano — serve qualcosa di più concreto?

Intervistato: Io penso che sia una mossa nella direzione giusta. Se sia sufficiente, secondo me è troppo presto per dirlo; chiunque dica di saperne di più... Io non sono un esperto, quindi è un parere da privato cittadino, non da avvocato. Per quanto sia ancora troppo presto, penso che le conseguenze... Sicuramente sono stati messi dei controlli e delle cose che molti vedono come burocrazia, ma che di fatto, secondo me, prima o poi dovrebbero essere in qualche modo chiarite, perché servono: questo è il breve periodo. L'altro aspetto che ritengo necessario, dal punto di vista dell'Unione Europea, è avere una visione più chiara di dove vogliamo investire su questi aspetti, perché per adesso siamo totalmente nelle mani degli americani, con qualche piccola eccezione. Siccome finora tutti i progressi nelle AI si sono basati fondamentalmente sulla potenza di calcolo e sulla quantità e qualità dei dati a disposizione, e queste due cose si traducono fondamentalmente in soldi — entrambe sono direttamente proporzionali a quanti miliardi di dollari voglio investire nel nuovo modello — ci sono dei modelli open source che seguono, diciamo, con uno o due anni di ritardo le stesse capacità. Quindi per il momento c'è ancora spazio, secondo me, per entrare e non essere lasciati troppo indietro. Ma il rischio è che noi abbiamo una legge ricchissima e poi chi sviluppa i modelli se ne fregghi, perché è dall'altra parte. La cosa che non è chiara — non ce l'ha nessuno, né noi né gli americani — è che, siccome i modelli stanno usando l'intelligenza diffusa degli umani, quindi quello che ha scritto lei magari in una chat senza sapere che sarebbe finita in pasto al modello, quello che ho scritto io, quello che hanno scritto gli autori... la vedono tutti come un problema di diritto d'autore, ma secondo me va oltre. Adesso, quando si tratta di diritto d'autore, chi è forte riesce a negoziare: la SIAE va da Anthropic e gli dice «dammi i 100 milioni», e

siamo a posto. Ma chi non è coperto da queste cose? Se lì dentro c'è una mia chat che ha addestrato le AI, in parte ho contribuito con un neurone, se dobbiamo metterla così. Allora il risultato dovrebbe essere che, se tu stai usando le AI, anche il mio neurone — cioè il bene pubblico della conoscenza diffusa dell'umanità — dovrebbe in qualche modo rientrare: ci dovrebbe essere un modo in cui parte dei profitti torna indietro. Di solito questo si fa con le tasse, però ancora, in tutta questa modernità, secondo me c'è un po' il famoso robber baron dell'Ottocento: quando si è diventati ricchi, ormai chi ha avuto ha avuto, chi ha dato ha dato. E come società rischiamo di trovarci in un momento in cui chi ha sfruttato la cosa di tutti poi non contribuisce a nessuno. Per di più ti porti dentro l'inquinamento, perché servono nuove centrali nucleari aggiuntive, e c'è l'acqua che viene drenata per il raffreddamento di queste cose. Questi fanno billion e billion di profitti, creando eventualmente posti di lavoro: bisogna fare un match. Ok, aumenta la produttività grazie alle AI, ma perché questo aumento di produttività deve necessariamente andare tutto ad Anthropic — dico Anthropic per dire, OpenAI, chiunque?

Intervistatore: Diciamo che Anthropic è il punto, quello di adesso.

Intervistato: Però OpenAI è lo stesso, Google con Gemini è lo stesso. Questi signori usano la conoscenza dell'essere umano, la conoscenza diffusa dell'umanità. Se c'è un aumento di produttività grazie a questo strumento, almeno dovremmo quantomeno dividerlo. E non sto parlando di un rientro in forma di denaro: se c'è un aumento, potrebbe essere che non sia più necessario lavorare tutti otto ore, o che si lavori meno, perché adesso lo fanno le AI — invece di avere tutti gli altri disoccupati. Come suddividere questo impatto è un problema che secondo me, al momento, non è stato completamente toccato.

Intervistatore: Quindi potrebbe essere un punto cardine, sia a livello europeo sia forse anche a livello italiano. Secondo lei potrebbe essere una

priorità che la nostra legislazione dovrebbe prendere sul serio, oppure ce ne sono altre prima, sempre legate all'AI e all'AI divide?

Intervistato: L'aspetto dell'impatto sulla produttività e sul lavoro, secondo me, dovrebbe essere una priorità dell'Unione Europea. Il mercato italiano, ormai, devo dire... l'Italia, o la Spagna, o la Francia da soli — tra virgolette — contano meno di Google, molto meno, in termini di potere. Non sto dicendo che sia giusto così, però è difficile avere un peso specifico. L'unica arma, tra virgolette, è: «vuoi accedere ai miei mercati, ti devi adeguare alle mie regole». Ma se lo fa solo l'Italia... Dall'altra parte la loro arma è: se i tuoi produttori di automobili non usano le AI, sono più costose dei miei produttori di automobili, e tra dieci anni non ne venderai più una. Bisogna bilanciare questi due aspetti.

