



GIULIA FERRARI
CHIARA FUSTINONI

Verso una nuova progettazione:
storia e metamorfosi del design
nell'era dell'intelligenza artificiale

© 2026 Politecnico di Torino
Tutti i diritti riservati.

Nessuna parte di questa pubblicazione può essere riprodotta, archiviata in sistemi di recupero dati o trasmessa in qualsiasi forma o con qualsiasi mezzo, elettronico, meccanico, fotocopiatura, registrazione o altro, senza il preventivo consenso scritto dell'autore, salvo i casi previsti dalla legge.

Corso di Laurea in Design e Comunicazione
Sessione di Laurea Luglio 2026

Autori: Giulia Ferrari. Chiara Fustinoni
Relatore: Per Paolo Peruccio
Corelatore: Gianluca Grigatti

Titolo: LIMINE - Verso una nuova progettazione: storia e metamorfosi
del design nell'era dell'intelligenza artificiale

Progetto grafico e impaginazione: Giulia Ferrari. Chiara Fustinoni

Stampato a Torino da Tipografia Ideal Snc.



**GIULIA FERRARI
CHIARA FUSTINONI**

Verso una nuova progettazione:
storia e metamorfosi del design
nell'era dell'intelligenza artificiale



**Politecnico
di Torino**

indice

abstract

introduzione

prima parte

capitolo 1: storia

1.1 L'evoluzione dell'Intelligenza Artificiale attraverso le sue fasi storiche

1.1.1 Boom e winter: le stagioni fondanti dell'AI

Primo ciclo: Origini e nascita della dicotomia (1940–1960)

Secondo ciclo: Crisi, rinascita e ridefinizione dei paradigmi (1970–1990)

Terzo ciclo: Dalla convergenza pragmatica al dominio dei dati: machine learning e svolta del deep learning (1990–2010)

1.1.2 Oltre i cicli: l'accelerazione contemporanea

Dall'AI statistica all'AI agentica: continuità e trasformazioni (2010–oggi)

Modelli multimodali (2023–2025)

La rivoluzione del ragionamento: Sistema 1 e 2 (2024–2025)

Intelligenza artificiale agentica basata su modelli LLM (2025)

1.2 Storia dell'intelligenza artificiale in applicazioni creative e progettuali

1.2.1 Il designer controlla la macchina

Cibernetica: Sistemi adattivi e reattività (anni '50)

Interfaccia: La nascita del medium grafico (anni '60)

Codifica: AARON e la generazione autonoma (anni '70)

1.2.2 Il designer programma la macchina

Automazione: il CAD come standard tecnico (anni '80)

Simulazione: Ottimizzazione e reti neurali (anni '90)

Parametrico: Il progetto come codice algoritmico (anni 2000)

1.2.3 Il designer dialoga con la macchina

Data-driven: L'AI come assistente analitico (2010-2013)

Generativo: La macchina come co-designer (2014-2015)

AI art: L'emergenza del nuovo autore (2016-2019)

Workflow: L'era della partnership creativa (2020-oggi)

seconda parte

CAPITOLO 2 - PANORAMA INTERNAZIONALE

2.1 Il panorama internazionale dell'intelligenza artificiale nel design

2.1.1 Evoluzione del mercato globale dell'AI nel design

2.1.2 Le forze trainanti del mercato

Smart Manufacturing e Industria 4.0

Accessibilità del Cloud Computing e HPC scalabile

2.1.3 Analisi di segmentazione del mercato

Segmentazione del mercato per tipologie: software e servizi

Segmentazione del mercato per settore di utilizzo finale

2.2 Key player globali e strategie nell'AI per il design

2.2.1 USA: L'AI come potenziamento tecnico

2.2.2 Europa: governance e verità scientifica

2.2.3 Cina: l'AI come infrastruttura totale

2.2.4 Corea del Sud: integrazione verticale

2.3 Il panorama accademico e la formazione del designer nel 2026

2.3.1 I poli internazionali della formazione: USA, Cina e Regno Unito

USA: ricerca e sperimentazione didattica

Cina: framework adattivo e artigianato aumentato

Regno Unito: design come leadership

2.3.2 Altre esperienze nel mondo

Europa continentale: tra ambiente costruito, Hybrid Intelligence e AI responsabile

Corea del Sud: un laboratorio nazionale per l'AI

America Latina: infrastrutture emergenti e prime sperimentazioni

2.3.3 Il contesto formativo italiano

Il dibattito italiano: regolamentare, non proibire

Il Politecnico di Torino: un hub nazionale per l'AI nel design

2.3.4 Cinque tratti convergenti dal panorama accademico

CAPITOLO 3 - LEGISLAZIONE

3.1 Pluralità, tensioni e pilastri: le coordinate del problema

3.1.1 Governance dell'AI e pluralità normativa

3.1.2 Le tensioni tra cooperazione globale e interessi privati

3.1.3 I tre pilastri della governance globale

3.2 L'Unione Europea e il modello regolatorio globale: ambizioni, compromessi e criticità

3.2.1 Il modello regolatorio europeo: sfide e prospettive

3.2.2 La classificazione del rischio come modello regolatorio

3.2.3 AI generativa e rischi sistemici: nuove frontiere normative

3.2.4 Governance, implementazione e l'impatto sulle PMI: un percorso irto di ostacoli

3.2.5 Diritti umani e AI: il ruolo del Consiglio d'Europa

3.2.6 Oltre l'AI Act: l'Europa può avere una sovranità normativa senza quella industriale?

Intervista 1 - Anne Asensio (designer e dirigente aziendale)

3.3 Nord America - Tra deregolamentazione, leadership tecnologica e nuovi orizzonti normativi

3.3.1 Stati Uniti d'America: un pendolo tra sicurezza nazionale e competitività economica

3.3.2 Canada: L'Artificial Intelligence and Data Act (AIDA) e la ricerca di un equilibrio prudente

3.4 America Latina: un mosaico di iniziative nazionali e la ricerca di un modello proprio

3.4.1 Brasile: protezione dei diritti come priorità

3.4.2 Cile: strategia nazionale e proposte legislative per un'AI al servizio

pubblico

3.4.3 Argentina: linee guida etiche e un dibattito aperto sull'ecosistema locale

3.5 L'Asia-Pacifico - Laboratorio di governance tra controllo statale, soft-law e ambizioni geopolitiche

3.5.1 Cina: la regolamentazione iterativa e l'imperativo del controllo statale

3.5.2 Giappone: l'agilità della soft-law e la promozione dell'innovazione con cautela

3.5.3 Corea del Sud: innovazione e regolamentazione in equilibrio dinamico

3.5.4 India: il Digital India Act e la ricerca di sovranità tecnologica e inclusione

3.6 Africa, Medio Oriente e la ricerca di una governance globale inclusiva: sfide e contraddizioni

3.6.1 Africa: in bilico tra sviluppo inclusivo e dipendenza tecnologica

3.6.2 Medio Oriente e AI: strategie di modernizzazione e controllo

3.6.2.1 Emirati Arabi Uniti: un modello di governance proattiva con ambiguità sui diritti

3.6.2.2 Arabia Saudita: Vision 2030 e l'AI come motore di trasformazione sotto egida statale

3.7 Organismi internazionali e la frammentazione della governance globale: un consenso elusivo

3.7.1 Nazioni Unite: Il Global Digital Compact e le aspirazioni di un consenso fragile

3.7.2 OCSE: i principi e la sfida dell'interoperabilità in un mondo multipolare

3.7.3 G7 e G20: tentativi di coordinamento tra grandi potenze e la prevalenza degli interessi nazionali

3.8 I limiti della regolazione: governare l'incertezza

CAPITOLO 4 - TECNOETICA

4.1 Limiti epistemici e opacità dell'intelligenza artificiale

4.1.1 I limiti cognitivi dell'intelligenza artificiale

Che cos'è l'intelligenza?

4.1.2 Black box e opacità progettuale

Il principio di explicability: punto di vista etico

Explainable AI (XAI): punto di vista tecnico

Opacità come problema progettuale

- 4.1.3 Allucinazioni
 - Scheming
 - Implicazioni per il design
- 4.1.4 Bias, discriminazione e giustizia algoritmica
 - Origine strutturale del bias
 - Il bias nei sistemi di AI generativa e nel design
 - Riconoscere il bias: una competenza del designer

4.2. La figura del designer nell'era dell'AI

- 4.2.1 L'agency del designer
- 4.2.2 Delega e automazione del processo progettuale
 - La delega come problema cognitivo
 - La delega come problema epistemico
 - La delega come problema di bias e disuguaglianza
 - Anche l'AI usa noi: il paradigma Human-Aided AI
- 4.2.3 Il capitale semantico
- 4.2.4 Alienazione e sostituzione
- 4.2.5 Responsibility gap e negligenza professionale
 - Quattro tipologie di lacuna nella responsabilità
- 4.2.6 Meaningful Human Control
 - Responsabilità culturale

Intervista 2 - Vera Tripodi (filosofa morale)

4.3. L'irriducibilità dell'umano

- 4.3.1 L'infrastruttura del potere algoritmico e la crisi della fiducia
- 4.3.2 La maschera dell'etica tecnologica: il fenomeno dell'ethics-washing
- 4.3.3 La funzione autore e la crisi dell'autorialità
- 4.3.4 Estetica frammentata e l'illusione della creatività post-antropocentrica
- 4.3.5 Il limite biologico e l'immaginazione morale

Intervista 3 - Maurizio Ferraris (filosofo teoretico)

4.4 Decidere è ancora un atto umano?

terza parte

CAPITOLO 5 - PROGETTO

- 5.1 Prima dell'algoritmo: cosa significava progettare
 - 5.1.1 L'evoluzione degli strumenti di progettazione
 - Il disegno come pensiero
 - L'avvento del Computer-Aided Design
 - Una transizione stratificata
 - 5.1.2 L'uomo al centro: il paradigma dello Human-Centered Design
 - L'essenza del design

- La nascita dello Human-Centered Design
- 5.1.3 Il Design Thinking come framework metodologico
- 5.1.4 Il Double Diamond: origine e architettura del modello
- 5.1.5 Il passaggio dallo Human-Centered Design allo Human-Centered Artificial Intelligence
 - I limiti del paradigma pre-algoritmico
 - Il passaggio allo Human-Centered Artificial Intelligence
- 5.1.6 La complessità come condizione permanente del progetto
 - La ricerca progettuale e la risoluzione dei problemi
 - AI maturity model
 - Il paradigma pre-algoritmico come punto di partenza
- 5.2 Dopo il modello generativo: cosa significa decidere
 - 5.2.1 Nuovi strumenti di intelligenza artificiale nel processo progettuale
 - L'AI nelle fasi del processo: una mappatura operativa
 - 5.2.2 Problemi emergenti dall'utilizzo dell'AI nelle fasi del processo
 - Difficoltà tecniche e progettuali
 - Il sovraccarico decisionale
 - Problemi comunicativi
 - 5.2.3 Lo Stingray Model
 - 5.2.4 L'AI non semplifica il progetto, ma lo modifica
 - Il plausibility gap: quando l'immagine precede la sostanza
 - 5.2.5 Il cambio di paradigma nell'interazione con le macchine: da command-based a intent-based
 - La nuova relazione con il sistema
 - Augment versus automate
 - 5.2.6 Il cambio di ruolo del designer: da creator a director
 - Dalla produzione alla curatela
 - Task affidati all'AI e decisioni strategiche riservate al designer
 - 5.2.7 Tensioni irrisolte
 - Il deskilling
 - I bias non intenzionali
 - L'omologazione estetica
 - L'atrofizzazione cognitiva
 - La disuguaglianza nell'accesso

Intervista 4 - Massimo Banzi (designer e imprenditore)

- 5.3 Nuovo sistema di competenze
 - 5.3.1 La formula del designer AI-native
 - Dai nuovi strumenti alla nuova identità professionale
 - 5.3.2 Le competenze umane che faranno la differenza
 - Le previsioni sulle competenze fondamentali nel 2030
 - Le domande che solo un umano può farsi
 - Il pensiero sistemico
 - 5.3.3 Nuove competenze da acquisire
 - Competenze operative
 - Competenze relazionali e critiche
 - Competenze cognitive e creative

Intervista 5 - Don Norman (ingegnere informatico)

- 5.3.4 La biforcazione delle competenze: due profili emergenti
 - Il designer come direttore creativo algoritmico
 - Il designer come garante della qualità e dell'etica
- 5.3.5 Da dove inizia il cambiamento?

CAPITOLO 6 - FORMAZIONE

- 6.1 Un sistema in ritardo: la formazione che non tiene il passo
 - 6.1.1 La differenza di velocità
 - 6.1.2 L'impreparazione istituzionale e il modello formativo obsoleto
- 6.2 Lo studente di design davanti all'AI: la dimensione psicologica
 - 6.2.1 I determinanti dell'adozione: tecnofobia, piacere percepito e conoscenza soggettiva
 - 6.2.2 La catena virtuosa: autoefficacia, riduzione dell'ansia, cognizione creativa

Intervista 6 - Gaetano Di Tondo (dirigente aziendale)

- 6.3 Principi per una nuova pedagogia del design
 - 6.3.1 Capire prima di usare: il sistema dietro lo strumento
 - 6.3.2 L'etica come parametro progettuale costante
 - 6.3.3 Costruire il giudizio prima di introdurre lo strumento
 - 6.3.4 L'apprendimento attraverso la pratica collaborativa
 - 6.3.5 Ripensare la valutazione: l'importanza del processo
 - 6.3.6 L'interdisciplinarietà come architettura del curriculum
 - 6.3.7 L'indole personale al centro della formazione

6.4 La rivincita delle humanities

Intervista 7 - Guido Saracco (già rettore del Politecnico di Torino)

- 6.5 Il docente di design nell'era generativa: una professione da reinventare?
 - 6.5.1 La condizione attuale: volontà senza strumenti
 - 6.5.2 Dall'erogatore di contenuti al regista della collaborazione
 - 6.5.3 La formazione dei formatori come pre-condizione
 - 6.5.4 ChatGPT nella progettazione didattica: potenzialità e limiti
- 6.6 L'università come luogo in cui si progetta il futuro
 - 6.6.1 La pedagogia semantica: formare creatori di significato

Conclusione

Bibliografia e Sitografia

Abstract

L'integrazione dell'Intelligenza Artificiale generativa nel processo progettuale segna una frattura epistemologica nella disciplina del design. Se il design moderno si è fondato su un approccio human-centered, oggi deve confrontarsi con modelli generativi che partecipano attivamente alla produzione di idee, immagini e decisioni. La nozione tradizionale di autore, storicamente legata all'intenzionalità individuale, entra in crisi, lasciando spazio a una dinamica di progetto in cui la capacità generativa algoritmica coesiste con la facoltà umana di interpretare, selezionare e attribuire senso.

Questa tesi indaga il fenomeno su tre livelli interconnessi. La prima parte ricostruisce le tappe storiche che hanno portato all'AI generativa, analizzando l'evoluzione dell'intelligenza artificiale e, in parallelo, le sue applicazioni creative e progettuali. La seconda parte analizza i diversi approcci internazionali, con un focus su Stati Uniti, Unione Europea e Cina, rivolgendo particolare attenzione alle legislazioni e ai rischi legati all'adozione incontrollata della tecnologia. Esplora poi le implicazioni tecno-etiche della nuova condizione progettuale, evidenziando come l'AI influenzi la distribuzione del potere e la governance nel design. La terza parte torna al nodo umano, interrogando cosa significhi oggi progettare, decidere e apprendere in un contesto in cui i modelli generativi sono co-autori dei processi creativi.

La ricerca sfocia in una riflessione sulla formazione del designer, che richiede un ripensamento di curricula, metodi e competenze, per preparare figure capaci di operare consapevolmente in un ecosistema progettuale sempre più algoritmico. La domanda non è se integrare l'AI, ormai ampiamente presente nella pratica, ma come trasformarla in uno strumento di consapevolezza e responsabilità. In questo modo, il designer conserva autonomia, pensiero critico e capacità di guidare il progetto con senso e valore, competenze fondamentali per crescere in una disciplina in continua trasformazione.

Introduzione

Da dove partiamo

Chi ha intrapreso un percorso di studi in design negli ultimi anni si è formato, nella maggior parte dei casi, in un sistema didattico in cui l'intelligenza artificiale generativa non era ancora, se non marginalmente, parte del curriculum. Si trova oggi a entrare in un contesto professionale in cui l'intelligenza artificiale è invece diffusa capillarmente, e in cui si presume che il designer ne possieda già un uso consapevole, pur non avendo ricevuto una formazione specifica in tal senso. Questo divario tra la velocità dell'innovazione tecnologica e la lentezza dei sistemi educativi non è un disagio isolato, ma una condizione strutturale della disciplina, ed è da questa condizione che nasce la domanda all'origine della ricerca.

Di fronte a questo scarto, sono state le parole dello storico del design Vanni Pasca, pronunciate al primo convegno internazionale di studi storici sul design dedicato al tema Design: storia e storiografia, a indirizzare la domanda di ricerca. In tale occasione, Pasca ha sostenuto che "Piuttosto che interrogarsi su cosa sia, il design dovrebbe avviare un lavoro di ricostruzione dei modi e delle norme con cui, nelle varie fasi storiche, il concetto di design si è rappresentato, tenendo conto di come esso implichi anche il definirsi della figura del designer e del suo ruolo."





1.1 L'evoluzione dell'Intelligenza Artificiale attraverso le sue fasi storiche

¹ M. Pasquinelli, *The Eye of the Master: A social history of artificial intelligence*, Verso Books, 2023.

² D. A. Grier, *When computers were human*, Princeton University Press, 2005.

³ Doron Swade, *The Cogwheel Brain: Charles Babbage and the Quest to Build the First Computer*, Little, Brown and Company, Londra, 2000.

⁴ A. Hyman, *Charles Babbage: Pioneer of the computer*, Oxford University Press, 1982.

Non molti sanno che l'idea del computer automatico, nel senso contemporaneo del termine, non nasce tanto dal tradizionale immaginario di automi pensanti, quanto piuttosto dall'esigenza concreta di meccanizzare il lavoro mentale degli impiegati¹.

Questa origine pragmatica, spesso mascherata nel corso della storia da narrazioni più affascinanti e quasi alchemiche, affonda le sue radici nell'Inghilterra dell'inizio del XIX secolo. In quel contesto, il termine "computer" non designava una macchina, bensì una figura professionale: impiegati, frequentemente donne, incaricati di eseguire manualmente lunghi e ripetitivi calcoli per istituzioni come il governo, la Società Astronomica o la Marina².

Fu proprio per rispondere all'inefficienza, alla lentezza e all'alta probabilità di errore di questo sistema che Charles Babbage concepì l'idea di sostituire il lavoro umano con una macchina automatizzata³. La sua Macchina Differenziale (1822), oggi riconosciuta come un precursore dei computer moderni, nacque da un'ambizione eminentemente commerciale: automatizzare la produzione di tavole logaritmiche prive di errori, strumenti fondamentali per l'astronomia e per il dominio britannico nei commerci marittimi⁴. Tra le problematiche che stimolarono lo sviluppo del calcolo meccanico vi era, in particolare, la determinazione della longitudine in mare aperto, una questione

cruciale per la navigazione.

Sebbene esistessero già calcolatrici meccaniche, queste erano limitate a operazioni elementari e non erano automatizzate. L'innovazione di Babbage consistette nell'immaginare un dispositivo capace di operare in modo continuo, sfruttando la potenza dei motori a vapore: non più una semplice macchina per calcolare, ma un vero e proprio "motore di calcolo", in grado di industrializzare il processo computazionale³.

È tuttavia riduttivo considerare Babbage come un genio isolato. Un contributo fondamentale fu offerto da Ada Lovelace, figura di straordinaria rilevanza nella storia dell'informatica. Figlia del poeta Lord Byron e della matematica Anne Isabella Milbanke, Lovelace sviluppò un profondo interesse per la notazione algebrica, allora nota come "Analisi", arrivando a definirsi "Analista"⁵. Collaborò attivamente alla progettazione della Macchina Analitica e, soprattutto, elaborò quello che è oggi riconosciuto come il primo programma informatico della storia: una sequenza di istruzioni pensata per essere eseguita da una macchina.

⁵ J. Essinger, *Ada's algorithm: How Lord Byron's daughter Ada Lovelace launched the digital age*, Melville House, 2014.

Sebbene il Motore Analitico non sia mai stato realizzato, il lavoro di Lovelace rappresenta una svolta concettuale decisiva. Come evidenziato dallo storico dell'informatica Doron Swade, Lovelace intuì ciò che lo stesso Babbage non colse pienamente: la possibilità che i numeri non rappresentassero soltanto quantità, ma anche entità simboliche di diversa natura, come lettere o note musicali³. Questa intuizione aprì la strada a una macchina capace non solo di eseguire calcoli numerici, ma di manipolare simboli secondo regole formali. Si tratta del passaggio cruciale dal calcolo alla computazione generalizzata, fondamento teorico dei moderni sistemi informatici⁶.

⁶ J. Copeland, *The Essential Turing*, Oxford University Press, 2004.

È in questo contesto che emergeranno le prime riflessioni sull'intelligenza artificiale, intesa come estensione della computazione simbolica verso forme di elaborazione sempre più sofisticate. La traiettoria che conduce dalle intuizioni di Lovelace alle macchine intelligenti del XX secolo non è quindi una rottura, ma una continuità: il passaggio da una macchina che calcola a una macchina che, almeno in parte, "pensa".

1.1.1 Boom e winter: le stagioni fondanti dell'AI

L'evoluzione dell'intelligenza artificiale non è stata un percorso lineare, ma un susseguirsi di cicli di estremo ottimismo e fasi di brusco rallentamento. Comprendere questa alternanza tra "primavere" e "inverni" è fondamentale per decifrare le radici dei paradigmi che ancora oggi dominano il campo computazionale.

Primo ciclo: Origini e nascita della dicotomia (1940–1960)

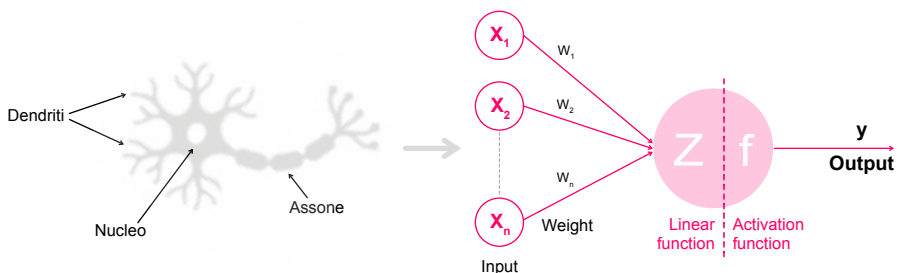
Dal secondo dopoguerra in avanti, quando la convergenza tra matematica, logica, cibernetica e teoria dell'informazione diede origine a una nuova visione computazionale della mente, la storia dell'intelligenza artificiale può essere letta come una tensione tra due paradigmi fondamentali: quello simbolico e quello connessionista⁷. Tali paradigmi riflettono non soltanto approcci tecnici differenti, ma anche concezioni divergenti della mente, della conoscenza e della razionalità⁸.

Il paradigma simbolico emerge in continuità con la tradizione logico-razionalista europea, in particolare con i lavori di Gottfried Wilhelm Leibniz e George Boole, i quali avevano già ipotizzato la possibilità di ridurre il ragionamento umano a manipolazione formale di simboli. Questa linea di pensiero trova una formalizzazione moderna nel lavoro di Alan Turing, la cui nozione di macchina universale introduce l'idea che ogni processo computabile possa essere espresso tramite regole simboliche. Il celebre test di Turing rappresenta inoltre un primo tentativo di definire l'intelligenza in termini comportamentali e linguistici.

⁷ M. A. Boden, *AI: Its nature and future*, Oxford University Press, 2016.

⁸ S. J. Russell, P. Norvig, *Artificial Intelligence: A modern approach (4th ed.)*, Pearson, 2021

Figura 1.1: Formalizzazione matematica e schematica del neurone artificiale a input e output binari sviluppato da Walter Pitts e Warren McCulloch nel 1943. Grafico rielaborato.



⁹ E. R. Kandel et al., *Principles of neural science (5th ed.)*, McGraw-Hill, 2013.

Parallelamente, il paradigma connessionista trae ispirazione da una tradizione biologica e neurofisiologica, influenzata dagli studi sul cervello e sul sistema nervoso⁹. Il lavoro pionieristico di Warren McCulloch e Walter Pitts propone un modello matematico di neurone artificiale, dimostrando che reti di tali unità possono implementare

¹⁰ W. S. McCulloch, W. Pitts, *A logical calculus of the ideas immanent in nervous activity*, The Bulletin of Mathematical Biophysics, 1943.

¹¹ A. Pickering, *The cybernetic brain: Sketches of another future*, University of Chicago Press, 2010.

¹² N. Wiener, *Cybernetics: Or control and communication in the animal and the machine*, MIT Press, 1948

¹³ J. McCarthy, M. L. Minsky, N. Rochester, C. E. Shannon, *A proposal for the Dartmouth summer research project on artificial intelligence*, Stanford University, 1956.

¹⁴ A. Newell, H. A. Simon, *Computer science as empirical inquiry: Symbols and search*, Communications of the ACM, 1976.

¹⁵ H. Gardner, *The mind's new science: A history of the cognitive revolution*, Basic Books, 1985.

funzioni logiche¹⁰. Questo modello gettò le basi per le reti neurali artificiali, aggregando input binari e prendendo decisioni sulla base di tale aggregazione tramite una funzione di attivazione a soglia, con un output binario {0, 1} come risultato.

Un risultato che segna il punto di contatto iniziale tra i due paradigmi, suggerendo che la computazione simbolica possa emergere da strutture distribuite.

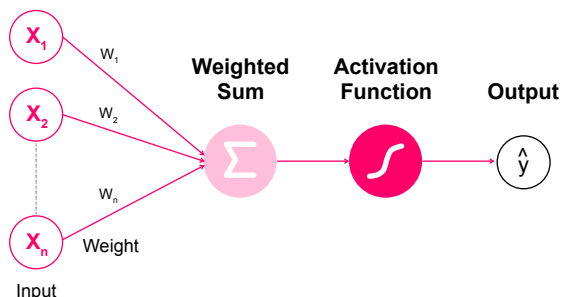
Nel contesto culturale degli anni Quaranta e Cinquanta, segnato dalla Guerra Fredda e dall'espansione della ricerca militare, la cibernetica assume un ruolo centrale come paradigma interdisciplinare¹¹. Norbert Wiener definisce i sistemi intelligenti in termini di feedback, controllo e comunicazione, influenzando sia il connessionismo sia approcci alternativi alla cognizione¹². Tuttavia, mentre la cibernetica enfatizza l'adattamento dinamico, il nascente campo dell'intelligenza artificiale tende rapidamente verso una formalizzazione simbolica.

La conferenza di Dartmouth, dove il termine "Intelligenza Artificiale" fu coniato per la prima volta nel 1956 da John McCarthy, rappresenta un momento fondativo per l'AI come disciplina autonoma, consolidando il paradigma simbolico come approccio dominante. In essa emerge un programma di ricerca ambizioso, basato sull'idea che ogni aspetto dell'intelligenza possa essere descritto formalmente e implementato in una macchina¹³. Questa ipotesi riflette una fiducia significativa nella capacità della logica formale di catturare la complessità dei processi cognitivi, dando origine al paradigma simbolico, che dominerà la ricerca nei decenni successivi.

Il cosiddetto "Physical Symbol System Hypothesis" formalizza questa posizione, sostenendo che un sistema simbolico sufficientemente complesso sia necessario e sufficiente per l'intelligenza generale¹⁴. Tale ipotesi riflette una visione cognitivista della mente, in cui i processi mentali sono analoghi a operazioni computazionali discrete. Questo paradigma si inserisce nel più ampio clima culturale del cognitivismo, che negli anni Sessanta sostituisce il comportamentismo come modello dominante nelle scienze cognitive¹⁵.

Nel frattempo, il connessionismo rimane in una posizione marginale, nonostante il contributo significativo di Frank Rosenblatt, che sviluppa nel 1957 il Perceptron, una rete neurale a singolo strato in grado di apprendere e riconoscere schemi. Il modello Perceptron è un modello com-

Figura 1.2: Schema di funzionamento del Perceptron sviluppato da Frank Rosenblatt nel 1958, evoluzione del primo neurone artificiale operante su input a valori reali. Grafico rielaborato.



Le ambiziose affermazioni di Rosenblatt sulle capacità del Perceptron, tra cui la sua capacità di riconoscere le persone e tradurre il parlato da una lingua all'altra, suscitavano in quel periodo un notevole interesse pubblico nei confronti dell'intelligenza artificiale. Tuttavia, ben presto emerse un limite fondamentale: la regola di apprendimento del Perceptron non era in grado di convergere quando veniva alimentata con dati di addestramento non separabili in modo lineare.

Il perceptron viene inizialmente accolto con entusiasmo, anche grazie al supporto militare e alle aspettative di applicazioni pratiche nel riconoscimento di pattern. Tuttavia, le limitazioni teoriche evidenziate da Marvin Minsky e Seymour Papert nel libro *Perceptrons* contribuiscono a un declino significativo del connessionismo.

Il fallimento della traduzione automatica, evidenziato dal rapporto ALPAC, segna un momento di crisi, mettendo in discussione la validità delle ipotesi di partenza e portando a una riduzione dei finanziamenti¹⁶. Ciò evidenzia già una delle caratteristiche fondamentali dell'evoluzione dell'IA: la tensione tra ambizione teorica e limiti pratici. Le promesse di una intelligenza artificiale generale si scontrano con la complessità dei problemi reali, mostrando come la formalizzazione della conoscenza sia un processo molto più difficile di quanto inizialmente previsto.

Questo episodio riflette dunque una dinamica più ampia: il predominio del paradigma simbolico è rafforzato non solo da risultati tecnici, ma anche da fattori istituzionali, culturali ed economici¹⁷. In un'epoca caratterizzata dalla fiducia nella razionalità formale e nella pianificazione centralizzata (tipica della modernità postbellica) l'approccio simbolico appare più compatibile con le aspettative di controllo e prevedibilità.

¹⁶ ALPAC, *Language and machines: Computers in translation and linguistics*, National Academy of Sciences, National Research Council, 1966.

¹⁷ D. Crevier, *AI: The tumultuous history of the search for artificial intelligence*, Basic Books, 1993.

Secondo ciclo: Crisi, rinascita e ridefinizione dei paradigmi (1970-1990)

Nel corso degli anni Settanta, l'ottimismo iniziale che aveva caratterizzato il paradigma simbolico subisce un progressivo ridimensionamento, dando origine a quello che è comunemente definito il primo "AI Winter". Le aspettative generate nei decenni precedenti, in particolare riguardo alla possibilità di costruire sistemi intelligenti generali basati su regole simboliche, si scontrano con confini computazionali, difficoltà di rappresentazione della conoscenza e fragilità nei contesti reali.

Dal punto di vista tecnico, i limiti dei sistemi simbolici diventano sempre più evidenti. La difficoltà di rappresentare la conoscenza in modo completo e coerente, unita alla rigidità delle regole logiche, rende questi sistemi poco adatti a contesti dinamici e incerti. Il problema del "knowledge acquisition bottleneck" emerge come uno degli ostacoli principali, evidenziando la difficoltà di estrarre e formalizzare la conoscenza degli esperti umani¹⁸.

Le critiche mosse da Hubert Dreyfus risultano particolarmente influenti, in quanto evidenziano l'incapacità dei sistemi simbolici di gestire il sapere tacito, contestuale e incarnato che caratterizza l'intelligenza umana. Ispirandosi alla fenomenologia di Martin Heidegger e Maurice Merleau-Ponty, Dreyfus sostiene che la cognizione non sia una mera manipolazione di simboli, ma un processo situato nel corpo e nell'ambiente¹⁹. Questa critica si inserisce in un più ampio cambiamento culturale degli anni Settanta, caratterizzato da una crescente sfiducia nei confronti dei modelli razionalisti e tecnocratici, anche in risposta alle crisi politiche ed economiche del periodo²⁰.

Parallelamente, il rapporto James Lighthill per il governo britannico rappresenta un momento istituzionale decisivo, poiché critica duramente i risultati dell'AI simbolica e contribuisce a una significativa riduzione dei finanziamenti²¹. Questo evento evidenzia come lo sviluppo dell'intelligenza artificiale sia strettamente legato a dinamiche politico-economiche, oltre che scientifiche.

Nonostante questa crisi, il paradigma simbolico non scompare, ma evolve verso forme più applicative, come i sistemi esperti, che raggiungono una certa maturità negli anni Ottanta²². Questi sistemi, basati su basi di conoscenza e motori inferenziali, trovano applicazione in ambiti industriali e medici, contribuendo a una temporanea rinascita dell'AI simbolica. Tuttavia, la loro dipendenza da co-

¹⁸ M. Polanyi, *The tacit dimension*, Doubleday & Co., 1966.

¹⁹ H. L. Dreyfus, *What computers still can't do: A critique of artificial reason*, MIT Press, 1992.

²⁰ D. Harvey, *A brief history of neoliberalism*, Oxford University Press, 2005.

²¹ J. Lighthill, *Artificial intelligence: A general survey*, Science Research Council, 1973.

²² E. A. Feigenbaum, *Expert systems in the 1980s*, State of the Art Report on Machine Intelligence, Pergamon Infotech, 1981.

noscenze esplicite e la difficoltà di scalabilità limitano il loro impatto a lungo termine.

Nel frattempo, il connessionismo, inizialmente marginalizzato, inizia una fase di rinascita grazie a sviluppi teorici e tecnologici cruciali. Il momento chiave è rappresentato dalla reintroduzione e formalizzazione dell'algoritmo di retropropagazione dell'errore (backpropagation), resa popolare dal lavoro di David Rumelhart, Geoffrey Hinton e Ronald Williams. Questo algoritmo consente l'addestramento efficiente di reti neurali multilivello, superando alcune delle limitazioni del perceptron evidenziate negli anni Sessanta.

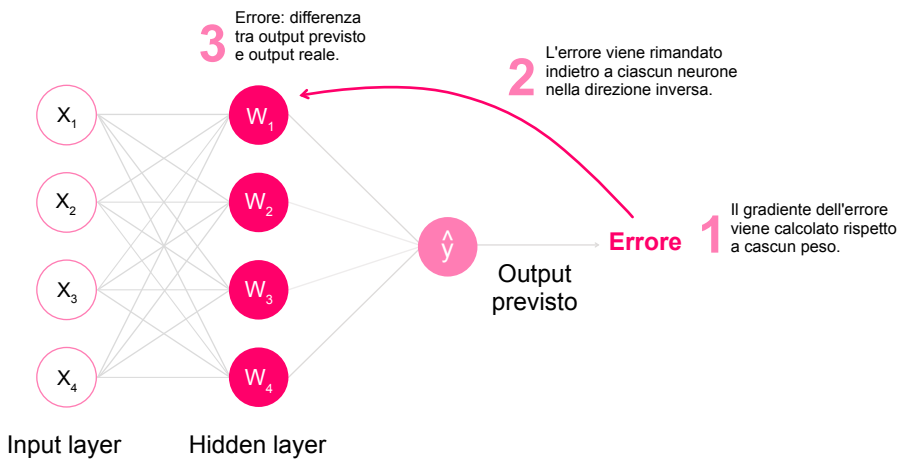
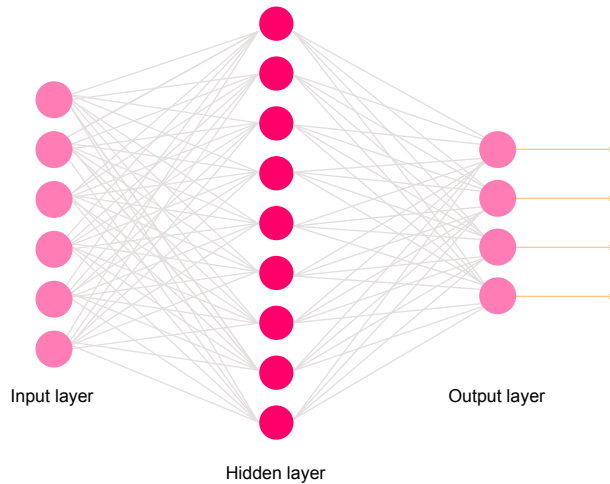


Figura 1.3: Rappresentazione schematica del processo di retropropagazione dell'errore (backpropagation) in una rete neurale multistrato (MLP). Grafico rielaborato.

Il Teorema dell'approssimazione universale, formulato da George Cybenko nel 1989, ha fornito una base matematica alle potenzialità delle reti neurali multistrato. Una rete neurale multistrato con un unico strato nascosto è in grado di approssimare qualsiasi funzione continua con qualsiasi grado di precisione desiderato, consentendo la risoluzione di problemi complessi in vari ambiti.

Durante tutti gli anni Ottanta, lo sviluppo della retropropagazione e del Teorema di approssimazione universale (UAT) ha segnato l'inizio di una seconda età dell'oro per le reti neurali. Tuttavia, il settore delle reti neurali ha attraversato una "seconda era buia" tra i primi anni '90 e i primi anni 2000 principalmente a causa dei limiti computazionali (l'addestramento delle reti neurali profonde richiedeva ancora molto tempo e risorse hardware ingenti) e problemi di overfitting e generalizzazione (le prime reti neurali davano buoni risultati sui dati di addestramento, ma scarsi su quelli non visti).

Figura 1.4: Rappresentazione grafica di una rete neurale feedforward dotata di un singolo livello nascosto (hidden layer). Grafico rielaborato.



²³ M. M. Waldrop, *Complexity: The emerging science at the edge of order and chaos*, Simon & Schuster, 1992.

²⁴ W. D. Nordhaus, *Two centuries of productivity growth in computing*, *The Journal of Economic History*, 2007.

²⁵ J. A. Fodor, Z. W. Pylyshyn, *Connectionism and cognitive architecture: A critical analysis*, *Cognition*, 1988.

²⁶ R. Sun, *Robust reasoning: Integrating rule-based and similarity-based reasoning*, *Artificial Intelligence*, 1995.

Dal punto di vista culturale, gli anni Ottanta segnano anche l'emergere di una nuova sensibilità verso sistemi complessi, auto-organizzazione e non linearità, influenzata da discipline come la teoria del caos e la biologia dei sistemi²³. Il connessionismo si inserisce in questo contesto come modello computazionale capace di catturare dinamiche emergenti e adattive, riflettendo una visione della conoscenza come processo dinamico piuttosto che come struttura statica. Inoltre, il crescente sviluppo dell'informatica e l'aumento della potenza computazionale contribuiscono a rendere praticabili approcci precedentemente considerati troppo onerosi²⁴. Questo fattore tecnologico è fondamentale per comprendere perché il connessionismo, pur essendo concettualmente già presente negli anni Quaranta, riesca ad affermarsi solo decenni dopo.

Un altro elemento chiave di questo periodo è il dibattito tra rappresentazioni simboliche e distribuite, che diventa centrale nelle scienze cognitive. Jerry Fodor e Zenon Pylyshyn criticano il connessionismo sostenendo che esso non sia in grado di spiegare la sistematicità e la produttività del linguaggio e del pensiero²⁵. Questo dibattito evidenzia come la contrapposizione tra i due paradigmi non sia soltanto tecnica, ma profondamente filosofica, riguardando la natura stessa della mente e della rappresentazione.

In questo contesto, alcuni ricercatori iniziano a esplorare approcci ibridi, cercando di combinare i punti di forza dei due paradigmi²⁶. Tuttavia, tali tentativi rimangono inizialmente marginali, poiché la comunità scientifica tende a organizzarsi attorno a scuole di pensiero distinte e spesso contrapposte.