Intervistatore: Perfetto, direi che ci siamo: le quattro macro-aree sono state toccate. La ringrazio molto, è stato un contributo importante per la tesi.

Intervistato 3

Marketing project manager per Scuola di Scienze Aziendali e Tecnologie Industriali Piero Baldesi

Chiamata online, 22 giugno 2026

L'intervista è stata condotta e registrata con il consenso dell'intervistata all'uso del nome, del ruolo e dei contenuti ai fini del presente lavoro.

Intervistatore:. Si parla sempre più di una nuova frattura legata all'intelligenza artificiale, un AI divide che non riguarderebbe più tanto l'accesso alla tecnologia, quanto la capacità di trarne beneficio. Dal punto di osservazione del suo lavoro, è un fenomeno che incontra concretamente? Mi può raccontare in che forme lo vede manifestarsi?

Intervistato: Mi interfaccio con questo fenomeno quotidianamente, ma in quanto responsabile marketing e formazione di una business school vedo una declinazione ben precisa: l'AI divide oggi non è una questione di accesso o di costi, bensì di integrazione e competenze. Come confermano anche i dati del Sistema Informativo Excelsior 2025, la vera barriera non è economica, ma culturale e di alfabetizzazione funzionale. Nel mio lavoro, sia nella formazione dei profili manageriali sia nel supporto ai team che vogliono lanciare una startup, vedo manifestarsi questo divario non tanto nell'incapacità di usare lo strumento in sé — aprire ChatGPT o generare un'immagine è ormai alla portata di tutti — quanto nella difficoltà strategica di capire come, quando e perché integrarlo nei processi aziendali. C'è un disallineamento tra la potenza tecnologica a disposizione e la consapevolezza del suo utilizzo. Mi piace usare una metafora: è come se fossimo abituati a spostarci a cavallo e improvvisamente ci venisse consegnata una Ferrari. Sappiamo che va fortissimo, ma non sappiamo come si guida e, paradossalmente, facciamo fatica a metterne a fuoco la reale utilità per le nostre mansioni specifiche di tutti i giorni. Il divario si

crea tra chi sa mappare i propri flussi di lavoro e ottimizzarli grazie all'AI e chi, invece, si limita a un utilizzo superficiale o, al contrario, la rifiuta per timore.

Intervistatore: L'intelligenza artificiale non sembra colpire tutti allo stesso modo: più che creare un divario nuovo, tende a riprodurre e talvolta ad amplificare diseguaglianze che già esistono. C'è anche un aspetto che appare controintuitivo, e cioè che con l'AI non sono necessariamente i soggetti più deboli o meno istruiti a restare indietro, come accadeva con il vecchio divario digitale, ma le linee di vantaggio e svantaggio si ridisegnano in modo nuovo. Nel suo campo, chi sono oggi i soggetti che rischiano di restare indietro e chi invece ne esce avvantaggiato?

Intervistato: Questo è un punto cruciale. Rispetto al vecchio digital divide, dove il discrimine era il titolo di studio o la disponibilità economica, oggi le linee di confine si ridisegnano sulla base della flessibilità cognitiva e della capacità di problem solving. Nella nostra esperienza, i soggetti che rischiano paradossalmente di restare indietro non sono necessariamente i meno istruiti, ma i profili professionali — soprattutto quelli con più anni di esperienza lavorativa — più rigidi e legati a mansionari standardizzati, che vedono l'AI come una minaccia o come un semplice «generatore di testi» senza comprenderne il potenziale di co-pilota. Chi è abituato a eseguire senza interrogarsi sul processo fa molta fatica. Non stiamo parlando di profili per forza anziani o prossimi alla pensione: la difficoltà appare già anche a distanza di cinque o dieci anni dall'ingresso nel mondo del lavoro. Al contrario, chi ne esce fortemente avvantaggiato sono le figure con forti competenze trasversali, capacità critiche e attitudine all'auto-apprendimento. Nel mondo delle startup, ad esempio, l'AI è un incredibile democratizzatore: permette a piccoli team o a singoli founder, spesso giovani o con risorse limitate, di validare un'idea di business, fare analisi di mercato o creare prototipi a una velocità

impensabile fino a pochi anni fa. Il vantaggio competitivo va a chi sa fare le domande giuste — il cosiddetto prompting, che richiede logica e visione d'insieme — e sa integrare l'output dell'AI nel proprio valore umano.

Intervistatore: L'Unione Europea ha approvato l'AI Act, che prova a regolare questi sistemi in base al loro livello di rischio. Secondo lei una risposta di questo tipo è sufficiente a ridurre l'AI divide, oppure siamo di fronte a un problema soprattutto sociale ed economico, su cui la sola norma può poco? E se dovesse indicare uno strumento concreto, normativo o di altro tipo, che oggi manca e che servirebbe davvero, quale sarebbe?

Intervistato: L'AI Act europeo è un passo fondamentale e necessario dal punto di vista etico e della sicurezza, ma regolamentare i rischi non significa automaticamente colmare i divari di competenza. Una norma per sua natura fissa dei limiti, ma non crea cultura d'impresa né capacità di innovazione. Siamo davanti a un problema intrinsecamente sociale ed economico, legato alla transizione del mercato del lavoro. Se dovessi indicare uno strumento concreto che oggi manca e che servirebbe davvero, non penserei a un'altra legge, ma a un Piano Nazionale di Alfabetizzazione Funzionale all'AI per le PMI e i giovani professionisti. Manca un «ponte» istituzionale che traduca la tecnologia in competenze pratiche per i singoli settori professionali. Servirebbero incentivi e infrastrutture diffuse — magari appoggiandosi proprio a incubatori, ITS e scuole di formazione sul territorio — per aiutare i professionisti e le micro-imprese a mappare le proprie attività e capire, concretamente, quale specifico strumento AI fa al caso loro. Dobbiamo insegnare a guidare quella famosa Ferrari, altrimenti rimarrà chiusa in garage per paura di fare incidenti, o regolata al punto da non poter accelerare.

Intervistatore: Guardando ai prossimi anni in Italia, qual è la cosa che la preoccupa di più rispetto al modo in cui l'intelligenza artificiale

ridistribuirà opportunità e diseguaglianze? E se potesse mettere una sola priorità sul tavolo di chi deve decidere le politiche pubbliche, quale indicherebbe?

Intervistato: La mia preoccupazione più grande per i prossimi anni in Italia è il rischio di una polarizzazione del tessuto produttivo. Da un lato, poche grandi aziende e startup native digitali capaci di correre a velocità esponenziale grazie all'automazione; dall'altro, la stragrande maggioranza delle nostre piccole e medie imprese tradizionali, che rischiano l'estraniamento competitivo per mancanza di competenze interne, aumentando il divario produttivo del Paese. Se potessi mettere una sola priorità sul tavolo del Governo, indicherei la riforma strutturale dei percorsi formativi post-diploma e professionalizzanti, unita al reskilling, cioè alla riqualificazione continua della forza lavoro attiva. La priorità assoluta deve essere l'inserimento nei programmi didattici — non solo tecnici, ma economici, umanistici e manageriali — di moduli dedicati all'integrazione dei sistemi di intelligenza artificiale nei processi di lavoro. Dobbiamo formare una generazione di professionisti che non subisca la tecnologia e che non si limiti a consumarla, ma che sappia orchestrarla per governare il cambiamento. La vera sfida politica non sarà proteggere i vecchi posti di lavoro dall'AI, ma formare le persone affinché siano in grado di ricoprire quelli nuovi.

Intervistatore: La ringrazio moltissimo. È stato un contributo prezioso per la tesi, soprattutto per il punto di vista di chi si occupa di formazione manageriale e di startup.

Intervistato: Grazie a lei, in bocca al lupo per il lavoro.

Intervistato 4

Sociologo dell'economia

Chiamata online, 23/06/2026

L'intervista è stata condotta e registrata con il consenso dell'intervistato all'uso del nome, del ruolo e dei contenuti ai fini del presente lavoro.

Intervistatore: Grazie mille. La prima domanda riguarda proprio il fenomeno dell'AI divide visto dal suo punto di vista: questa nuova frattura legata all'intelligenza artificiale, che non riguarda più tanto l'accesso alla tecnologia, quanto il trarne un beneficio effettivo. Dal punto di vista del suo lavoro, è un fenomeno che incontra concretamente? E, in caso, mi può raccontare come lo incontra?

Intervistato: Lo incontro solo come studioso, non come practitioner. Chiarisco subito che non ho da fornire evidenze empiriche sul divide che tu stai indagando, ma sul piano teorico sì. La prima cosa che sottolineerei è la natura di lungo periodo del divide: non è un fenomeno recente, non è il prodotto degli ultimi anni, ha alle spalle almeno ottant'anni. Non so che tipo di data iniziale tu e Flavio usiate per parlare di AI, ma normalmente si fa riferimento a quella famosa Summer School che si tenne negli Stati Uniti nel 1956, in cui per la prima volta fu usato il termine, fra l'altro da parte di uno che era evidentemente più brillante dal punto di vista del marketing dell'idea, e che si inventò: «in fondo questa la possiamo chiamare Artificial Intelligence». Chi erano i protagonisti di quell'incontro? Erano tutti anglosassoni, o inglesi, o inglesi americanizzati che poi insegnavano al MIT. Quindi il divide è originario, nasce già negli anni Cinquanta e si accresce — o comunque non si riduce in nessun modo — ed è legato a fenomeni che sono polarizzanti per definizione: le grandi università americane e la spesa militare. Non dimentichiamo che stiamo parlando di tecnologie che nascono in ambito militare e tali restano, come