Terzo ciclo: Dalla convergenza pragmatica al dominio dei dati: machine learning e svolta del deep learning (1990–2010)

A partire dagli anni Novanta, il campo dell'intelligenza artificiale entra in una fase di trasformazione profonda, caratterizzata da un progressivo spostamento dall'ambizione di costruire sistemi intelligenti generali verso approcci più pragmatici e orientati ai dati²⁷. Questa transizione segna una parziale ricomposizione della tensione tra paradigma simbolico e connessionista, non attraverso una sintesi teorica, ma guidata da esigenze applicative e vincoli computazionali.

Il machine learning emerge come nuovo paradigma dominante, ponendosi come disciplina autonoma che privilegia l'apprendimento dai dati rispetto alla programmazione esplicita di regole²⁸. In questo contesto, sia approcci simbolici (come gli alberi decisionali e i sistemi basati su regole apprese) sia approcci connessionisti (come le reti neurali) coesistono all'interno di un ecosistema metodologico più ampio²⁹. Tuttavia, il focus si sposta progressivamente verso metodi statistici e probabilistici, come le reti bayesiane e i modelli grafici, che introducono una nuova metodologia basata sull'incertezza e sull'inferenza probabilistica³⁰.

Se il connessionismo continua a evolversi, ma rimane inizialmente in una posizione subordinata rispetto ad altri approcci di machine learning, la vittoria di Deep Blue contro Garry Kasparov nel 1997 segna un trionfo dell'approccio simbolico e della potenza computazionale combinata con tecniche di ricerca euristica³¹. Tuttavia, questo successo evidenzia anche i limiti dell'AI simbolica, in quanto il sistema non possiede una comprensione semantica del gioco, ma si basa su esplorazioni combinatorie estremamente efficienti.

Nel frattempo, lo sviluppo di internet e la digitalizzazione globale trasformano radicalmente il contesto in cui l'AI opera. La crescente disponibilità di dati, unita all'aumento della capacità di calcolo, crea le condizioni per una nuova fase evolutiva. Questo cambiamento è profondamente intrecciato con l'emergere del capitalismo digitale e delle grandi piattaforme tecnologiche, come Google, Amazon e Facebook, che basano i loro modelli di business sull'analisi dei dati su larga scala³².

In questo nuovo ecosistema, il valore si sposta dalla rappresentazione esplicita della conoscenza alla capacità di

²⁷ P. Domingos, *A few useful things to know about machine learning*, Communications of the ACM, 2012.

²⁸ T. M. Mitchell, *Machine learning*, McGraw-Hill, 1997.

²⁹ E. Alpaydin, *Introduction to machine learning (2nd ed.)*, MIT Press, 2010.

³⁰ J. Pearl, *Probabilistic reasoning in intelligent systems: Networks of plausible inference*, Morgan Kaufmann, 1988.

³¹ M. Campbell, A. J. Hoane, F. H. Hsu, *Deep Blue*, Artificial Intelligence, 2002.

³² N. Srnicek, *Platform capitalism*, Polity Press, 2017.

³³ A. Halevy, P. Norvig, F. Pereira, *The unreasonable effectiveness of data*, IEEE Intelligent Systems, 2009.

estrarre pattern da grandi quantità di dati³³. Questo cambiamento segna una rottura significativa con il paradigma simbolico classico, che presupponeva una modellazione esplicita del mondo. Al contrario, il machine learning contemporaneo privilegia approcci empirici, spesso opachi, ma altamente performanti.

La vera svolta avviene nei primi anni 2000 con la rinascita del connessionismo sotto forma di deep learning. Il lavoro di Geoffrey Hinton, Yann LeCun e Yoshua Bengio contribuisce a sviluppare architetture neurali profonde capaci di apprendere rappresentazioni gerarchiche dai dati. Questi modelli superano molte delle limitazioni precedenti grazie a innovazioni come le reti convoluzionali, le tecniche di pre-training e l'uso di GPU per l'addestramento.

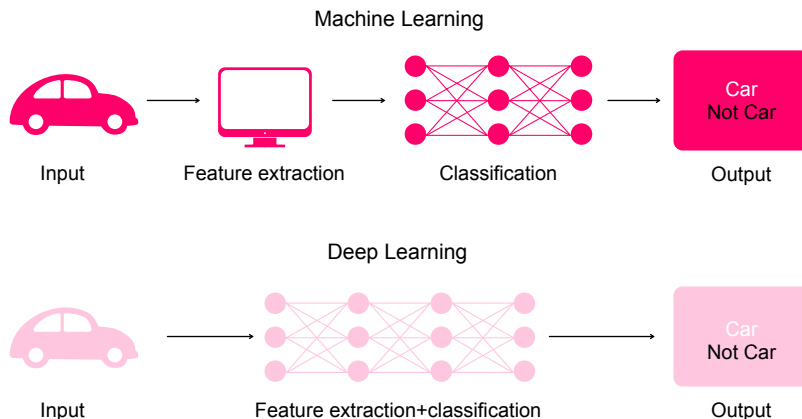


Figura 1.5: Analisi comparativa tra Machine Learning tradizionale e del Deep Learning. Grafico rielaborato.

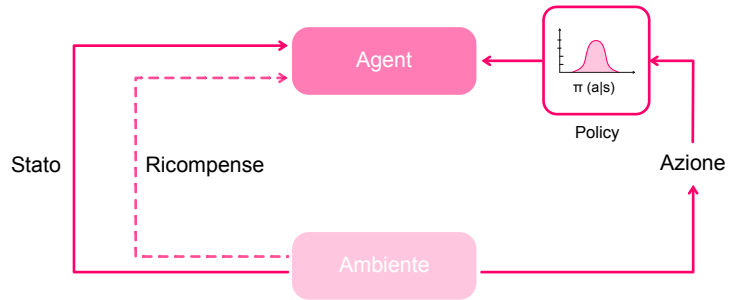
L'apprendimento automatico utilizza sia modelli discriminativi che modelli generativi per prevedere i dati. A partire dalla fine degli anni 2000 e dall'inizio degli anni 2010, i progressi nel deep learning hanno portato a notevoli miglioramenti nella classificazione delle immagini, nel riconoscimento vocale e nell'elaborazione del linguaggio naturale. In questo periodo, le reti neurali venivano solitamente addestrate come modelli discriminativi a causa della relativa difficoltà di addestrare modelli generativi³⁴.

Un punto di svolta simbolico è rappresentato dalla competizione ImageNet Challenge, in cui nel 2012 un modello basato su deep learning, AlexNet, ottiene risultati significativamente superiori rispetto agli approcci tradizionali³⁵. Il successo di AlexNet ha segnato una svolta nello sviluppo delle reti neurali convoluzionali (CNN), aprendo la strada a ulteriori progressi nella classificazione delle immagini e nel rilevamento degli oggetti.

³⁴ Tony Jebara, *Machine Learning: Discriminative and Generative*, Springer Science+Business Media, New York, 2012.

³⁵ A. Krizhevsky, I. Sutskever, G. E. Hinton, *ImageNet classification with deep convolutional neural networks*, *Advances in Neural Information Processing Systems*, NeurIPS, 2012.

Figura 1.6: Diagramma di flusso e architettura fondamentale del Reinforcement Learning (RL). Grafico rielaborato.



L'era moderna del DRL è iniziata nel 2013, quando DeepMind (acquisita da Google nel 2014) ha introdotto le Deep Q-Networks (DQN). Il team, guidato da Volodymyr Mnih, ha presentato un agente in grado di imparare a giocare ai giochi per Atari 2600 direttamente dagli input di pixel grezzi, senza regole specifiche del gioco o caratteristiche definite manualmente. Ciò ha dimostrato che una singola architettura di rete neurale poteva apprendere una vasta gamma di compiti esclusivamente attraverso tentativi ed errori, un tratto distintivo dell'intelligenza generale.

Figura 1.7: Schema architetturale di un sistema di Deep Reinforcement Learning (DRL). Grafico rielaborato.

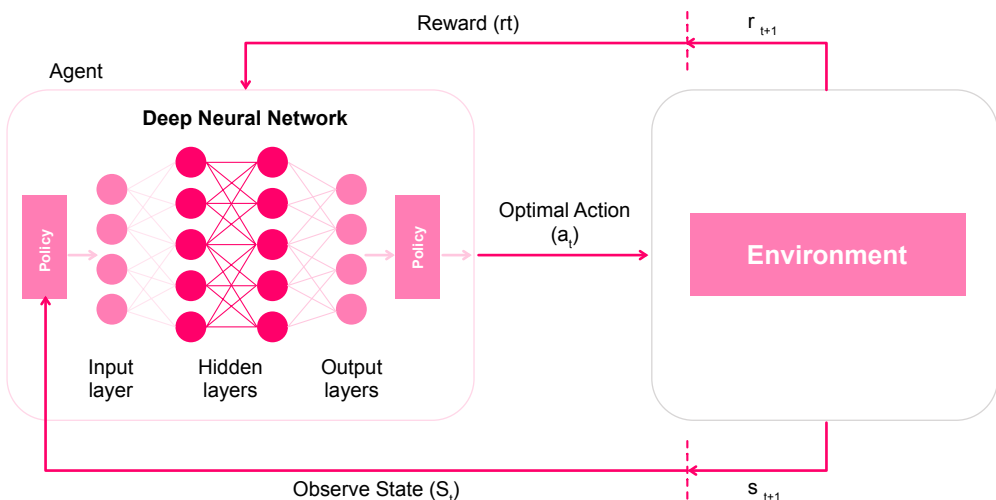
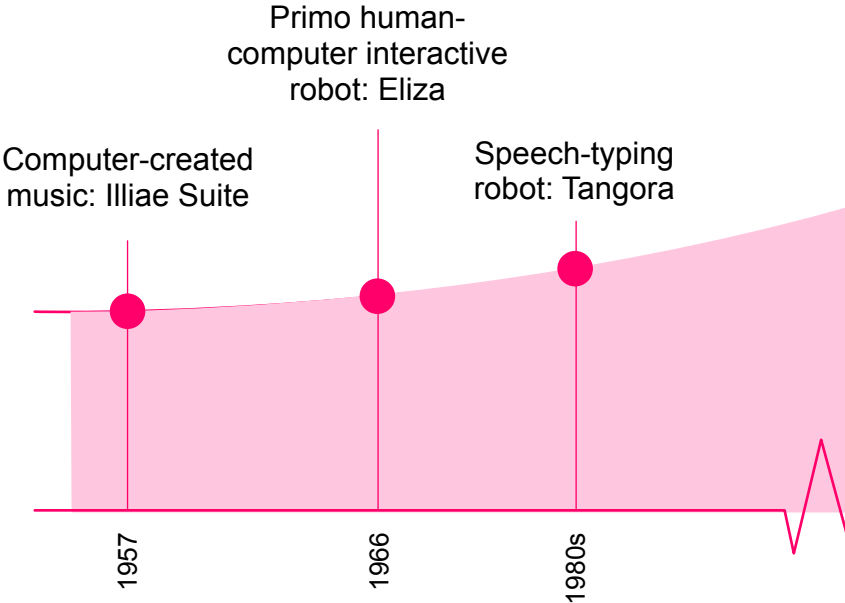
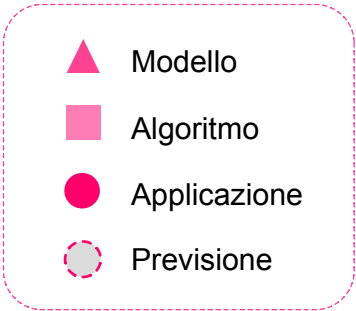
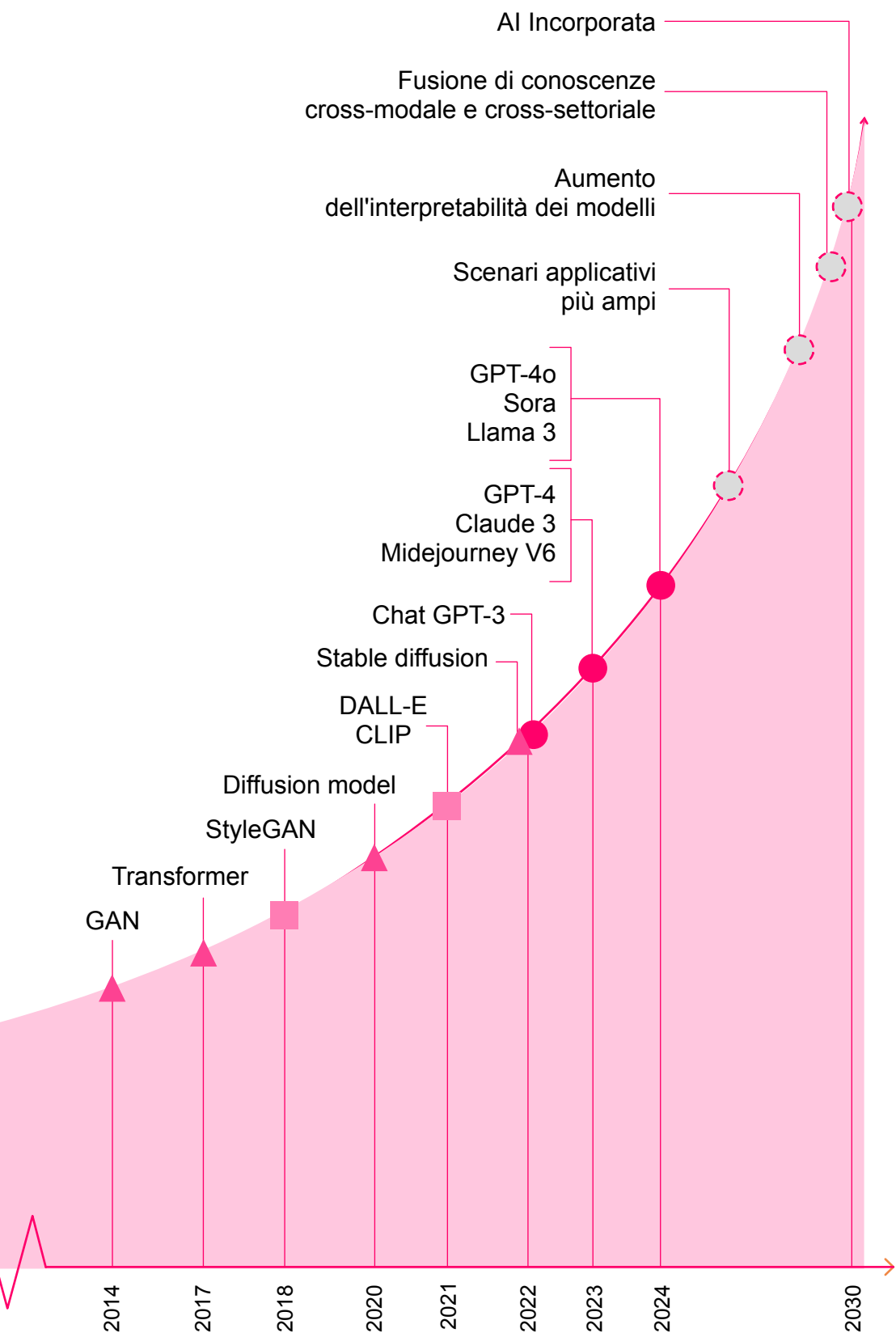


Figura 1.8: Linea del tempo dello sviluppo storico e dell'evoluzione dei modelli di Intelligenza Artificiale dal 1957 alle proiezioni future. Grafico rielaborato.





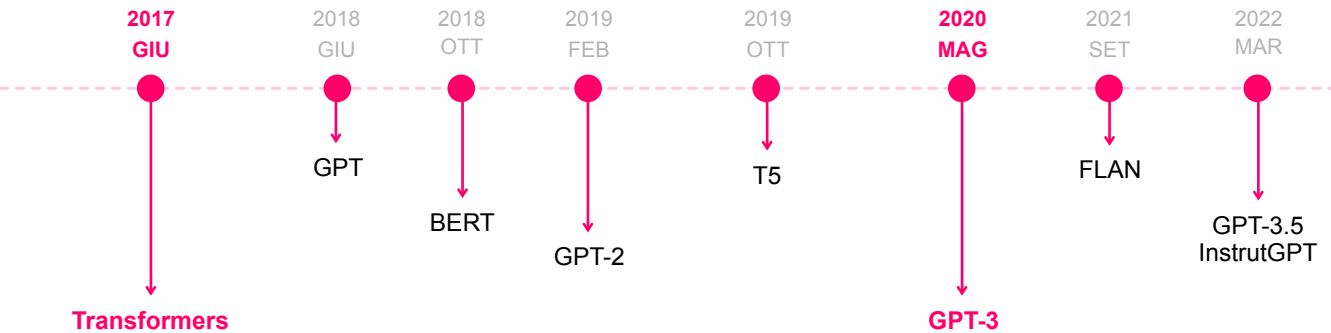
Dopo AlphaGo, DeepMind ha lanciato AlphaZero nel 2017. A differenza del suo predecessore, che era stato addestrato su migliaia di partite giocate da dilettanti e professionisti, AlphaZero ha imparato partendo da zero. Ha disputato milioni di partite contro se stesso, imparando esclusivamente dalle regole del gioco. Dopo sole 24 ore di addestramento, AlphaZero ha sconfitto i programmi campioni del mondo di scacchi (Stockfish), shogi (Elmo) e go (AlphaGo Zero). Questo traguardo ha dimostrato che l'intelligenza artificiale è in grado di assimilare secoli di conoscenza strategica umana (e di superarla) esclusivamente attraverso l'auto-ottimizzazione.

Ma forse l'applicazione più significativa del Deep RL negli anni '20 è il Reinforcement Learning from Human Feedback (RLHF). Questa tecnica utilizza il RL per allineare i modelli linguistici di grandi dimensioni (come GPT-4) alle intenzioni umane. I modelli linguistici di grandi dimensioni (LLM) addestrati esclusivamente sulla previsione del token successivo possono essere dannosi o inaffidabili. L'RLHF addestra un "modello di ricompensa" basato sulle preferenze umane e utilizza la PPO (Proximal Policy Optimization) per mettere a punto l'LLM in modo da massimizzare tale ricompensa. Questa integrazione tra l'apprendimento a rinforzo (RL) e l'IA generativa ha rappresentato il passo fondamentale che ha trasformato i modelli GPT grezzi in assistenti utili come ChatGPT.

Nonostante il predominio del deep learning, il paradigma simbolico non scompare, ma continua a evolversi in ambiti specifici, come la rappresentazione della conoscenza, il ragionamento automatico e i sistemi ibridi³⁶. Questa persistenza suggerisce che la dicotomia tra simbolico e connessionista non sia una semplice alternanza storica, ma una tensione strutturale destinata a riemergere in forme nuove.

³⁶ R. J. Brachman, H. J. Levesque, *Knowledge representation and reasoning*, Elsevier/Morgan Kaufmann, 2004.

Figura 1.9: Cronologia dell'evoluzione dei Grandi Modelli Linguistici (Large Language Models, LLM). Grafico rielaborato.



1.1.2 Oltre i cicli: l'accelerazione contemporanea

Superata la fase delle dicotomie storiche, l'ultimo decennio ha impresso una velocità senza precedenti allo sviluppo tecnologico. Questa accelerazione non rappresenta solo un incremento di potenza, ma una vera e propria mutazione qualitativa nel modo in cui le macchine elaborano l'informazione e interagiscono con il mondo.

Dall'AI statistica all'AI generativa e agentica: continuità e trasformazioni (2010-oggi)

A partire dal secondo decennio del XXI secolo, l'intelligenza artificiale entra in una fase di espansione senza precedenti, caratterizzata dal consolidamento del deep learning come paradigma dominante e dalla sua progressiva estensione verso modelli sempre più generali e multimodali.

Uno degli sviluppi più rilevanti è rappresentato dall'introduzione delle architetture transformer, formalizzate nel lavoro "Attention Is All You Need" di Ashish Vaswani e colleghi³⁷. I transformer superano i limiti delle reti neurali ricorrenti, permettendo una gestione più efficiente delle dipendenze a lungo raggio nei dati sequenziali, e diventano rapidamente la base per una nuova generazione di modelli linguistici³⁸.