il mio libro su Peter Thiel ti avrà ampiamente dimostrato. Mi pare chiaro che i fattori che spingono verso il divide ci sono, sono quelli, e sono probabilmente incolmabili almeno nel periodo ragionevolmente immaginabile. C'è un solo paese che sta tentando il catch up, cioè la Cina, per il momento mantenendo però un relativo lag e soprattutto, mi sembra, scegliendo una strategia diversa rispetto all'impero americano: per il momento non è tanto nella competizione sulla fase inventiva. Se prendiamo la distinzione classica tra invenzione e innovazione — un conto è l'invenzione, un conto è l'innovazione, un conto ancora è la diffusione dell'innovazione — la Cina sceglie per il momento di attestarsi piuttosto sulla fascia dell'innovation che non su quella dell'invention. Sta di fatto che, dopo essere stata un grande imitatore, ormai è anche un grande produttore. Quindi staremo a vedere come questa gara — da cui dipende non solo il futuro delle AI, ma il futuro del pianeta nel suo complesso — andrà a finire; andrà avanti, comunque. Il vero attore debole in tutto questo scenario è l'Europa, siamo noi: non esistiamo, per motivi complessi, storici, politici, economici, militari, aggiungi quello che vuoi, però la nostra nullità è abbastanza evidente. E anche qui ha origini lontane: non ha origini solo nelle AI, ma nel fatto che noi sostanzialmente usiamo solo le tecnologie degli Stati Uniti, in tutti i campi, in tutte le piattaforme. Siamo totalmente dipendenti, non abbiamo sviluppato nessuna capacità autonoma, pur avendo ottime università, un'industria militare potenzialmente in crescita, e pur avendo uno degli attori globali di questo mondo che sta in Olanda — quindi non ci sono soltanto le aziende produttrici, tra virgolette, c'è anche un'azienda decisiva che sta in Olanda perché fa soltanto lei quel particolare chip. Quindi noi non esistiamo dal punto di vista della capacità tecnologica, ma non esistiamo nemmeno dal punto di vista della strategia: non abbiamo nessuna strategia che non sia quella, pur importante, di regolare le AI — ma di questo eventualmente parleremo dopo — non di produrre in proprio AI, magari con

caratteristiche diverse da quelle dei grandi player di cui abbiamo appena parlato.

Intervistatore: Per quanto riguarda le regole e le normative, le riprenderemo dopo, perché ho proprio una domanda sull'AI Act su cui sono curioso di sapere come la pensa. Prima, però, volevo farle una domanda sul suo campo, quello universitario e della ricerca. Nel suo campo, secondo lei chi rischia di rimanere indietro in questa evoluzione così veloce — tralasciando l'Europa che, come diceva giustamente, è un po' il fanalino di coda e a cui arriva tutto in ritardo rispetto alle innovazioni che ci sono in America o in altri contesti? Chi rimarrà indietro, e chi invece può trarre realmente vantaggio da questo utilizzo coscienzioso dell'AI?

Intervistato: Si possono per ora fare soltanto delle stime, delle valutazioni qualitative; non mi pare che possiamo rispondere alla tua domanda con dei dati — Flavio è molto più bravo e quantitativo di me — non c'è ancora uno storico. Dal punto di vista qualitativo, la cosa che colpisce di più è che sono a rischio le fasce alte del nostro mercato del lavoro: questo mi sembra indubbio, c'è una convergenza su questa valutazione da parte di quasi tutti gli studiosi. Per la prima volta non sono a rischio le fasce di qualificazione medio-bassa, ma quelle di qualificazione medio-alta, e questo cambia tutto. Se immagini una situazione in cui i nostri nipoti, parlando di quella macchina lì, diranno «il mio medico» — e magari anche «il mio amico», ma limitiamoci al mio medico — cioè se andiamo verso un utilizzo massiccio in tutte le professioni qualificate, si apre una vera e propria voragine dal punto di vista del mercato del lavoro qualificato, per la prima volta, perché non mi pare ci siano confronti possibili con un evento di questa natura. Proprio questo richiederebbe una risposta europea sul rapporto fra AI, struttura delle professioni, qualificazioni medio-alte e quindi produzione di conoscenza, e ruolo delle istituzioni a partire dall'università. Non vedo nessuna mobilitazione di questo tipo, che invece