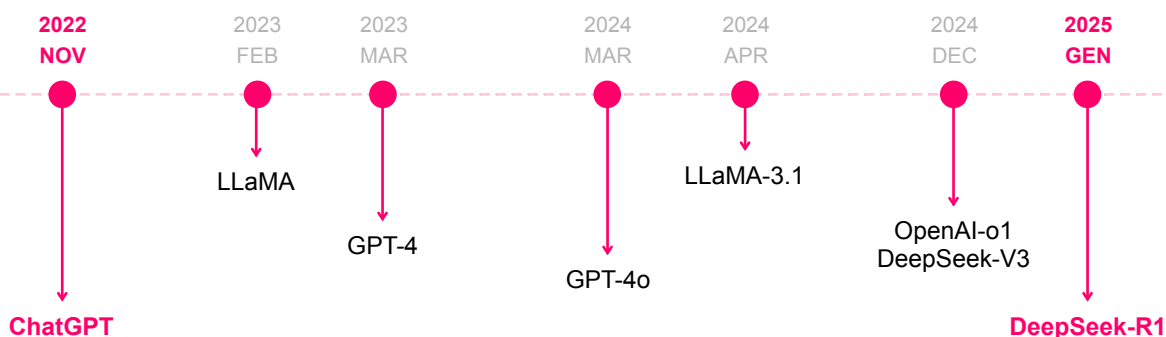
Questa innovazione porta alla nascita dei cosiddetti large language models (LLM), sistemi addestrati su enormi quantità di dati testuali capaci di generare linguaggio naturale, tradurre, riassumere e rispondere a domande con un alto grado di coerenza³⁹.

Modelli sviluppati da aziende come OpenAI e Google dimostrano capacità emergenti che non erano esplicitamente programmate, riaprendo il dibattito sulla natura dell'intelligenza artificiale e sulla possibilità di forme di generaliz-

³⁷ A. Vaswani et al., *Attention is all you need*, Advances in Neural Information Processing Systems, 2017.

³⁸ J. Devlin, M. W. Chang, K. Lee, K. Toutanova, *BERT: Pre-training of deep bidirectional transformers for language understanding*, Proceedings of NAACL-HLT, 2019.

³⁹ T. Brown et al., *Language models are few-shot learners*, Advances in Neural Information Processing Systems, 2020.



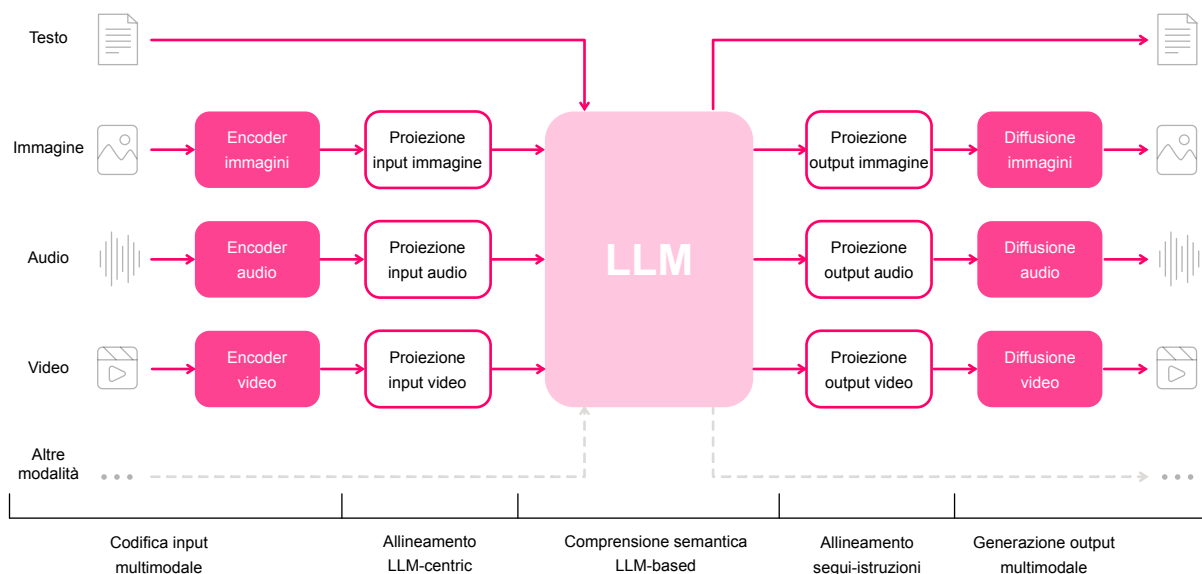
⁴⁰ R. Bommasani et al., *On the opportunities and risks of foundation models*, arXiv, 2021.

zazione più ampie⁴⁰.

L'esempio emblematico è il lancio di Chat GPT che nel novembre 2022 ha segnato una svolta epocale nella storia dell'intelligenza artificiale, portando i modelli linguistici avanzati all'attenzione del grande pubblico. Chat GPT ha raggiunto i 100 milioni di utenti attivi nel giro di due mesi dal suo lancio, diventando così l'applicazione di consumo con la crescita più rapida della storia e suscitando un interesse diffuso per l'intelligenza artificiale in tutti i settori e tra il grande pubblico.

Modelli multimodali (2023–2025)

L'evoluzione dai modelli linguistici basati esclusivamente sul testo ai sistemi di IA multimodali rappresenta un passo fondamentale verso un'intelligenza artificiale più generale. Questi modelli sono in grado di elaborare e generare contenuti attraverso diverse modalità, tra cui testo, immagini, audio e video, consentendo un'interazione uomo-computer più naturale e completa.



La rivoluzione del ragionamento: dal Sistema 1 al Sistema 2 (2024–2025)

Figura 1.10: Diagramma di flusso dell'architettura e della pipeline di elaborazione di un Modello Linguistico Multimodale (Multimodal Large Language Model). Grafico rielaborato.

L'ultima fase pone l'accento sul potenziamento delle capacità di ragionamento, andando oltre il semplice riconoscimento di schemi per arrivare a un pensiero analitico strutturato. Questo cambiamento trae ispirazione dalla teoria dei due processi della psicologia cognitiva, che distingue tra risposte rapide e intuitive (Sistema 1) e ragionamento lento e analitico (Sistema 2).

Figura 1.11:
Rappresentazione schematica dei due processi di pensiero teorizzati dallo psicologo Daniel Kahneman. Grafico rielaborato.

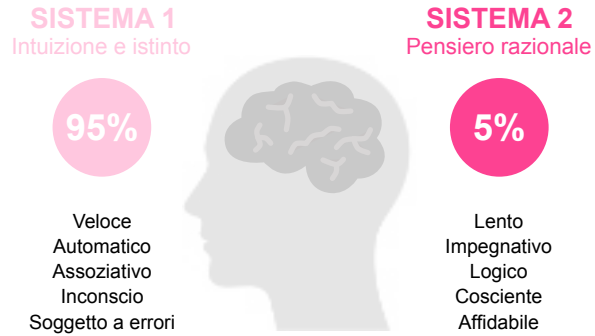


Figura 1.12:
Rappresentazione schematica dei paradigmi dell'IA Simbolica e Connessionista e della loro convergenza nell'IA Neuro-simbolica. Grafico rielaborato.

Il successo dei modelli di ragionamento ha portato a una loro diffusione capillare tra i principali fornitori di IA, la maggior parte dei quali ha implementato sistemi a doppia modalità che combinano risposte rapide con capacità analitiche più approfondite.

Infatti, i LLM sono in grado di manipolare strutture linguistiche e logiche in modo apparentemente simbolico, pur essendo basati su rappresentazioni distribuite⁴¹. Questo fenomeno ha portato alcuni studiosi a parlare di una “riemersione implicita del simbolico” all’interno di architetture connessioniste, suggerendo una possibile convergenza

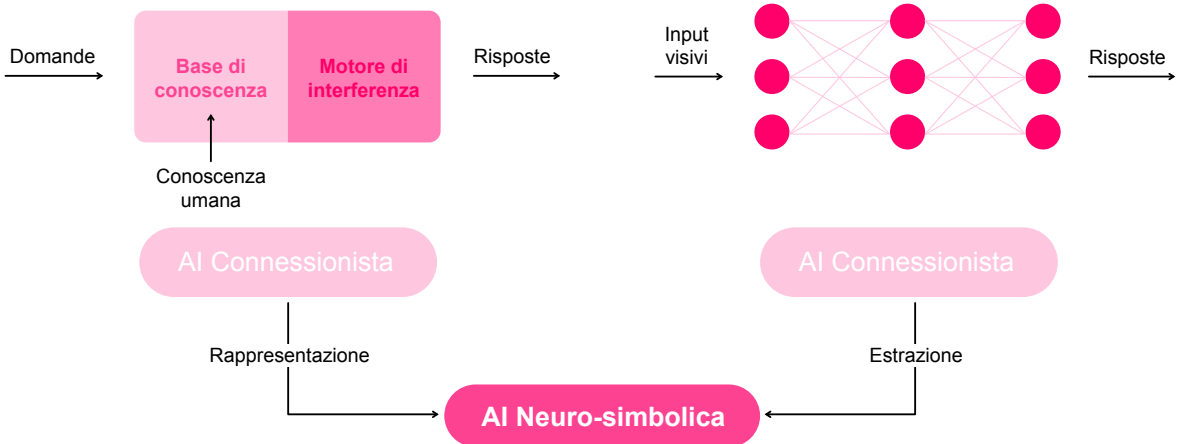
⁴¹ B. M. Lake, T. D. Ullman, J. B. Tenenbaum, S. J. Gershman, *Building machines that learn and think like people*, Behavioral and Brain Sciences, 2017.

AI Simbolica

Operazione analitica
Ragionamento
Sistema 2
Intervento umano
Pochi dati

AI Connessionista

Operazione sintetica
Apprendimento
Sistema 1
Apprendimento dai Big Data
Big Data



⁴² A. D. Garcez et al., *Neural-symbolic computing: An effective methodology for informed machine learning*, Pattern Recognition Letters, 2019.

⁴³ E. M. Bender, T. Gebru, A. McMillan-Major, S. Shmitchell, *On the dangers of stochastic parrots: Can language models be too big?*, Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency, 2021.

⁴⁴ A. D. Garcez, L. C. Lamb, *Neurosymbolic AI: The 3rd wave*, arXiv, 2020.

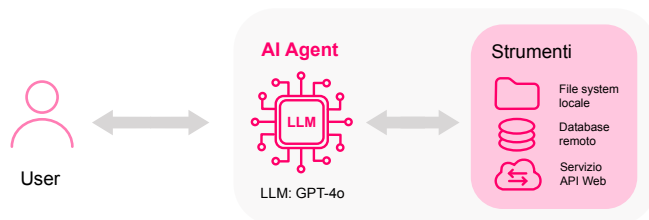
⁴⁵ M. Wooldridge, *An introduction to multiagent systems (2nd ed.)*, John Wiley & Sons, 2009.

Figura 1.13: Diagramma architetturale di un agente basato su Intelligenza Artificiale (AI Agent). Grafico rielaborato.

tra i due paradigmi⁴². Tuttavia, questa convergenza rimane oggetto di dibattito, in quanto tali modelli non possiedono una rappresentazione esplicita e verificabile della conoscenza⁴³.

Parallelamente, si assiste allo sviluppo della cosiddetta neural-symbolic AI, che mira a integrare capacità di apprendimento statistico con meccanismi di ragionamento simbolico⁴⁴. Questo approccio rappresenta un tentativo esplicito di superare la dicotomia storica tra i due paradigmi, recuperando elementi della tradizione simbolica (come la logica e la struttura) senza rinunciare alla flessibilità del connessionismo.

Un ulteriore sviluppo cruciale degli ultimi anni è rappresentato dall'emergere dell'agent AI, ovvero sistemi progettati non soltanto per elaborare informazioni, ma per agire autonomamente in ambienti complessi, perseguendo obiettivi attraverso sequenze di decisioni⁸. Questo concetto, pur avendo radici nella teoria degli agenti degli anni Novanta, assume una nuova rilevanza nel contesto contemporaneo grazie all'integrazione con modelli generativi e capacità di pianificazione⁴⁵.



Intelligenza artificiale agentic basata su modelli LLM (2025)

L'ultima frontiera nello sviluppo dell'IA consiste dunque nella creazione di agenti autonomi in grado di pianificare ed eseguire compiti complessi, nonché di interagire con sistemi esterni utilizzando i modelli di linguaggio di grandi dimensioni (LLM) come base cognitiva. Questi sistemi agentici rappresentano un passaggio dalla risposta passiva alle domande alla risoluzione proattiva dei problemi e all'esecuzione dei compiti.

Tra gli esempi degni di nota figurano Manus, un agente IA generico lanciato nel marzo 2025 dalla startup cinese Monica.im, che trasforma autonomamente i pensieri degli utenti in azioni nell'ambito di attività basate sul web; Gen-spark, uno spazio di lavoro IA all-in-one che nell'aprile 2025

si è evoluto in una piattaforma di super agenti no-code, consentendo la creazione senza soluzione di continuità di strumenti IA come diapositive, documenti e immagini con un unico prompt; e ChatGPT Agent di OpenAI, introdotto nel luglio 2025, che seleziona in modo proattivo da una cassetta degli attrezzi di competenze per completare le attività nel proprio ambiente informatico simulato.

Dal punto di vista storico-epistemologico, l'agentic AI può essere interpretata come un ritorno a una concezione più ampia dell'intelligenza, intesa non solo come capacità di rappresentazione o apprendimento, ma come azione situata e orientata a obiettivi⁴⁶. Questo orientamento richiama alcune intuizioni della cibernetica e delle teorie dell'azione sviluppate nel secondo dopoguerra, suggerendo una continuità profonda tra le diverse fasi evolutive dell'AI¹¹.

Allo stesso tempo, l'integrazione di modelli generativi all'interno di architetture agentiche solleva nuove questioni relative al controllo, all'affidabilità e alla prevedibilità dei sistemi intelligenti⁴⁷. La capacità degli agenti di operare in modo autonomo amplifica infatti le problematiche già presenti nel deep learning, come l'opacità e la difficoltà di interpretazione⁴⁸.

Sul piano culturale, l'ascesa dell'AI generativa e agentic si inserisce in un contesto caratterizzato da una crescente automazione delle attività cognitive e creative, trasformando profondamente il rapporto tra esseri umani e tecnologie⁴⁹. L'intelligenza artificiale non è più soltanto uno strumento di supporto, ma diventa un attore attivo nei processi decisionali e produttivi, ridefinendo i confini tra umano e artificiale.

In questo scenario, la storica dicotomia tra symbolic AI e connectionist AI non scompare, ma si trasforma in una tensione interna ai sistemi stessi. Le architetture contemporanee incorporano elementi di entrambi i paradigmi, dando origine a configurazioni ibride che sfidano le categorie tradizionali⁵⁰. Piuttosto che una sostituzione lineare, la storia dell'AI appare quindi come un processo cumulativo e stratificato, in cui idee e approcci del passato vengono continuamente rielaborati alla luce di nuove condizioni tecnologiche e culturali.

⁴⁶ M. Bratman, *Intention, plans, and practical reason*, Harvard University Press, 1987.

⁴⁷ D. Amodei et al., *Concrete problems in AI safety*, arXiv, 2016.

⁴⁸ F. Doshi-Velez, B. Kim, *Towards a rigorous science of interpretable machine learning*, arXiv preprint, 2017.

⁴⁹ E. Brynjolfsson, A. McAfee, *The second machine age: Work, progress, and prosperity in a time of brilliant technologies*, W. W. Norton & Company, 2014.

⁵⁰ G. Marcus, *The next decade in AI: Four steps towards robust artificial intelligence*, arXiv preprint, 2020.

1.2 Storia dell'intelligenza artificiale in applicazioni creative e progettuali

⁵¹ H. A. Simon, *The sciences of the artificial* (3rd ed.), MIT Press, 1996.

Come osservava Herbert Simon già nel 1969 in *The Sciences of the Artificial*⁵¹, mentre le scienze naturali descrivono fenomeni esistenti, il design opera in un dominio normativo e trasformativo, orientato alla configurazione intenzionale di stati futuri. Progettare implica la costruzione e la selezione di alternative possibili all'interno di uno spazio di soluzioni, definite in relazione a obiettivi, vincoli e criteri di valutazione.

La complessità del design deriva precisamente da questa condizione: il processo progettuale non consiste nella ricerca di una soluzione univoca, ma nella negoziazione tra molteplici variabili eterogenee che spesso risultano tra loro parzialmente incompatibili.

Nel contesto contemporaneo, tale complessità si intensifica ulteriormente a causa della crescente articolazione dei sistemi progettati: prodotti connessi, servizi di elevata dimensionalità, dinamiche non lineari e filiere globali. È in questo scenario che l'AI entra nel design non come semplice innovazione tecnologica, ma come risposta alla crescente complessità progettuale. Tecniche di machine learning e generative design permettono di analizzare grandi quantità di dati e simulare scenari alternativi, ampliando la capacità del progettista di operare su sistemi complessi. L'AI non sostituisce il processo decisionale umano, ma ne ristrutturata le condizioni,

⁵² IBM, *AI product design: The future of intelligent experiences*, n.d.

spostando il focus dalla definizione diretta della soluzione alla costruzione dei sistemi e dei criteri attraverso cui le soluzioni emergono⁵².

Il design occupa una posizione intermedia tra arte, ingegneria e tecnologia, funzionando come uno spazio di traduzione in cui le capacità computazionali dell'intelligenza artificiale vengono trasformate in forme culturali e funzionali. Per questo motivo, la sua storia non può essere separata nettamente da quella dell'arte o dell'innovazione tecnologica, ma emerge precisamente dalla loro sovrapposizione, come risultato di una progressiva convergenza tra pratiche espressive, sistemi computazionali e processi progettuali.

Alla luce di queste premesse, il presente capitolo ricostruisce, in una prospettiva cronologica, le principali tappe attraverso cui l'intelligenza artificiale è entrata nei processi di design, evidenziando come il suo ruolo si sia trasformato nel tempo: da supporto tecnico e computazionale a infrastruttura operativa, fino a configurarsi come agente attivo nelle pratiche progettuali contemporanee.

1.2.1 Il designer controlla la macchina

Nelle prime fasi dell'informatica applicata al progetto, il computer veniva inteso principalmente come uno strumento esecutivo di estrema precisione. In questo paradigma, la gerarchia è chiara: il designer mantiene il controllo totale sulla logica e sulla forma, utilizzando la macchina per estendere le proprie capacità manuali e di calcolo.

Cibernetica: Sistemi adattivi e reattività (anni '50)

Radicata nell'intreccio tra cibernetica e teoria dei sistemi, la progettazione computazionale affonda le sue origini nella metà del XX secolo, quando studiosi come Norbert Wiener iniziano a indagare i meccanismi di comunicazione e controllo nei sistemi complessi, siano essi biologici o artificiali. La cibernetica introduce infatti concetti fondamentali come feedback, adattamento e autoregolazione, che diventeranno centrali per comprendere le dinamiche dei sistemi progettuali contemporanei.

Queste idee gettarono le basi per la comprensione dei sistemi complessi e ispirarono le prime esplorazioni degli

approcci computazionali nella progettazione.

Tra queste, la macchina Musicolour sviluppata nel 1953 da Gordon Pask rappresenta un caso emblematico: progettata per generare risposte luminose a stimoli sonori, essa instaurava un vero e proprio dialogo con il performer, configurandosi come un sistema capace di apprendere e adattarsi nel tempo. Come sottolinea lo stesso Pask qualche anno più tardi, l'aspetto più rilevante non risiedeva nella componente sinestetica, ma nella possibilità di costruire un'interazione in cui l'esecutore addestra la macchina e, al contempo, viene influenzato da essa, dando luogo a una relazione cooperativa, *"per ottenere effetti che [l'esecutore] non avrebbe potuto ottenere da solo"*⁵³.

⁵³ A. I. Miller, *The artist in the machine: The world of AI-powered creativity*, MIT Press, 2019.

Contemporaneamente, esperimenti come le tartarughe Elmer ed Elsie della Machina Speculatrix di W. Grey Walter e l'Homeostat di Ross Ashby esploravano forme di comportamento autonomo e adattivo, anticipando l'idea di sistemi artificiali capaci di reagire all'ambiente⁵⁴.

⁵⁴ Artnet News, *A brief history of AI art, from its humble 1960s beginnings to the present day*, 2021.

Interfaccia: La nascita del medium grafico (anni '60)

Negli anni Sessanta, le prime sperimentazioni computazionali segnano un passaggio fondamentale dall'interazione reattiva di matrice cibernetica a forme più complesse di generazione e progettazione basate su regole⁵⁵. In ambito artistico, questo cambiamento si manifesta con la nascita della computer art, in cui il computer diventa un sistema capace di produrre autonomamente configurazioni formali. In Germania, pionieri come Georg Nees e Frieder Nake, ispirati anche dalle teorie estetiche di Max Bense, sviluppano algoritmi in grado di generare pattern geometrici composti da linee e curve, successivamente tracciati da plotter.

⁵⁵ R. M. Auer, D. Elflein, *Artificial intelligence: Intelligent art, human-machine interaction and creative practice*, Transcript Verlag, 2018.

Le prime opere del 1965 furono esposte presso la Studiengalerie dell'Università di Stoccarda, in occasione della mostra denominata "Computer Grafik" tenutasi dal 5 al 19 Febbraio. Tale evento rappresenta un momento di rottura: non si tratta più di immagini singole create manualmente, ma di risultati derivati da programmi che descrivono intere classi di forme. Come sottolinea Nake, il processo creativo si sposta così dall'oggetto al sistema generativo, iniziando a mettere in crisi l'idea tradizionale di autorialità e suscitando forti reazioni da parte degli artisti dell'epoca⁵⁶.

⁵⁶ F. Nake, *Computer art: A personal recollection*, Proceedings of the 5th Conference on Creativity & Cognition, 2005.



Figura 1.14: Il mainframe IBM 7094 (qui in una foto d'archivio della Columbia University), tipico esempio dei calcolatori degli anni '60⁵³.

Nel frattempo, negli Stati Uniti, A. Michael Noll conduce esperimenti analoghi presso i Bell Labs, utilizzando l'IBM 7094, un computer che occupava un'intera stanza, per generare composizioni visive basate su distribuzioni matematiche⁵³. I suoi lavori, inizialmente nati anche da errori tecnici dei plotter, dimostrano come output computazionali possano essere percepiti come immagini dotate di valore estetico. Noll stesso mette in relazione queste produzioni con opere della storia dell'arte, paragonando ad esempio un suo lavoro all'opera "Ma Jolie" di Picasso, contribuendo a legittimare il computer come mezzo creativo.

Queste sperimentazioni confluiscono simbolicamente nella mostra *Cybernetic Serendipity* del 1968 presso l'Institute of Contemporary Art di Londra, che rappresenta uno dei primi momenti di sintesi tra arte, tecnologia e sistemi interattivi. L'esposizione include opere di artisti come Nam June Paik, con il suo Robot K-456, Bruce Lacey, con un gufo fotosensibile, e Jean Tinguely, che presenta macchine per dipingere in grado di produrre disegni attraverso l'interazione con il pubblico. In questi lavori, il comportamento del sistema e la partecipazione dell'utente diventano elementi centrali dell'opera, se-



gnando il passaggio da un'arte oggettuale a un'arte basata su interazione e autonomia operativa⁵⁴.

Mentre in ambito artistico il computer viene esplorato come generatore di forme e nuovi comportamenti, nel campo della progettazione emerge un cambiamento altrettanto significativo nella relazione tra progettista e macchina. Nel 1963 Ivan E. Sutherland presenta Sketchpad: A Man-Machine Graphical Communication System, considerato il primo software interattivo di Computer-Aided Design⁵⁷. Il sistema introduce un'interazione diretta tra progettista e computer, permettendo di creare e modificare figure geometriche direttamente su schermo attraverso una manipolazione immediata delle forme. Sketchpad non si limita a registrare tracciati grafici, ma opera su entità strutturate, gestendo automaticamente vincoli e relazioni geometriche come parallelismo, perpendicolarità e uguaglianza di lunghezza anche durante le modifiche successive. Il disegno viene così concepito come un sistema di relazioni formali piuttosto che come una rappresentazione statica. Pur non includendo forme di apprendimento o decisione autonoma, Sketchpad dimostra per la prima volta che il computer può partecipare attivamente al processo progettuale, anticipando i principi fondamentali dei sistemi di progettazione assistita successivi⁵⁸.

⁵⁷ W. Hosh, *Ivan Sutherland: American computer scientist*, Encyclopaedia Britannica, 2009.

⁵⁸ I. E. Sutherland, *Sketchpad: A man-machine graphical communication system* (Technical Report No. 574), University of Cambridge Computer Laboratory, 2003.

Figura 1.15: Ivan Sutherland mentre opera sul sistema Sketchpad nel 1963 presso l'MIT, utilizzando una penna ottica per interagire direttamente con lo schermo⁵⁸.



Anche sul piano teorico, emergono approcci che contribuiscono a formalizzare il design come processo strutturato. In *Notes on the Synthesis of Form* (1964), Christopher Alexander introduce l'idea di utilizzare tecniche matematiche e computazionali per risolvere problemi di progettazione. Il suo lavoro sottolinea l'importanza di

⁵⁹ *Parametric Architecture, The evolution of computational design: From algorithms to AI*, n.d.

scomporre le complesse sfide progettuali in componenti più semplici e gestibili, che possano essere affrontate in modo sistematico: un principio che risuona profondamente nelle metodologie di progettazione computazionale odierne⁵⁹.

Codifica: AARON e la generazione autonoma (anni '70)

Negli anni Settanta, la ricerca artistica computazionale compie un ulteriore passo avanti, spostandosi dalla semplice generazione di forme alla simulazione di processi decisionali creativi. Figura centrale di questa fase è Harold Cohen, pioniere nell'uso dell'intelligenza artificiale in ambito artistico, che a partire dal 1973 sviluppa AARON, uno dei primi sistemi in grado di generare autonomamente immagini. A differenza delle precedenti esperienze di computer art, basate su pattern geometrici e regole relativamente semplici, AARON è progettato per produrre immagini riconoscibili, costruendo figure e composizioni attraverso un insieme complesso di istruzioni logiche e matematiche.

Il funzionamento del sistema non consiste nell'esecuzione meccanica di un disegno predefinito, ma nella definizione di regole generative che descrivono come costruire forme, proporzioni e relazioni spaziali. Come Cohen stesso sottolinea in *What is an Image?* (1979), il programma non è concepito come uno strumento di supporto diretto all'artista, ma come un sistema autonomo capace di operare all'interno di un dominio di possibilità definito. All'interno di questo processo, elementi di casualità ven-

Figura 1.16: Harold Cohen con un'installazione di AARON presso il San Diego Museum⁶⁰.



⁶⁰ H. Cohen, *What is an image?*, Proceedings of the 6th International Joint Conference on Artificial Intelligence, Tokyo, Japan, 1979.

gono integrati per introdurre variazione e imprevedibilità, mentre le strutture logiche garantiscono coerenza formale, permettendo al sistema di generare intere classi di immagini piuttosto che singole opere⁶⁰.

Un aspetto cruciale dell'esperimento di Cohen riguarda proprio il rapporto tra controllo e autonomia. Durante lo sviluppo di AARON, l'artista osserva come alcune delle regole implementate producano risultati inattesi, forme che non erano state esplicitamente previste, suggerendo che il sistema fosse in grado di prendere decisioni che si avvicinano a quelle artistiche. È per questo che Cohen afferma che *"la creatività non risiede né nel programmatore né nel programma, ma nel dialogo tra programma e programmatore"*⁶⁵.

È necessario fare una distinzione fondamentale: a differenza dei moderni sistemi di conversione da testo a immagine che assemblano immagini preesistenti da vasti insiemi di dati, AARON crea opere originali basandosi su istruzioni di composizione estetica fornite dal suo creatore. Questo lo rende unico, poiché emula il processo decisionale umano nella creazione artistica⁶¹.

⁶¹ Artribune, *Arte e intelligenza artificiale: Storia e sviluppi*, 2023.

1.2.2 Il designer programma la macchina

Con l'avvento dei sistemi CAD e delle prime logiche parametriche, il ruolo del progettista evolve verso una dimensione più astratta e sistematica. Non si disegna più l'oggetto singolo, ma si definiscono le regole e le istruzioni che lo generano, trasformando la programmazione in un nuovo linguaggio espressivo del design.

Automazione: Il CAD come standard tecnico (anni '80)

Negli anni Ottanta, le tecnologie di progettazione assistita da computer entrano in una fase di consolidamento e diffusione, segnando il passaggio da sperimentazioni pionieristiche a strumenti operativi ampiamente adottati nel mondo del design e dell'ingegneria⁵⁹.

Un ruolo fondamentale in questa transizione è svolto dalla diffusione dei sistemi CAD (Computer-Aided Design), che permettono di produrre disegni tecnici con maggiore precisione, controllo e velocità. In ambito industriale avanzato, un caso emblematico è rappresentato da CATIA, sviluppata inizialmente alla fine degli anni Settanta all'interno di Dassault Aviation e commercializzata a par-



⁶² Dassault Systèmes, *La nostra storia*, n.d.

tire dal 1981 da Dassault Systèmes. Nata per rispondere alle esigenze dell'industria aeronautica, CATIA introduce un salto qualitativo rispetto ai sistemi precedenti, consentendo la modellazione tridimensionale di superfici complesse e la gestione integrata di geometrie altamente sofisticate⁶².

⁶³ D. E. Weisberg, *Autodesk and AutoCAD: A history of CAD innovation*, Shapr3D, 2023.

È con la fondazione di Autodesk nel 1982 e il lancio di AutoCAD che il CAD viene portato ai personal computer. Presentato al COMDEX nello stesso anno, AutoCAD derivava da un software sviluppato da Michael Riddle alla fine degli anni Settanta e si distingue per la sua accessibilità economica e la compatibilità con hardware diffuso. Questo abbattimento dei costi rappresenta un punto di svolta fondamentale: il CAD non è più uno strumento riservato a grandi aziende o centri di ricerca, ma diventa progressivamente accessibile a una vasta comunità di progettisti, architetti e ingegneri⁶³.

⁶⁴ Robots Authority, *The rise and fall of expert systems in the 1980s*, 2025.

Tuttavia, la rilevanza degli anni Ottanta non risiede soltanto nella diffusione degli strumenti, ma anche nella stabilizzazione di un paradigma progettuale. Il computer viene ormai concepito come un'estensione delle capacità tecniche del progettista, orientata principalmente alla rappresentazione, alla precisione e all'efficienza operativa. A differenza delle sperimentazioni artistiche e generative dei decenni precedenti, il focus si sposta su applicazioni deterministiche, in cui il controllo rimane saldamente nelle mani dell'utente e il sistema esegue istruzioni esplicite senza autonomia decisionale⁶⁴.

Simulazione: Ottimizzazione e reti neurali (anni '90)

⁶⁵ O. Sigmund, *Topology optimization: A tool for the tailoring of structures and materials*, Technical University of Denmark, 2000.

Negli anni Novanta, il design computazionale si sposta da una logica prevalentemente rappresentativa e operativa verso un paradigma basato su simulazione, analisi e generazione di alternative. Se negli anni Ottanta il CAD si era affermato come strumento per la produzione e gestione del disegno tecnico, ora il computer inizia a essere utilizzato per prevedere il comportamento dei sistemi progettati, integrando tecniche di simulazione come il Finite Element Method (FEM) e modelli di analisi strutturale e ambientale⁶⁵.

In questo decennio emergono le prime applicazioni di reti neurali artificiali, utilizzate per apprendere relazioni tra parametri progettuali e prestazioni del sistema. Sebbene ancora limitate rispetto agli sviluppi contemporanei, queste tecniche introducono un elemento fondamentale: la possibilità di costruire modelli adattativi ca-

⁶⁶ J. Rhee, H. Jung, T. Kim, *Three decades of machine learning with neural networks in computer-aided architectural design (1990–2021)*, Design Science, 2023.

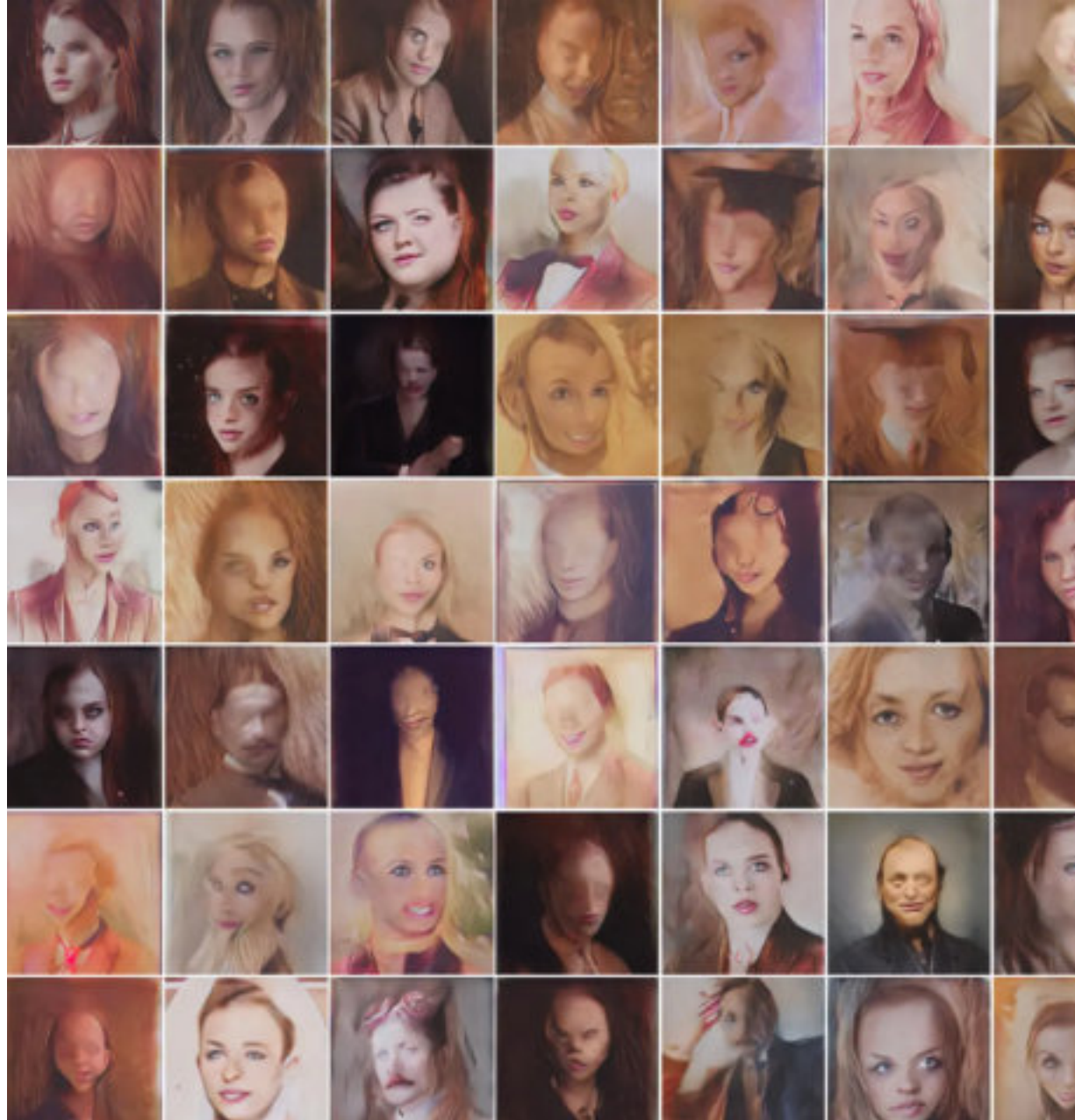
paci di apprendere dai dati, riducendo la necessità di definire esplicitamente tutte le regole del processo progettuale. Ad esempio, le reti neurali vengono impiegate per prevedere il comportamento strutturale di un componente o le prestazioni aerodinamiche di una forma a partire da un insieme di parametri geometrici, consentendo di accelerare i processi di analisi e ottimizzazione all'interno degli ambienti CAD⁶⁶.

Un ulteriore sviluppo rilevante è rappresentato dall'introduzione di algoritmi evolutivi e genetici, come i Genetic Algorithms sviluppati da John Holland, utilizzati per esplorare automaticamente molteplici soluzioni progettuali all'interno di uno spazio di vincoli. Applicati all'ottimizzazione strutturale, alla definizione di geometrie aerodinamiche e alla generazione di configurazioni spaziali, questi metodi permettono di individuare soluzioni ad alte prestazioni attraverso processi iterativi di selezione e variazione.

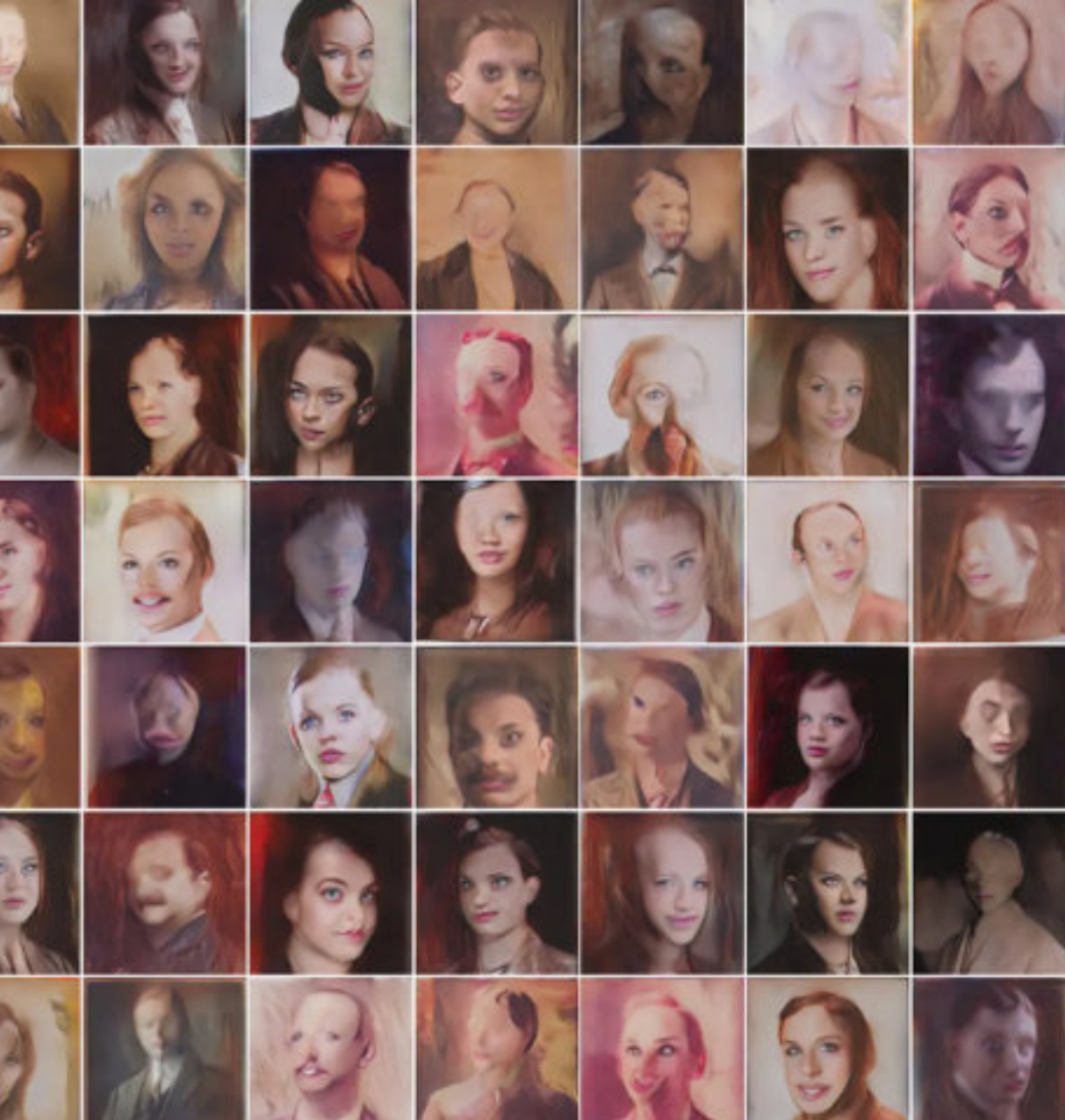
Questi approcci segnano un primo passaggio verso il generative design: il progettista non definisce più una singola soluzione, ma imposta un sistema capace di generare e valutare alternative. Nel campo dell'architettura, tali trasformazioni si traducono in una crescente attenzione alla generazione computazionale della forma. Progettisti come Greg Lynn esplorano geometrie complesse e superfici continue attraverso software di animazione e modellazione avanzata, concependo la forma come risultato di processi dinamici. Progetti come Embryological House (1997) mostrano come variazioni parametriche possano generare intere famiglie di soluzioni, anticipando i principi del design parametrico⁶⁵.