ritengo estremamente urgente; per quel poco che posso vedere io — che non sono uno studioso di quelle discipline, ovvio, però ho frequentazione con loro — non mi sembra che ci sia ancora in corso una mobilitazione di questo tipo. A cosa potrebbe puntare? Proprio a profilare diversamente questo che si annuncia come una specie di grande emorragia di lavoro dello spirito, se posso usare questo termine weberiano: ricollocare il «lavoro dello spirito» weberiano nella fase dell'artificial intelligence. È un'impresa a cui devono concorrere non solo le facoltà e le discipline ingegneristiche, ma anche le scienze sociali, anche le nostre scienze. Auspicherei una specie di mobilitazione generale — uso il termine volutamente enfatico — perché è un'emergenza, un'emergenza i cui effetti si produrranno in poco tempo. La cosa riguarda tutte le discipline, in realtà, anche quelle apparentemente più lontane: mi vengono in mente quelle giuridiche, che sono fra le più direttamente toccate. Non solo quelle ingegneristiche e tecnologiche, non solo la medicina e la sanità, ma anche quelle apparentemente più lontane, le umanistiche, le giuridiche. Ripeto, mi pare che lì ci sia un vuoto di iniziativa e di riflessione assai grave. Dentro questa diversa profilazione ci metto anche, visto che magari poi non ci saranno domande su questo, gli aspetti di tipo etico. Siamo noi che potremmo insegnare l'etica a Claude, all'intelligenza artificiale; siamo noi, perché l'abbiamo inventata noi, non l'hanno inventata gli americani. Quindi quando i ragazzi di Anthropic insegnano l'etica a Claude stanno semplicemente usando il nostro sapere — e quando dico nostro intendo europeo. Quindi, se noi usiamo il loro sapere nel momento in cui usiamo le piattaforme, i bot, le chat eccetera, loro usano noi, il nostro sapere, nel campo dell'etica dell'intelligenza artificiale, che sarà il problema più esplosivo. Su questo ti sottolineo la lettura dell'enciclica, lettura indispensabile: non so se l'hai ancora vista, ma se non l'hai vista leggila in fretta. «Magnifica Humanitas» è un pezzo importante di questa riflessione, perché indica esattamente una strada di questo tipo: non di

demonizzare il fenomeno dal punto di vista di un'autorità spirituale come la Chiesa Cattolica, ma proprio di entrarci dentro portando i contenuti etico-filosofici necessari.

Intervistatore: La ringrazio, perché sull'etica sinceramente non ci avevamo ancora pensato, io e il professore: non era stato pianificato, ma sicuramente una digressione su questo sarà dovuta. Mi ha già risposto bene o male sia alla terza sia alla quarta domanda, quindi cercherò di fare un'ultima domanda che riassume i concetti che mancano. Sull'AI Act a livello europeo mi è parso di capire, dal suo punto di vista, che come strumento non sia sufficiente ad aiutare le persone nell'Unione Europea a far entrare sempre di più l'intelligenza artificiale nei processi quotidiani, non demonizzandola ma facendo capire come funziona. La mia domanda allora è: a livello italiano, territoriale, se dovessero interpellarla, quale sarebbe il tema principale, il pilastro fondamentale su cui bisognerebbe regolare in via prioritaria l'intelligenza artificiale in Italia?

Intervistato: L'importanza dell'atto europeo è dimostrata dalle reazioni che ha provocato. Avrai visto, se hai letto bene il libro su Peter Thiel, che l'Europa è in fondo l'anticristo per Peter Thiel. Naturalmente lui per anticristo intende altra cosa rispetto a quello che intendiamo noi, perché noi filosofi europei, che abbiamo letto bene l'Apocalisse, sappiamo che non è questo l'anticristo: semmai l'Europa è il potere che frena l'avvento dell'anticristo, quello che Paolo, nella famosa lettera ai Tessalonicesi che cito nel libro, chiamava *katéchon*, che in greco vuol dire qualcuno che rallenta, che frena l'avvento dell'iniquità, l'avvento dell'anticristo. Quindi semmai noi siamo questo *katéchon*: noi freniamo, noi stiamo con l'act, e questo dà molto fastidio agli attori principali, agli attori chiave dell'intelligenza artificiale alla Peter Thiel — questo nostro voler frenare e regolare. L'obiettivo, in questo momento, dei poteri decisivi dell'intelligenza artificiale americana è combattere la regolazione

europea; quindi già questo ti dice che l'act ha fatto bene, ha visto giusto. Non è sufficiente, come si diceva, però certe reazioni di Elon Musk la dicono lunga: che un'autorità come quella europea possa multare un'azienda come la sua nel momento in cui viola alcune regole elementari di trasparenza e di accesso... La mia posizione su questo è molto chiara: l'Europa può essere un efficace attore regolativo, ma non deve essere solo questo, deve anche dire la sua, avere una propria voce. E allora vengo al nostro contributo possibile. Tu dici: in Italia cosa potremmo introdurre nel menù? Dove siamo forti noi, in quali ambiti siamo più forti e quindi possiamo dire la nostra da una posizione di forza e non di debolezza? È una bella domanda, perché se guardiamo al sistema produttivo siamo forti soltanto in settori talmente a valle, che utilizzano, consumano e incorporano a valle l'intelligenza artificiale, che il made in Italy dell'intelligenza artificiale potrebbe eventualmente essere usato per essere fra i primi a sperimentare quello che succede in settori in cui siamo leader. Sappiamo che in alcuni settori siamo leader mondiali — naturalmente settori di beni di consumo, di qualità — e lì quello che potrebbe succedere sarebbe interessante. Anche su questo, però, non mi pare, per quanto non abbia un quadro preciso: sarebbe interessante che magari voi faceste qualche riflessione sul fatto che le aziende leader mondiali del made in Italy — e sai che alludo a settori come design, moda, alimentare eccetera — stiano facendo qualcosa in questo campo, stiano riflettendo sul fatto che proprio loro potrebbero dettare degli standard avanzati di tipo applicativo nei loro settori. Sarebbe auspicabile; non mi pare che stia avvenendo, ma lo auspicherei. Non mi sembra che le varie Confindustria stiano muovendosi, almeno da osservatore ormai più che da direttore di aziende. L'altro grande tema che vedo, visto che noi siamo forti nel pensiero filosofico — non solo in quello tecnologico a guida americana, ma in quello delle scienze filosofiche e sociali — è che lì sarebbe interessante approfondire un punto che sta emergendo: i limiti ontologici

dell'intelligenza artificiale. Quali sono i limiti ontologici, che riguardano l'essere in sé della cosa di cui stiamo parlando? Probabilmente riguardano il suo non essere incorporata, il non avere corpo. Naturalmente si potranno studiare forme di sensibilità neurofisiologica, modi per cui l'intelligenza artificiale simuli una corporeità; ma sappiamo — ed entro in un campo che non è il mio, lo dico più da un punto di vista teorico che analitico — che le interazioni fra i processi mentali, neurologici e fisiologici dell'uomo sono talmente complesse e talmente inesplorate che probabilmente nessuna intelligenza artificiale arriverà mai a possedere quel fascio di relazioni fra mente e corpo che è ciò che rende la nostra intelligenza umana tale. Questo è un campo enorme, è un campo del futuro, perché anche qui, leggendo il mio libro su Peter Thiel, avrai visto quanto loro — Thiel è usato lì come la punta più avanzata di un movimento, che è anche un movimento culturale, non solo tecnologico e politico — spingano per esempio nella direzione delle scienze della vita. Non so se hai visto i capitoli finali del mio libro, su tutte le possibilità, attraverso startup e finanziamenti massicci in questo campo, che permettono loro di forzare le tappe verso interventi delle tecnologie...

Intervistato 5

CEO e Advertising specialist di Piano MAKE

Intervista condotta in forma scritta via mail

L'intervista è stata condotta in forma scritta, e non telefonica, con il consenso dell'intervistato all'uso del nome, del ruolo e dei contenuti ai soli fini del presente lavoro.

1. Si parla sempre più di una nuova frattura legata all'intelligenza artificiale, un «AI divide» che non riguarderebbe più tanto l'accesso alla tecnologia, quanto la capacità di trarne beneficio. Dal punto di osservazione del suo lavoro, questo fenomeno è qualcosa che incontra concretamente? Mi può raccontare in che forme lo vede manifestarsi?

L'unica frattura che vedo è tra chi usa e chi non usa le tecnologie AI. Oggi usare questi strumenti anche solo al 5% della loro piena potenzialità scava un baratro con chi li sta rigettando o ignorando. Ma non vedo un problema di adozione, anzi. Ogni ondata di innovazione tecnologica ha impiegato meno tempo della precedente per raggiungere il miliardo di user. Gli strumenti AI hanno immediatezza di adozione senza precedenti. Per quanto il singolo user potrebbe non utilizzarlo in modo avanzato, come collegamenti API e RAG personalizzate, fatico a vedere il contrario, ovvero come le aziende in 2-3 anni possano sopravvivere senza inserirli nei normali flussi di lavoro.

2. L'intelligenza artificiale non sembra colpire tutti allo stesso modo: più che creare un divario nuovo, tende a riprodurre e talvolta ad amplificare diseguaglianze che già esistono. C'è anche un aspetto che appare controintuitivo, e cioè che con l'AI non sono necessariamente i soggetti più deboli o meno istruiti a restare indietro, come accadeva con il vecchio divario digitale, ma le linee di vantaggio e svantaggio si

ridisegnano in modo nuovo. Nel suo campo, chi sono oggi i soggetti che rischiano di restare indietro e chi invece ne esce avvantaggiato?

Il vantaggio è per tutti quei settori dove le barriere di ingresso si abbassano. Ad esempio, se c'è stato un momento in cui per aprire una vetrina su internet dovevi necessariamente rivolgerti a un programmatore skillato, oggi potresti pubblicare il tuo sito web in vibe coding, letteralmente in poche ore. Il rovescio della medaglia si ha per tutta quella fascia di mercato scarsamente specializzata. Tutti quei consulenti e professionisti che non hanno una conoscenza approfondita della propria materia possono essere facilmente sostituiti da strumenti AI.