Figura 1.17: Modelli fisici tridimensionali relativi al progetto Embryological House (1997) dell'architetto Greg Lynn⁶⁵.





Sempre in questo decennio si consolida il concetto di "Artificial Intelligence in Design" come disciplina autonoma, volta a studiare l'applicazione dei metodi di intelligenza artificiale ai processi progettuali⁶⁷. A partire dal 1991 vengono organizzate conferenze internazionali dedicate, note come AI in Design, che raccolgono contributi su sistemi esperti, apprendimento automatico, reti neurali e altri approcci cognitivi applicati alla progettazione ingegneristica e architettonica. La First International Conference on Artificial Intelligence in Design (AID'91) fu curata da John S. Gero, un importante ricercatore nel campo dell'IA applicata alla progettazione, che agì come editore e coordinatore delle pubblicazioni della confe-



⁶⁹ Grasshopper3D,
*Grasshopper: Algorithmic
 modeling for Rhino, 2007.*

producibili con metodi tradizionali⁶⁹.

Parallelamente si diffonde un ecosistema di strumenti e linguaggi che favoriscono la democratizzazione della programmazione creativa, come Processing sviluppato da Casey Reas e Ben Fry, che introduce un approccio accessibile alla generazione visiva computazionale. Nel frattempo, i ricercatori creano e rendono pubblici enormi insiemi di dati, come ImageNet, utilizzabili per addestrare algoritmi in grado di catalogare fotografie e identificare oggetti. Infine, programmi di visione artificiale già pronti, come Google DeepDream, permettono ad artisti e pubblico di sperimentare rappresentazioni visive di

come i computer interpretano immagini specifiche⁵⁴.

Sul piano teorico, il contributo di Patrik Schumacher è centrale nella formalizzazione di questo cambiamento. Nel 2008, con il saggio *Parametricism - A New Global Style for Architecture and Urban Design*, Schumacher definisce il "parametricismo" come una nuova epoca stilistica dell'architettura contemporanea, basata sull'interdipendenza sistemica tra elementi progettuali e sull'uso di parametri variabili come principio ordinatore della forma⁷⁰.

⁷⁰ P. Schumacher, *Parametricism: A new global style for architecture and urban design*, Patrik Schumacher Website, 2008.

1.2.3 Il designer dialoga con la macchina

L'ultima frontiera della progettazione computazionale è segnata dal passaggio dal comando al dialogo. Grazie all'integrazione di modelli generativi e sistemi data-driven, il rapporto tra umano e artificiale diventa una negoziazione continua, in cui la macchina agisce come un interlocutore attivo capace di proporre soluzioni e reagire agli intenti progettuali.

Data-driven: L'AI come assistente analitico (2010-2013)

All'inizio degli anni 2010 si assiste a una fase di transizione fondamentale nell'evoluzione della progettazione computazionale, caratterizzata dalla progressiva affermazione del paradigma data-driven e dal ritorno delle tecniche di apprendimento automatico⁷¹. Sebbene le reti neurali fossero già state esplorate nei decenni precedenti, è in questo periodo che, grazie all'aumento della capacità computazionale e alla disponibilità di grandi quantità di dati, il machine learning inizia a dimostrare un'efficacia concreta in ambiti applicativi complessi. Allo stesso tempo, la diffusione delle piattaforme digitali, dei social media e dei sistemi connessi porta a una crescita esponenziale della quantità di dati disponibili. Di conseguenza anche il design inizia a orientarsi verso modelli basati sull'analisi dei comportamenti degli utenti, introducendo pratiche come user analytics, tracciamento delle interazioni e test comparativi⁵².

Si afferma così una prima forma di progettazione data-driven, in cui i dati diventano materia prima del processo e le decisioni progettuali vengono supportate da evidenze quantitative piuttosto che esclusivamente da intuizioni qualitative. Il comportamento degli utenti, le metriche di utilizzo e le informazioni raccolte durante l'interazione

⁷¹ DMG MORI, *Intelligenza artificiale nella produzione*, n.d.

con sistemi digitali diventano elementi fondamentali per guidare le decisioni dei designer. Tuttavia, in questa fase l'intelligenza artificiale assume ancora prevalentemente il ruolo di assistente: analizza dati, identifica pattern e supporta il processo decisionale, senza intervenire direttamente nella generazione delle soluzioni progettuali.

Generativo: La macchina come co-designer (2014-2015)

A partire dalla metà degli anni 2010 si assiste all'affermazione della progettazione generativa come evoluzione diretta della progettazione parametrica e algoritmica⁷².

⁷² J. Lee, *The evolution of AI in design: From CAD to generative models*, Number Analytics Blog, 2025.

Un passaggio chiave in questa trasformazione è rappresentato dallo sviluppo del progetto Autodesk Dreamcatcher (2014), successivamente integrato nei sistemi di generative design di Autodesk Fusion 360. In questi ambienti, il progettista inserisce parametri quali carichi strutturali, materiali, vincoli geometrici, processi produttivi e obiettivi prestazionali, mentre il sistema esplora automaticamente lo spazio delle possibili soluzioni, generando un elevato numero di alternative ottimizzate⁷³.

⁷³ Prisma Tech, *Come l'intelligenza artificiale generativa sta cambiando il mondo del design*, 2024.

In questi anni la progettazione generativa si inserisce pienamente nel contesto dell'Industria 4.0, ridefinendo il rapporto tra progettazione, simulazione e produzione. Le applicazioni industriali di questi strumenti, adottati da aziende come Airbus, General Motors e BMW, dimostrano il loro impatto concreto nella progettazione di componenti strutturali ottimizzati, caratterizzati da riduzione del peso, efficienza dei materiali e miglioramento delle prestazioni⁷¹.

AI art: L'emergenza del nuovo autore (2016-2019)

Tra il 2016 e il 2019 si sviluppa una nuova fase nella relazione tra intelligenza artificiale e progettazione, in cui le tecnologie di machine learning generativo iniziano a essere utilizzate non solo come strumenti operativi, ma come veri e propri dispositivi espressivi (Artnet News, 2021). In particolare, la diffusione delle Generative Adversarial Networks (GAN), introdotte nel 2014, rende possibile la produzione autonoma di immagini complesse e visivamente coerenti, aprendo il campo a una nuova forma di creatività computazionale.

In questi anni si diffondono strumenti come DeepDream e Pix2Pix che permettono la manipolazione e trasformazione neurale delle immagini, ampliando ulteriormente il repertorio espressivo disponibile. Il dataset assume così un ruolo centrale: l'estetica dell'opera non deriva più soltanto dall'intenzione progettuale, ma dalla struttura e dai bias dei dati utilizzati per l'addestramento⁶¹. In questo contesto emerge il fenomeno dell'AI art, in cui l'intelligenza artificiale non si limita ad assistere il processo creativo, ma partecipa attivamente alla generazione dell'opera. L'immagine non è più progettata direttamente dall'autore umano, ma deriva dall'interazione tra algoritmo, dataset e processo di addestramento.

Artisti come Mario Klingemann lavorano da oltre un decennio con reti neurali e GAN, esplorando la produzione di immagini che mettono in discussione i concetti tradizionali di autorialità, originalità e intenzionalità artistica. Klingemann rappresenta una figura centrale della cosiddetta nuova avanguardia computazionale, in cui l'opera emerge dall'interazione tra algoritmo, dataset e intervento umano⁵³.

Mike Tyka è tra i primi a riconoscere il potenziale espressivo di sistemi come DeepDream e, successivamente, delle GAN, utilizzandoli per la realizzazione di opere come *Portraits of Imaginary People*. In queste produzioni, la generazione di volti sintetici evidenzia la capacità delle reti neurali di costruire nuove identità visive a partire dai dati, rafforzando l'idea di una creatività computazionale autonoma. Come afferma lo stesso Tyka, in una prospettiva computazionale della mente, se il cervello può essere inteso come una macchina, allora anche le macchine possono essere considerate capaci di produrre forme di creatività⁵³.

Figura 1.18: Dettaglio del progetto *Neurography* (2018) dell'artista algoritmico Mario Klingemann, esposto presso The Photographers' Gallery⁵³.

Sougwen Chung invece sviluppa pratiche artistiche basate sulla collaborazione tra essere umano e intelligenza artificiale, indagando forme di co-creazione in cui il gesto umano e l'elaborazione algoritmica si intrecciano in tempo reale⁵⁴.

Nel frattempo, collettivi come *Obvious* contribuiscono alla diffusione mediatica del fenomeno, rendendo visibili al grande pubblico il potenziale estetico delle reti neurali.

Un momento cruciale di questa fase è rappresentato

dall'ottobre del 2018, quando per la prima volta nella sua storia, la casa d'aste Christie's a New York mette all'asta un'opera d'arte creata da un'intelligenza artificiale, nello specifico da una GAN⁶¹.

L'opera Portrait of Edmond de Belamy è stata messa all'asta con una stima tra i 7.000 e i 10.000 dollari, per poi essere venduta per la cifra sorprendente di 432.500 dollari. L'opera, firmata non da un artista ma dalla formula matematica dell'algoritmo, rende esplicita la ridefinizione del concetto di autore.

Workflow: L'era della partnership creativa (2020-oggi)

Negli anni 2020 si entra in una fase ulteriore di integrazione dell'intelligenza artificiale nei processi progettuali, definibile come era della Generative Artificial Intelligence (GenAI) applicata al design. In questo contesto, l'AI non si configura più soltanto come strumento di supporto o ottimizzazione, ma come componente attiva e diffusa all'interno dell'intero workflow, contribuendo in modo diretto alle fasi di ideazione, sviluppo, validazione e produzione⁷⁴.

⁷⁴ K. Chen, Y. Wang, L. Zhang, *How generative AI supports human in conceptual design*, Design Science, 2025.

Nella fase di ricerca, strumenti di analisi basati su AI permettono di elaborare grandi dataset provenienti da piattaforme digitali e social, come Facebook, X e LinkedIn, generando insight utili alla definizione del problema. Durante il brainstorming e l'ideazione, modelli generativi vengono utilizzati per produrre rapidamente testi, immagini, mockup e wireframe, facilitando l'esplorazione concettuale e riducendo il tempo necessario alla visualizzazione delle idee⁷⁴. Strumenti come Figma integrano funzionalità di automazione e prototipazione assistita, consentendo la creazione di interfacce e sistemi interattivi in tempo reale. In questo contesto, si afferma quindi il concetto di AI come partner creativo, in cui l'intelligenza artificiale non viene utilizzata esclusivamente per automatizzare compiti, ma per stimolare il pensiero progettuale, generare alternative inattese e ampliare le possibilità espressive⁷⁵. Ad esempio DALL-E, Stable Diffusion o Adobe Firefly sono usati per creare moodboard visuali, esplorazioni cromatiche o scenari alternativi.

⁷⁵ M. Ozsoy, *The impact of generative AI on creative design processes*, DergiPark, 2025.

Nella fase di prototipazione, l'intelligenza artificiale permette di simulare il comportamento degli utenti e prevedere scenari d'uso prima della realizzazione fisica del prodotto⁵². Le interfacce possono essere testate virtual-

mente attraverso modelli che replicano diversi profili di utilizzo, migliorando la qualità delle decisioni progettuali. Durante l'iterazione, strumenti AI accelerano significativamente i cicli di sviluppo, automatizzando operazioni complesse e riducendo tempi e costi. Tecniche come l'A/B testing automatizzato consentono di confrontare diverse soluzioni progettuali in condizioni controllate, identificando rapidamente le configurazioni più efficaci.

Nella fase di distribuzione, l'AI supporta la scalabilità dei sistemi e la coerenza delle esperienze utente su scala

Figura 1.19: L'opera Edmond de Belamy (2018) creata dal collettivo parigino Obvious utilizzando un algoritmo di tipo Generative Adversarial Network (GAN)⁶¹.



globale, automatizzando l'adattamento dei prodotti a contesti e mercati differenti⁷¹. Questo contribuisce a una crescente integrazione tra design, sviluppo e gestione del prodotto, che non sono più fasi separate ma componenti interconnesse di un unico sistema.

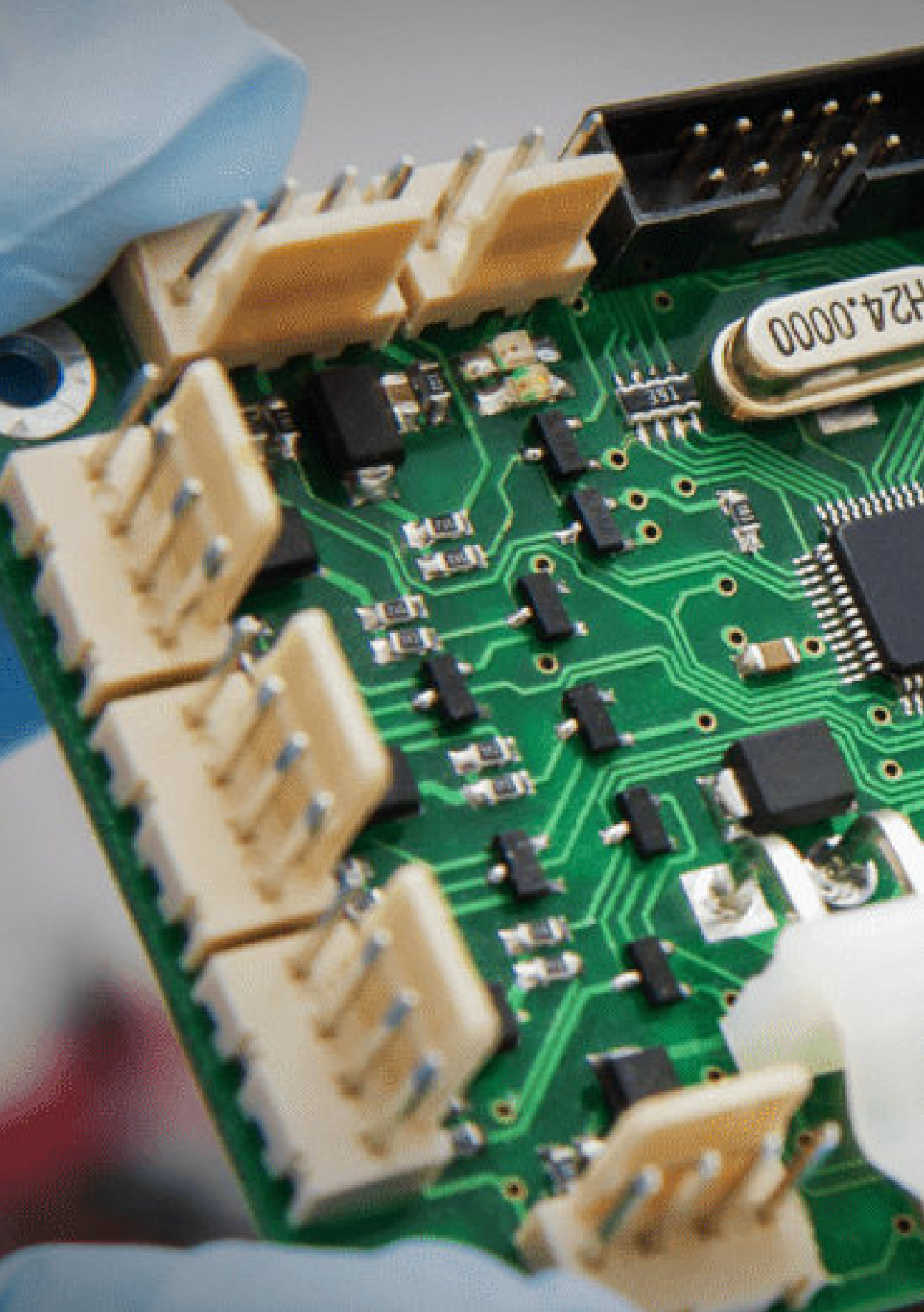
Si riscontrano notevoli progressi anche in ambito artistico: nel 2021 viene creato quello che può essere definito il degno successore di AARON, Botto. Ideato dall'artista tedesco Mario Klingemann, dall'imprenditore Simon Hudson e dal programmatore Ziv Epstein, Botto utilizza un generatore di immagini basato sull'intelligenza artificiale, simile a DALL-E o MidJourney. Tuttavia, ciò che lo distingue è il suo "modello di gusto". L'artista IA è in grado di adattarsi continuamente alle preferenze dei suoi col-

lezionisti, modificando l'estetica delle sue opere in base al feedback della sua comunità di oltre 5.000 partecipanti. Allo stesso tempo, Botto rappresenta un'evoluzione significativa verso modelli sempre più autonomi: l'integrazione di modelli linguistici avanzati e basi di conoscenza consente al sistema non solo di generare immagini, ma anche di articolare discorsi sulle proprie produzioni, avvicinandosi a forme embrionali di coerenza espressiva. Botto segna un ulteriore passaggio nella trasformazione dell'intelligenza artificiale da strumento a soggetto operativo, suggerendo la possibilità di sistemi capaci di sviluppare nel tempo una propria identità estetica⁷⁶.

⁷⁶ Domus, *The evolution of AI artists: From AARON to Botto*, 2025.

Ciò che emerge come dato certo è che la visione più matura e contemporanea non interpreta l'intelligenza artificiale come una tecnologia sostitutiva, bensì come un agente radicalmente potenziante, capace di estendere i confini della pratica progettuale attraverso una collaborazione costante tra intuizione umana e potenza computazionale⁷⁷.

⁷⁷ Officina dei Segni, *L'intelligenza artificiale nel design: Collaborazione creativa o minaccia?*, 2025.





panorama internazionale

2.1 Il panorama internazionale dell'intelligenza artificiale nel design

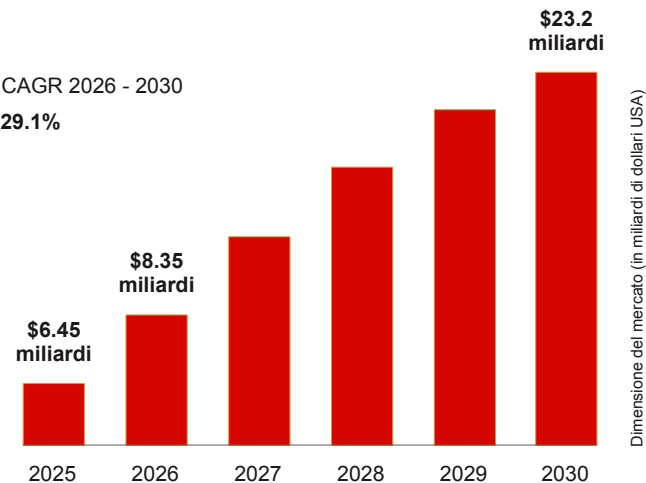
L'adozione dell'IA nel settore creativo sta ridisegnando gli equilibri economici e professionali su scala globale. Analizzare lo stato dell'arte significa osservare come diverse aree geografiche e settori industriali stiano integrando queste tecnologie, definendo nuove traiettorie di crescita e competitività.

2.1.1 Evoluzione del mercato globale dell'AI nel design

L'integrazione dell'Intelligenza Artificiale nel design sta vivendo una fase di crescita esponenziale, che sta trasformando radicalmente i processi creativi e produttivi. Le stime di mercato indicano che il valore globale dell'AI in questo settore, che ammontava a 6,45 miliardi di dollari nel 2025, è destinato a salire a 8,35 miliardi di dollari nel 2026, con un tasso di crescita annuale composto (CAGR) del 29,5%. Questa rapida espansione nel periodo storico è stata alimentata dall'adozione diffusa di strumenti CAD (Computer-Aided Design) nei flussi di lavoro, dall'introduzione precoce di software di simulazione e dalla crescente domanda di cicli di sviluppo prodotto più rapidi ed efficienti. Guardando al futuro, le proiezioni indicano che il mercato raggiungerà i 23,2 miliardi di dollari entro il 2030, mantenendo un CAGR del 29,1%. Questa crescita sostenuta sarà guidata da molteplici fattori, tra

cui l'integrazione dell'AI con i sistemi PLM (Product Life-cycle Management), la crescente richiesta di design leggeri e ottimizzati, l'adozione di piattaforme di progettazione basate sul cloud e l'espansione delle tecnologie di digital twin. Le tendenze principali del prossimo decennio includeranno l'ottimizzazione tramite design generativo, la prototipazione virtuale basata sull'AI, la simulazione automatizzata del design e la personalizzazione assistita dall'intelligenza artificiale.

Figura 2.1: Rapporto sul mercato dell'Intelligenza Artificiale nel Design 2026⁷⁸. Grafico rielaborato.



Nello specifico, il mercato dell'intelligenza artificiale applicata al design comprende i ricavi generati da soggetti che offrono servizi quali l'ottimizzazione progettuale basata su AI, la generazione automatizzata di prototipi, la previsione delle prestazioni, l'adattamento dinamico dei progetti e l'analisi dei dati finalizzata al miglioramento continuo del design. Tale mercato include anche la commercializzazione di software di progettazione basati su AI, strumenti CAD intelligenti, piattaforme di design automatizzato e soluzioni avanzate di simulazione per lo sviluppo prodotto. I valori economici sono calcolati "franco fabbrica", ossia riferiti al valore dei beni venduti dai produttori o creatori, sia lungo la filiera (produttori a valle, grossisti, distributori, rivenditori) sia direttamente al cliente finale, includendo inoltre i servizi correlati offerti dai fornitori⁷⁸.

⁷⁸ The Business Research Company, *Artificial Intelligence (AI) In Industrial Design Global Market Report, 2026*.

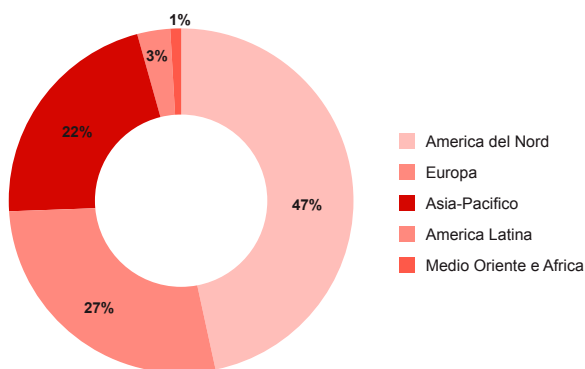
Nel contesto progettuale, l'intelligenza artificiale generativa si propone di potenziare in modo significativo la creatività, l'efficienza e la capacità di problem solving. Essa supporta i designer lungo diverse fasi del processo, intervenendo nel flusso di lavoro attraverso l'automazione di attività ripetitive, quali la produzione di elaborati tecnici e la ricerca sui materiali. Parallelamente, gli algo-

⁷⁹ Precedence Research, *Generative AI in Industrial Design Market, 2025*.

Figura 2.2: Quota di mercato dell'Intelligenza Artificiale Generativa nel Design Industriale, per area geografica, 2025 (%)⁷⁹. Grafico rielaborato.

ritmi generativi sono in grado di produrre varianti e concept progettuali sulla base di parametri e vincoli definiti, fungendo da catalizzatori di nuove idee. Questo scenario, caratterizzato da una crescente domanda di soluzioni progettuali complesse e personalizzate, consente ai designer di focalizzarsi su dimensioni più strategiche e ad alto valore aggiunto. Inoltre, la possibilità di valutare prestazioni, efficienza ed estetica dei progetti in ambienti virtuali avanzati contribuisce a una significativa riduzione di tempi e risorse nel ciclo di sviluppo⁷⁹.

Quota di mercato dell'IA generativa nel design, per regione, 2025 (%)



2.1.2 Le forze trainanti del mercato

Dietro l'espansione esponenziale dell'IA nel design agiscono motori strutturali che vanno oltre la semplice innovazione tecnologica. Dalla trasformazione dei modelli produttivi alla necessità di cicli di sviluppo sempre più rapidi, diverse spinte sistemiche stanno accelerando la transizione verso flussi di lavoro interamente digitalizzati.

Smart Manufacturing e Industria 4.0

L'intelligenza artificiale applicata al design trova una delle sue principali aree di sviluppo nella produzione intelligente, dove contribuisce in modo determinante all'ottimizzazione dei processi produttivi, alla manutenzione predittiva delle apparecchiature e al miglioramento della qualità del prodotto attraverso l'analisi dei dati in tempo reale e l'automazione avanzata. In questo contesto, l'AI non si limita a supportare le fasi esecutive, ma interviene anche a monte, influenzando le decisioni progettuali sulla base di feedback continui provenienti dal sistema pro-

duttivo.

Un indicatore significativo della diffusione di tali paradigmi è fornito dalla Society of Operations Engineers, che nel gennaio 2025 ha stimato come l'80% dei produttori britannici preveda di implementare soluzioni riconducibili all'Industria 4.0 entro il 2025. Attualmente, il 67% delle imprese dichiara di essere a conoscenza di tali tecnologie, mentre il 23% ha già adottato sistemi di digitalizzazione e automazione e il 62% si trova in fase di pianificazione della transizione. Questi dati evidenziano come l'evoluzione verso modelli produttivi intelligenti rappresenti un driver strutturale per l'adozione dell'intelligenza artificiale nel design⁷⁸.

L'integrazione dell'intelligenza artificiale e degli strumenti digitali avanzati sta fondamentalmente ridefinendo il flusso di lavoro del design, migliorando sia la velocità che la capacità creativa. Le piattaforme di AI generativa consentono ai designer industriali di iterare rapidamente i concetti, ottimizzare la topologia strutturale per la produzione e visualizzare i prodotti finali con un realismo senza precedenti prima della prototipazione fisica. A conferma di questo cambiamento metodologico, il Canva Visual Economy Report del giugno 2024 ha rilevato che l'82% dei leader aziendali globali ha utilizzato strumenti basati sull'AI per produrre contenuti visivi nell'ultimo anno, segnalando un'adozione diffusa di queste tecnologie per migliorare le prestazioni aziendali. Questo supporto tecnologico contribuisce al più ampio contributo economico del settore; secondo il National Endowment for the Arts nel marzo 2024, il settore artistico e culturale, che include il design, ha contribuito con un record di 1,1 trilioni di dollari all'economia statunitense, consentendo ai designer di soddisfare le crescenti pressioni sui tempi di commercializzazione e di offrire soluzioni altamente differenziate⁸⁰.

⁸⁰ Yahoo Finance, *Industrial Design Market Report 2026*, 2026.

Accessibilità del Cloud Computing e HPC scalabile

Un ulteriore driver fondamentale per la diffusione dell'AI generativa nel design è rappresentato dalla crescente accessibilità delle infrastrutture di calcolo avanzato, in particolare attraverso il cloud computing e l'High-Performance Computing (HPC) scalabile. I flussi di lavoro tipici del design computazionale e dell'ingegneria dei prodotti, infatti, richiedono ingenti risorse computazionali per l'esecuzione di simulazioni complesse, l'ottimizzazione multi-obiettivo e l'analisi di grandi volumi di dati.

Le piattaforme cloud consentono di accedere on-demand a GPU ad alte prestazioni e a cluster di calcolo distribuiti, riducendo significativamente i costi iniziali di investimento (CAPEX) e trasformandoli in costi operativi (OPEX) più flessibili. Questo paradigma abilita non solo una maggiore scalabilità, ma anche forme di collaborazione distribuita in tempo reale, con impatti diretti sulla riduzione dei cicli di sviluppo e sull'accelerazione del processo di innovazione.

Inoltre, gli aggiornamenti continui dell'infrastruttura da parte degli hyperscaler garantiscono l'accesso costante a hardware e framework AI di ultima generazione, senza interruzioni operative o necessità di upgrade interni. La progressiva riduzione dei costi delle GPU e la maturazione degli ecosistemi cloud contribuiscono ulteriormente ad abbassare le barriere all'ingresso, rendendo queste tecnologie accessibili a un numero sempre più ampio di attori, incluse PMI e studi di design. Ad esempio, nel novembre 2023, Autodesk ha introdotto Autodesk AI, un assistente basato sull'AI integrato nella sua piattaforma Fusion e nell'ecosistema di design più ampio per supportare i flussi di lavoro di design generativo, modellazione e ingegneria.

2.1.3 Analisi di segmentazione del mercato

Per comprendere la portata reale della trasformazione in atto, è necessario scomporre il mercato dell'IA in base alle sue diverse declinazioni operative. La distinzione tra software, servizi e settori di applicazione finale rivela dove si concentrano gli investimenti e quali ambiti stanno guidando l'innovazione progettuale.

Segmentazione del mercato per tipologie: software e servizi

Il mercato si articola in due segmenti principali: software e servizi. Il segmento software comprende piattaforme di design autonome, moduli AI integrati in CAD/CAE/PLM, strumenti di automazione della simulazione, soluzioni di ottimizzazione basate su AI, co-piloti per ingegneri e motori generativi 3D. Il segmento dei servizi include invece la personalizzazione dei modelli AI, l'integrazione di sistema, l'implementazione e l'ottimizzazione in ambiente cloud, nonché lo sviluppo di flussi di lavoro di ingegneria basati su AI.

Il segmento software domina il mercato grazie alla diffusione capillare di strumenti CAD/CAE potenziati dall'AI, piattaforme di design generativo e moduli integrati negli ecosistemi di ingegneria esistenti. Le imprese tendono a privilegiare gli investimenti in software, in quanto incidono direttamente sull'automazione del design, sull'efficienza delle simulazioni e sulle capacità di ottimizzazione. Inoltre, i modelli di licenza SaaS rafforzano la concentrazione dei ricavi in questo segmento, mentre l'integrazione di co-piloti AI e motori generativi 3D nelle piattaforme tradizionali ne accelera ulteriormente la domanda.

Il segmento dei servizi è invece destinato a registrare il tasso di crescita più elevato, trainato dalla crescente necessità di personalizzazione dei modelli AI, supporto all'integrazione e ottimizzazione delle implementazioni cloud. L'introduzione dell'AI generativa in ambienti di ingegneria legacy richiede spesso attività di consulenza e trasformazione dei flussi di lavoro. Con il passaggio dalle fasi pilota a un'adozione su scala aziendale, aumenta significativamente la domanda di integrazione di sistema e implementazione di workflow di ingegneria AI. Questa esigenza di supporto operativo e soluzioni su misura sostiene la crescita del segmento. Un esempio è rappresentato dall'ampliamento, nel giugno 2023, delle capacità di design generativo e AI di PTC nella piattaforma Creo, con un potenziamento degli strumenti di ottimizzazione e simulazione per semplificare i flussi di lavoro e accelerare l'adozione dell'AI nel design.

Segmentazione del mercato per settore di utilizzo finale

Per quanto riguarda i settori di utilizzo finale, il mercato si suddivide in automobilistico, aerospaziale e difesa, macchinari industriali e manifattura, elettronica di consumo, energia e potenza, nautica e costruzione navale, robotica e automazione, oltre ad altri ambiti applicativi. Il segmento automobilistico detiene la quota di mercato maggiore, grazie all'adozione estesa di tecnologie di ottimizzazione del design basate su AI, alleggerimento strutturale e prototipazione rapida. Gli OEM (Original Equipment Manufacturer) e i fornitori di primo livello utilizzano l'AI generativa per la progettazione di componenti, l'ottimizzazione dei pacchi batteria nei veicoli elettrici, la modellazione aerodinamica e il miglioramento delle simulazioni di crash. La forte pressione su time-to-market, efficienza energetica e sviluppo di veicoli elettrici e autonomi continua a sostenere elevati investimenti in piattaforme di ingegneria basate su AI, consolidando la

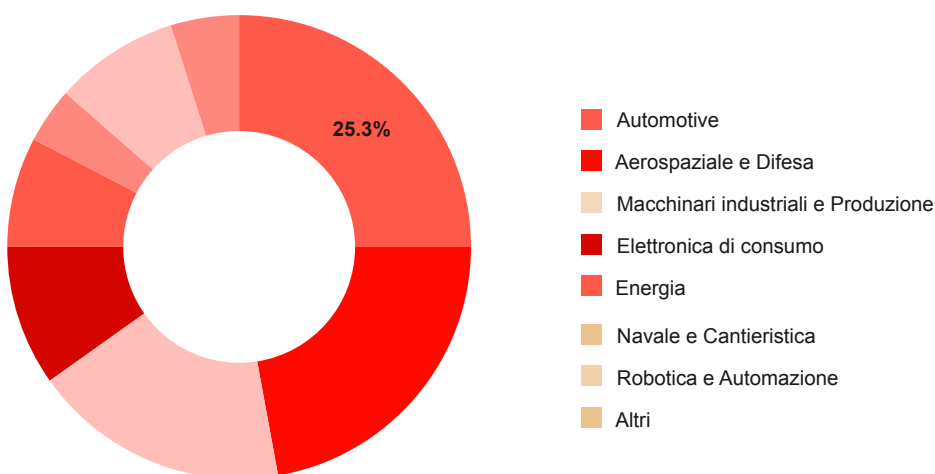
leadership del settore.

Il segmento della robotica e dell'automazione è invece previsto in più rapida crescita, spinto dalla crescente diffusione di robot intelligenti, cobot e sistemi di automazione avanzata in ambiti quali manifattura, logistica e sanità. In questo contesto, l'AI generativa abilita l'iterazione rapida del design meccanico, l'ottimizzazione dei percorsi di movimento, lo sviluppo di componenti leggeri e l'ingegneria dei sistemi embedded. L'evoluzione verso Industria 4.0, fabbriche intelligenti e sistemi autonomi sta accelerando la domanda di soluzioni robotiche altamente personalizzate e ottimizzate tramite AI, rendendo questo segmento il più dinamico in termini di crescita⁸¹.

⁸¹ Fortune Business Insights, *Generative AI in Product Design & Engineering Market, 2026*.

Figura 2.3: Quota di mercato globale dell'IA generativa nel design e nell'ingegneria di prodotto, per settore di utilizzo finale, 2025⁸¹. Grafico rielaborato.

Quota di mercato globale dell'IA generativa nel design e ingegneria del prodotto, per settore di utilizzo finale, 2025



2.2 Key player globali e strategie nell'ai per il design

L'evoluzione dell'Intelligenza Artificiale sta ridisegnando il panorama del design su scala globale. In diverse aree geografiche emergono attori chiave portatori di approcci e filosofie progettuali distintive: dal potenziamento tecnico nordamericano alla governance responsabile europea, fino all'infrastruttura digitale totale cinese e all'integrazione verticale di fabbrica sudcoreana. Il presente capitolo analizza le strategie aziendali dei principali player, esaminando le modalità di integrazione dell'AI nei processi di design e produzione e le tendenze emergenti che stanno ridefinendo il futuro della disciplina. Comprendere in che modo il profilo del designer sia destinato a evolversi permetterà di orientare in maniera più consapevole lo sviluppo dei percorsi formativi, affinché rispondano in modo coerente alle esigenze reali del mercato e alle traiettorie future della professione.

2.2.1 Strategie aziendali per area geografica

L'approccio all'intelligenza artificiale non è uniforme, ma riflette filosofie culturali e politiche industriali differenti a seconda del contesto geografico. Dalla visione nordamericana del potenziamento tecnico a quella europea della regolamentazione etica, ogni regione sta definendo una propria via allo sviluppo tecnologico.

USA: L'AI come potenziamento tecnico

Negli Stati Uniti l'Intelligenza Artificiale è concepita anzitutto come un acceleratore

tecnologico, orientato alla velocità di calcolo e alla generazione di geometrie complesse. Tale impostazione si riflette con chiarezza nelle strategie dei principali attori dell'industria software.

Autodesk, con il lancio di Fusion 2026, ha segnato una discontinuità significativa rispetto alla modellazione geometrica tradizionale, inaugurando il paradigma del cosiddetto Neural CAD. L'innovazione centrale consiste nell'integrazione di modelli di Intelligenza Artificiale nativi che abilitano un'interazione bimodale con il progettista: attraverso il linguaggio naturale è possibile tradurre intenti progettuali in geometrie solide editabili, mentre la funzionalità AutoConstrain impiega algoritmi di machine learning per assegnare automaticamente i vincoli meccanici, come tangenze, parallelismi e quote, a uno schizzo. Il software si configura così come un vero e proprio "collaboratore tecnico", capace di ridurre in modo sostanziale il tempo dedicato ad attività ripetitive, quali la generazione automatica di tavole 2D a partire dai modelli tridimensionali, preservando al contempo la piena modificabilità dei parametri⁸². La direzione strategica è inequivocabile: abbassare la soglia di competenza tecnica necessaria per operare nel CAD, spostando il valore del designer dalla padronanza dello strumento alla capacità di istruirlo.

⁸² Autodesk, *Autodesk Fusion 2026 / Neural CAD*, Prisma Tech, 2026.

In parallelo, la partnership tra NVIDIA e Synopsys, annunciata al GTC 2026, ha esteso i confini del design ingegneristico attraverso il concetto di Agentic AI. Sfruttando la piattaforma Omniverse, le due aziende hanno introdotto uno stack tecnologico nel quale agenti AI autonomi orchestrano flussi di lavoro complessi, dalla progettazione di microchip alla verifica di sistemi fisici. L'integrazione di modelli di fisica reale riduce il divario tra simulazione e realtà, consentendo la generazione di dati sintetici ad alta fedeltà che accelerano lo sviluppo di sistemi robotici e veicoli autonomi. La progettazione non avviene più in isolamento, bensì all'interno di un "gemello digitale" fotorealistico e fisicamente accurato. L'obiettivo è l'automazione dei flussi decisionali: gli agenti AI non si limitano a eseguire compiti, ma validano il progetto simulando milioni di ore di funzionamento in pochi secondi all'interno di Omniverse⁸³.

⁸³ NVIDIA + Synopsys, *Synopsys Showcases NVIDIA Partnership Impact and Ecosystem Innovation at GTC 2026*, PR Newswire, 2026.