3. L'Unione Europea ha approvato l'AI Act, che prova a regolare questi sistemi in base al loro livello di rischio. Secondo lei una risposta di questo tipo è sufficiente a ridurre l'AI divide, oppure siamo di fronte a un problema soprattutto sociale ed economico, su cui la sola norma può poco? E se dovesse indicare uno strumento concreto, normativo o di altro tipo, che oggi manca e che servirebbe davvero, quale sarebbe?

Ovviamente l'Europa, come suo solito, tenta di regolamentare un fenomeno globale difficilmente arginabile. Il rischio lo vedo soprattutto sulla privacy e reputazione personale, in quanto adesso chiunque può costruire affermazioni o fatti completamente falsi ma apparentemente veritieri sul conto di un'altra persona. Non ho gli strumenti per trovare una soluzione a questa potenziale deriva, ma sicuramente è un tema da affrontare al più presto.

4. Guardando ai prossimi anni in Italia, qual è la cosa che la preoccupa di più rispetto al modo in cui l'intelligenza artificiale ridistribuirà opportunità e diseguaglianze? E se potesse mettere una sola priorità

sul tavolo di chi deve decidere le politiche pubbliche, quale indicherebbe?

Come detto prima, ogni tecnologia che distrugge le barriere di ingresso in un mercato ne favoriscano la libertà di accesso e la scala sociale. La priorità delle politiche è, a mio avviso, favorire l'accesso a queste tecnologie ma anche regolamentare gli aspetti privacy e reputazione per arginare il rischio di diffamazione o notizie false.

Bibliografia

Ahmmad, M., Shahzad, K., Iqbal, A., & Latif, M. (2025). Trap of social media algorithms: A systematic review of research on filter bubbles, echo chambers, and their impact on youth. *Societies*, 15(11), 301. <https://doi.org/10.3390/soc15110301>

Aizenberg, E., & van den Hoven, J. (2020). Designing for human rights in AI. *Big Data & Society*, 7(2). <https://doi.org/10.1177/2053951720949566>

Bender, E. M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021). On the dangers of stochastic parrots: Can language models be too big? In *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency* (pp. 610-623). Association for Computing Machinery. <https://doi.org/10.1145/3442188.3445922>

Brunetti, I., & Mussida, C. (2025). *AI occupational exposure and wage distribution: The case of Italy* (Inapp Working Paper n. 135). Inapp.

Castells, M. (2007). Communication, power and counter-power in the network society. *International Journal of Communication*, 1, 238-266.

Cheong, B. C. (2024). Transparency and accountability in AI systems: Safeguarding wellbeing in the age of algorithmic decision-making. *Frontiers in Human Dynamics*, 6, 1421273. <https://doi.org/10.3389/fhumd.2024.1421273>

Collins, R. (1992). *Teorie sociologiche*. Il Mulino.

Croteau, D., & Hoynes, W. (2022). *Sociologia generale. Teorie, metodo, concetti* (3^a ed. a cura di F. Antonelli & E. Rossi). McGraw-Hill.

Eder, M., & Sjøvaag, H. (2024). Artificial intelligence and the dawn of an algorithmic divide. *Frontiers in Communication*, 9. <https://doi.org/10.3389/fcomm.2024.1453251>

European Social Survey European Research Infrastructure (ESS ERIC). (2025). *CRONOS-3 Wave 5* (Edition 1.0) [Data set]. Sikt - Norwegian Agency for Shared Services in Education and Research. <https://doi.org/10.21338/cron3w5e01>

Felten, E. W., Raj, M., & Seamans, R. (2021). Occupational, industry, and geographic exposure to artificial intelligence: A novel dataset and its potential uses. *Strategic Management Journal*, 42(12), 2195-2217. <https://doi.org/10.1002/smj.3286>

Ferri, V., Porcelli, R., & Fenoaltea, E. M. (2024). *Lavoro e Intelligenza artificiale in Italia: tra opportunità e rischio di sostituzione* (Inapp Working Paper n. 125). Inapp.

Giannelli, M. (2021). Poteri dell'AGCOM e uso degli algoritmi. Tra innovazioni tecnologiche ed evoluzioni del quadro regolamentare. *Osservatorio sulle fonti*, 2/2021.

Gran, A.-B., Booth, P., & Bucher, T. (2021). To be or not to be algorithm aware: A question of a new digital divide? *Information, Communication & Society*, 24(12), 1779-1796. <https://doi.org/10.1080/1369118X.2020.1736124>

Hargittai, E. (2002). Second-level digital divide: Differences in people's online skills. *First Monday*, 7(4). <https://doi.org/10.5210/fm.v7i4.942>

Hartmann, D., Wang, S. M., Pohlmann, L., & Berendt, B. (2025). A systematic review of echo chamber research: Comparative analysis of conceptualizations, operationalizations, and varying outcomes. *Journal of Computational Social Science*, 8(2). <https://doi.org/10.1007/s42001-025-00381-z>

Helsper, E. J., & van Dijk, J. A. G. M. (2012). A corresponding fields model for the links between social and digital exclusion. *Communication Theory*, 22(4), 403-426.

Kahneman, D. (2017). *Pensieri lenti e veloci*. Mondadori. (Edizione originale: Thinking, fast and slow, Farrar, Straus and Giroux, 2011)

Morozov, E. (2013). *To save everything, click here: The folly of technological solutionism*. PublicAffairs.

Norris, P. (2001). *Digital divide: Civic engagement, information poverty, and the Internet worldwide*. Cambridge University Press.

OECD. (2023). *OECD Employment Outlook 2023: Artificial Intelligence and the Labour Market*. OECD Publishing.

Pariser, E. (2011). *The filter bubble: What the internet is hiding from you*. Penguin Press.

Petrovčič, A., Reisdorf, B. C., Grošelj, D., & Dolničar, V. (2024). Disentangling the role of algorithm awareness and knowledge in digital inequalities: An empirical validation of an explanatory model. *Information, Communication & Society*. <https://doi.org/10.1080/1369118X.2024.2363896>

Pizzinelli, C., Panton, A. J., Mendes Tavares, M., Cazzaniga, M., & Li, L. (2023). *Labor market exposure to AI: Cross-country differences and distributional implications* (IMF Working Paper n. 216). International Monetary Fund.

Shi, W., & Li, J. (2024). New digital divide shaped by algorithm? Evidence from agent-based testing on Douyin's health-related video recommendation. *Communication Research*.

Sunstein, C. R. (2001). *Republic.com*. Princeton University Press.

Teutloff, O., Einsiedler, J., Kässi, O., Braesemann, F., Mishkin, P., & del Rio-Chanona, R. M. (2025). Winners and losers of generative AI: Early evidence of shifts in freelancer demand. *Journal of Economic Behavior & Organization*, 235, 106845. <https://doi.org/10.1016/j.jebo.2024.106845>

Tichenor, P. J., Donohue, G. A., & Olien, C. N. (1970). Mass media flow and differential growth in knowledge. *Public Opinion Quarterly*, 34(2), 159-170. <https://doi.org/10.1086/267786>

van Deursen, A. J. A. M., & Helsper, E. J. (2015). The third-level digital divide: Who benefits most from being online? In L. Robinson, S. R. Cotten, & J. Schulz (Eds.), *Communication and information technologies*

annual (pp. 29-52). Emerald. <https://doi.org/10.1108/S2050-206020150000010002>

van Deursen, A. J. A. M., & van Dijk, J. A. G. M. (2011). Internet skills and the digital divide. *New Media & Society*, 13(6), 893-911. <https://doi.org/10.1177/1461444810386774>

van Dijk, J. A. G. M. (2020). *The digital divide*. Polity Press.

Wang, Y., Boerman, S. C., Kroon, A. C., Möller, J., & de Vreese, C. H. (2025). The artificial intelligence divide: Who is the most vulnerable? *New Media & Society*, 27(7), 3867-3889. <https://doi.org/10.1177/14614448241232345>

Zuboff, S. (2015). Big other: Surveillance capitalism and the prospects of an information civilization. *Journal of Information Technology*, 30(1), 75-89. <https://doi.org/10.1057/jit.2015.5>

Sitografia

AGCOM. (2020). *Le strategie di disinformazione online e la filiera dei contenuti fake (Indagine conoscitiva «Piattaforme digitali e sistema dell'informazione»)*. Autorità per le Garanzie nelle Comunicazioni. <https://www.agcom.it/node/10941>

European Commission. (2023). *Report on the state of the Digital Decade 2023*. <https://digital-strategy.ec.europa.eu/en/library/2023-report-state-digital-decade>

ISTAT. (2025, 17 Gennaio). *Imprese e Ict - Anno 2024* (Comunicato Stampa). Istituto Nazionale di Statistica. <https://www.istat.it/comunicato-stampa/imprese-e-ict-anno-2024/>

ISTAT. (2025, 15 Dicembre). *Imprese e Ict - Anno 2025* (Comunicato Stampa). Istituto Nazionale di Statistica. <https://www.istat.it/comunicato-stampa/imprese-e-ict-anno-2025/>

Pariser, E. (2011). *Beware online «filter bubbles»* [Video]. TED. https://www.ted.com/talks/eli_pariser_beware_online_filter_bubbles

van Dijk, J. A. G. M. (2020). *Closing the digital divide: The role of digital technologies on social development, well-being of all and the approach of the Covid-19 pandemic*. United Nations Department of Economic and Social Affairs. <https://www.un.org/development/desa/dspd/wp->

[content/uploads/sites/22/2020/07/Closing-the-Digital-Divide-by-Jan-A.G.M-van-Dijk-.pdf](#)