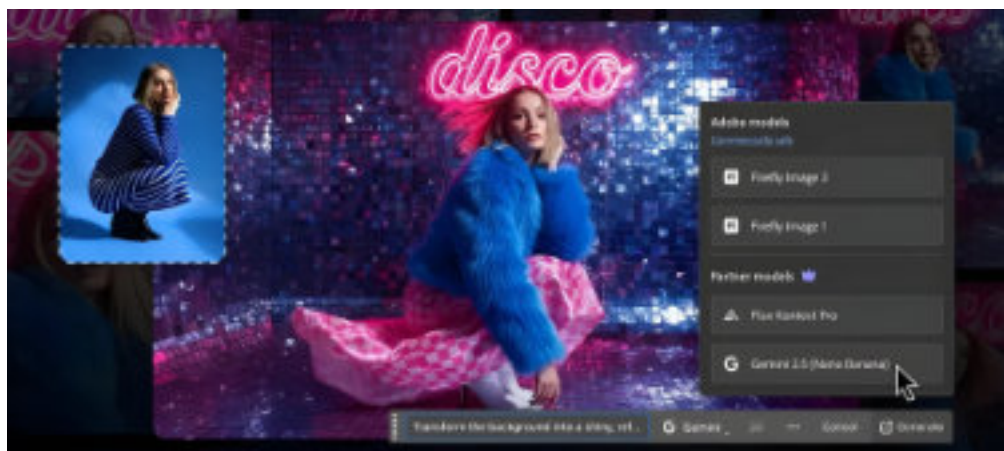
⁸⁴ Ansys, *Ansys 2025 R2: Engineering Copilot, 2025.*

⁸⁵ Ansys, *Ansys 2025 R2 Accelerates Innovation with AI, 2025.*

Un ulteriore caso emblematico è rappresentato da Ansys, che con la versione 2025 R2 ha introdotto l'Ansys Engineering Copilot, un assistente virtuale addestrato su decenni di dati tecnici proprietari⁸⁴. Lo strumento consente ai progettisti di navigare processi di simulazione complessi attraverso un'interfaccia conversazionale. Il dato più significativo riguarda le prestazioni: calcoli critici nei settori del 5G/6G e dell'aerospazio risultano fino a 17 volte più veloci grazie all'accelerazione GPU e all'AI.⁸⁵ Ne consegue una democratizzazione della simulazione: anche progettisti meno esperti possono validare le proprie intuizioni con un rigore scientifico precedentemente riservato a specialisti, riducendo i rischi progettuali e il time-to-market. Il focus strategico è la velocità di iterazione, che abilita un approccio al design "per tentativi infiniti" in cui l'errore non genera costi temporali, ma dati utili all'ottimizzazione.

Un altro caso di particolare rilevanza nel panorama nordamericano è rappresentato da Adobe, che con l'ecosistema Firefly ha ridefinito il ruolo dell'AI nella produzione creativa. A partire dal lancio iniziale come generatore di immagini, Firefly si è evoluto in una piattaforma unificata che integra strumenti AI per la generazione e l'editing di immagini, video, audio e grafica vettoriale all'interno di un unico studio creativo. Con il rilascio dell'Image Model 5, nell'ottobre 2025, Adobe ha raggiunto la generazione nativa in risoluzione 4 MP senza upscaling, con un realismo fotografico che eccelle nella resa di illuminazione, texture e ritratti anatomicamente accurati. Il modello introduce inoltre la funzionalità Prompt to Edit, che permette di descrivere le modifiche desiderate in linguaggio naturale. La piattaforma si distingue per il suo approccio "commercially safe", poiché i modelli proprietari sono addestrati su contenuti con licenza, garantendo la sicurezza legale per l'uso commerciale, e per l'integrazione di modelli di terze parti (Google Imagen, OpenAI GPT Image,

Figura 2.4: Adobe lancia Firefly Boards e introduce nuovi modelli AI in Photoshop⁸⁶.



⁸⁶ Adobe, *Adobe Firefly Delivers Groundbreaking AI Audio, Video and Imaging Innovations and New Models in All-In-One Creative AI Studio*, 2025.

⁸⁷ Adobe, *Adobe and NVIDIA Announce Strategic Partnership to Deliver the Next Generation of Firefly Models and Creative, Marketing and Agentic Workflows*, 2026.

⁸⁸ Figma, *Your creativity, unblocked with Figma AI*, 2025.

Black Forest Labs Flux) direttamente nell'interfaccia, offrendo ai creativi la libertà di scegliere l'estetica più adatta al proprio progetto⁸⁶.

Nel marzo 2026, la partnership strategica tra Adobe e NVIDIA ha ulteriormente consolidato questa visione. La collaborazione prevede lo sviluppo della prossima generazione di modelli Firefly basati sulle tecnologie di calcolo avanzato NVIDIA, l'introduzione di flussi di lavoro agentic per l'automazione delle pipeline creative e di marketing, e il lancio di una soluzione cloud-native per la creazione di "gemelli digitali 3D" che preservano l'identità visiva del brand, costruita sulle librerie NVIDIA Omniverse e sullo standard OpenUSD⁸⁷. Tale sinergia esprime una tendenza chiave dell'approccio nordamericano: l'AI non come sostituto del creativo, ma come infrastruttura intelligente che scala la produzione mantenendo il controllo autoriale e la conformità del brand.

Un contributo altrettanto significativo proviene da Figma, che ha integrato l'AI in modo capillare nella piattaforma di design collaborativo. Il lancio di Figma Make consente di trasformare descrizioni in linguaggio naturale in prototipi funzionali con codice generato automaticamente, abbattendo la barriera tra design e sviluppo. L'AI Assistant offre suggerimenti contestuali su layout, spaziatura e gerarchia visiva, genera scaffolding di prototipi a partire da prompt, e produce automaticamente specifiche tecniche dai file di progetto, funzionalità che riducono sensibilmente il carico di lavoro ripetitivo e accelerano il passaggio dall'ideazione alla produzione. Il sistema Visual Search, basato sull'analisi delle caratteristiche visive, consente inoltre di individuare componenti simili all'interno di librerie complesse caricando semplicemente uno screenshot, promuovendo la coerenza dei design system nelle grandi organizzazioni⁸⁸. L'approccio di Figma si allinea alla tendenza nordamericana della democratizzazione degli strumenti: l'AI non richiede competenze tecniche avanzate, ma risponde a intenti espressi in forma naturale, rendendo il design accessibile a una platea più ampia di professionisti.

Tra i principali player nordamericani, PTC rappresenta un altro caso particolarmente significativo, in quanto combina un'impostazione orientata alla performance tipica del contesto statunitense con un approccio progettuale che, per certi aspetti, si avvicina alla tradizione ingegneristica europea. Con il rilascio di Creo 12, l'azienda ha rafforzato le proprie capacità di design guidato dall'intelligenza artificiale integrando la fisica termica nella Gene-

rative Topology Optimization (GTO), un motore di design generativo che opera direttamente all'interno dell'albero parametrico del modello CAD. A differenza di molte piattaforme concorrenti, che producono mesh astratte da ricostruire manualmente, Creo mantiene ogni dimensione, equazione e family table pienamente editabile anche dopo l'ottimizzazione, garantendo continuità operativa e controllo diretto da parte del progettista.

⁸⁹ Engineering.com, *Leveraging simulation-driven design with PTC Creo, 2025.*

⁹⁰ Tristar Digital Thread Solutions, *PTC Creo AI Optimization: How to Make Faster Design Decisions, 2026.*

La simulazione in tempo reale (Creo Simulation Live), sviluppata con Ansys, consente inoltre di visualizzare istantaneamente stress, deformazioni e comportamento termico durante la modellazione, trasformando l'analisi da fase di validazione a feedback continuo integrato nel processo progettuale⁸⁹. Con l'introduzione del Creo AI Assistant nella versione Creo+ 13.0, PTC ha infine abilitato un supporto conversazionale direttamente nell'ambiente di lavoro, facilitando la risoluzione degli errori e l'ottimizzazione dei flussi, con un miglioramento del 15% nella rigenerazione dei modelli e un'accelerazione fino a 10 volte nei controlli di interferenza globale⁹⁰. Pur inserendosi pienamente nel paradigma nordamericano, PTC si distingue per una forte attenzione alla tracciabilità, alla producibilità e alla coerenza fisica delle soluzioni progettuali, mostrando una chiara convergenza con i principi europei della "verità ingegneristica" (esplorati di seguito).

Europa: governance e verità scientifica

Nel contesto europeo, l'approccio all'AI nel design si distingue per una forte impostazione regolatoria e sistemica. L'attenzione è rivolta in particolare ai Digital Twin come strumenti di validazione e controllo, alla protezione della proprietà intellettuale e all'uso efficiente delle risorse lungo l'intero ciclo di vita del prodotto. L'intelligenza artificiale non è solo un abilitatore tecnologico, ma un presidio di affidabilità e conformità, orientato al rispetto delle normative e dei vincoli fisici e materiali. La sostenibilità rappresenta un asse centrale: non solo in termini di riduzione dei materiali impiegati, ma anche di gestione responsabile dei dati e coerenza con i principi dell'ingegneria reale. Di conseguenza, il design viene interpretato come un'attività intrinsecamente responsabile, inserita in un sistema complesso in cui ogni decisione progettuale ha implicazioni tecniche, ambientali e normative interconnesse.

Dassault Systèmes propone una visione dell'AI rigorosamente ancorata alle leggi fisiche e alle normative indu-

⁹¹ Dassault Systèmes, *IA e innovazione protagonista al 3DEXPERIENCE World 2026*, Tecnelab, 2026.

Figura 2.5: Presentazione dei tre Compagni Virtuali di Dassault al 3DEXPERIENCE World 2026. Immagine da "New Solidworks AI agents added at 3DExperience World", DEVELOP3D.

striali, denominata Scientific AI. In occasione del 3DEXPERIENCE World 2026, l'azienda ha presentato tre "Compagni Virtuali" specializzati: Aura, preposto alla gestione del contesto progettuale e dei requisiti; Leo, dedicato al ragionamento ingegneristico e alla logica di produzione; Marie, garante del rigore scientifico e della conformità normativa in settori quali l'aerospaziale e il medicale. Questi assistenti non operano come "scatole nere", bensì collaborano con l'utente assicurando che ogni suggerimento generativo sia fabbricabile, sicuro e conforme agli standard globali⁹¹. A differenza dell'approccio statunitense, orientato alla velocità, l'Europa persegue la "verità scientifica": i compagni virtuali fungono da "filtri di realtà", garantendo che il design sia conforme alle leggi dei



materiali e alle stringenti normative dell'Unione Europea.

Siemens, in occasione del DAC 2025, ha potenziato il proprio portafoglio EDA (Electronic Design Automation) con soluzioni di AI generativa e agentica per la progettazione di semiconduttori e circuiti stampati (PCB). L'approccio si distingue per l'enfasi sulla sicurezza e sulla sovranità tecnologica: gli ingegneri possono avvalersi di modelli avanzati, come NVIDIA Nemotron, all'interno di ambienti protetti, preservando il know-how aziendale da esposizioni indesiderate. L'AI interviene nell'ottimizzazione dei layout elettrici complessi, riducendo drasticamente gli errori di integrità del segnale e accelerando la verifica dei sistemi chip-to-system⁹².

⁹² Siemens, *Siemens EDA AI DAC 2025*, 2025.

Schaeffler, al CES 2026, ha dimostrato come l'AI stia trasformando la robotica umanoide attraverso lo sviluppo di attuatori e sensori "intelligenti". Il design non riguarda più soltanto la forma esterna, ma la capacità del robot di muoversi con una qualità prossima a quella umana e di

⁹³ Schaeffler, *Schaeffler at CES 2026: Technologies That Move the Future*, 2026.

adattarsi autonomamente ad ambienti produttivi flessibili. Integrando la meccanica di precisione con il monitoraggio digitale, Schaeffler realizza sistemi che "percepiscono" e "reagiscono", rendendo i processi manifatturieri non soltanto automatizzati, ma autonomi⁹³. Si assiste quindi all'umanizzazione della meccanica: l'AI consente ai componenti fisici di "apprendere" dall'ambiente, inaugurando un design genuinamente adattivo.

Cina: l'AI come infrastruttura totale

In Cina l'Intelligenza Artificiale è concepita come un'infrastruttura totale: un tessuto connettivo che unifica design, produzione e marketing all'interno di modelli multimodali.

⁹⁴ Baidu, *ERNIE 5.0: Unified Multimodal Model*, 2026.

L'evoluzione tecnologica promossa da Baidu, culminata in un modello multimodale unificato da 2,4 trilioni di parametri, ha definitivamente abbattuto le barriere tra testo, immagine, video e audio. Nel campo del design, questo progresso si traduce nella nascita di Agenti Intelligenti dotati di un avanzato ragionamento cross-modale: oggi l'intelligenza artificiale è in grado di interpretare un documento tecnico, generare il relativo modello 3D e produrre simultaneamente una simulazione video del suo funzionamento. In parallelo, Baidu sta investendo massicciamente nei Digital Humans per lo sviluppo di interfacce umane virtuali. Queste figure ridefiniscono il design dell'interazione, trasformando l'identità digitale in un vero e proprio prodotto industriale scalabile⁹⁴. La visione strategica cinese punta, dunque, a un'integrazione totale in cui l'AI agisce da collante tra file CAD, simulazioni e identità dell'utente: il design evolve così in un flusso continuo di dati che connette il concept virtuale direttamente alla gestione produttiva della fabbrica.

⁹⁵ Tech in Asia, *DeepSeek and Tencent researchers launch AI tool for CAD design*, 2026.

A rafforzare questa spinta verso l'autonomia tecnologica è la collaborazione tra Tencent e la startup DeepSeek (marzo 2026), che ha portato alla creazione di un innovativo strumento AI per il design CAD. L'obiettivo dichiarato è sfidare il dominio dei software occidentali attraverso una modellazione 3D più agile e nativamente integrata con algoritmi di deep learning. Questo progetto rappresenta un caso emblematico della strategia cinese volta allo sviluppo di soluzioni software "sovrane", progettate per garantire ai designer nazionali strumenti capaci di massimizzare la velocità di iterazione nel design di prodotto⁹⁵. La tendenza emergente è la rottura dei monopoli esistenti: l'obiettivo è sostituire il "peso" dei software legacy con tool CAD leggeri e basati sull'AI, capaci di acce-

lerare i tempi di progettazione e produzione ben oltre le attuali possibilità dei competitor occidentali.

Corea del Sud: integrazione verticale

La Corea del Sud rappresenta un caso singolare nel panorama globale, dove l'intelligenza artificiale viene applicata non tanto alla fase concettuale del design, quanto all'intero ciclo produttivo, dalla progettazione dei semiconduttori alla manifattura su scala industriale. Samsung, nell'ambito di un piano di investimenti quinquennale da circa 310 miliardi di dollari, ha avviato nel 2025 il progetto AI Megafactory in collaborazione con NVIDIA. Sfruttando la piattaforma Omniverse, l'iniziativa prevede la creazione di digital twin completi degli stabilimenti di fabbricazione a livello globale, con l'obiettivo di abilitare la manutenzione predittiva, ottimizzare le operazioni dal design iniziale al controllo qualità e ridurre significativa-



Figura 2.6: (Da sinistra) Ronnie Vasishtha, vicepresidente senior di NVIDIA per il settore delle telecomunicazioni, e June Moon, vicepresidente esecutivo e responsabile della ricerca e sviluppo di Samsung Electronics. Immagine da "Samsung Advances AI in Mobile Networks With NVIDIA",

mente i tempi di ramp-up dei nuovi impianti.

L'approccio sudcoreano si distingue per una spiccata verticalità: Samsung controlla infatti l'intera catena del valore, dalla progettazione e fabbricazione dei propri chip di memoria HBM4 fino alla gestione logistica delle fabbriche. Questa integrazione strutturale conferisce un vantaggio competitivo nell'implementazione dell'AI manifatturiera, poiché consente di ottimizzare simultaneamente il design del prodotto, il processo produttivo e l'infrastruttura di fabbrica. La strategia geografica rimane fortemente concentrata nel territorio nazionale, dove nel 2023 è stato annunciato un investimento di 230 mi-

⁹⁶ Enki AI, *Samsung AI in Manufacturing 2025: How the NVIDIA Partnership Drives Industrial Efficiency*, 2026.

liardi di dollari in vent'anni per la costruzione del più grande polo mondiale di semiconduttori, destinato a fungere da base per queste nuove strutture AI-driven⁹⁶. Per il design, la lezione sudcoreana appare dunque la seguente: l'intelligenza artificiale non è soltanto uno strumento per generare forme, ma un sistema nervoso digitale che permea l'intero ecosistema, legando indissolubilmente il progetto alla produzione e alla consegna.

2.3 Il panorama accademico e la formazione del designer nel 2026

Nel 2026 il mondo accademico si trova in una posizione ambivalente, sospeso tra l'entusiasmo per le potenzialità degli strumenti generativi e la necessità di preservare il pensiero critico, l'autonomia intellettuale e la responsabilità etica dei propri studenti. Le tendenze analizzate nel capitolo precedente pongono interrogativi urgenti ai percorsi formativi. Se il valore si sposta dall'esecuzione all'ideazione, dalla tecnica al giudizio, dalla competenza strumentale alla visione sistemica, come devono evolversi i curricula?

La sfida non è semplicemente aggiungere corsi sull'AI agli ordinamenti esistenti, ma ripensare le competenze fondamentali che definiscono il designer. La capacità di lavorare con l'incertezza e l'ambiguità, di integrare saperi eterogenei, di esercitare un pensiero critico sugli output algoritmici, di comprendere le implicazioni etiche e sociali delle scelte progettuali: queste diventano competenze core, non complementari.

Al tempo stesso, la formazione tecnica non può essere abbandonata. Comprendere come funzionano i modelli generativi, quali sono i loro limiti e bias, come si struttura un problema perché l'AI possa affrontarlo efficacemente: tutto questo richiede una base solida di conoscenze che va oltre la superficie dell'interfaccia utente.

Il capitolo che segue esplora come le istituzioni formative stiano affrontando questa transizione, analizzando le strategie adottate per preparare i designer a un futuro in cui la collaborazione con l'intelligenza artificiale non sarà più un'opzione, ma una condizione strutturale della pratica professionale.

2.3.1 I poli internazionali della formazione: USA, Cina e Regno Unito

Le grandi istituzioni accademiche mondiali stanno agendo come laboratori d'avanguardia per la definizione dei futuri curricula del design. Osservare i modelli formativi dei principali poli internazionali permette di individuare le competenze che verranno richieste alla prossima generazione di professionisti.

USA: ricerca e sperimentazione didattica

Il panorama accademico statunitense rappresenta il punto di riferimento mondiale per la ricerca sull'AI applicata al design, sia per la densità delle istituzioni coinvolte sia per la profondità degli approcci elaborati. Al centro di questo ecosistema si trova lo Stanford Institute for Human-Centered Artificial Intelligence (HAI), il caso più emblematico di ricerca interdisciplinare sull'AI centrata sui valori umani. Fondato nel 2019, il centro integra competenze provenienti da informatica, filosofia, scienze politiche, design e medicina, spostando la questione fondamentale da "cosa può fare l'AI" a "come deve essere progettata". La missione del portafoglio formativo di HAI è fornire ai leader le conoscenze necessarie in materia di AI, aiutandoli ad affrontare le implicazioni sociali ed etiche nel settore pubblico, privato e dell'istruzione⁹⁷.

⁹⁷ Stanford University, *Institute for Human-Centered Artificial Intelligence (HAI)*, 2026.

L'AI Index Report 2025, prodotto dall'istituto, conferma che la presenza dell'intelligenza artificiale nei sistemi scientifici, tecnologici ed economici su scala globale ha raggiunto una fase di accelerazione senza precedenti, con una crescita significativa di pubblicazioni, brevetti e investimenti⁹⁸. Gli esperti di Stanford prevedono che nel 2026 l'attenzione si sposterà dalla promessa alla valutazione rigorosa dell'impatto reale dell'AI, ponendo domande come "quanto bene funziona, a quale costo e per chi?" anziché "può farlo?". Tra le previsioni principali figurano un maggiore realismo rispetto all'hype, con consapevolezza dei limiti dell'AI generica, tra cui costi ambientali e rischi di deskilling, e lo sviluppo di un'interazione umano-AI più centrata sull'essere umano, con sistemi progettati

⁹⁸ Stanford HAI, *AI Index Report 2025*, Stanford University, 2025.

⁹⁹ Stanford University, *Stanford AI experts predict what will happen in 2026, 2025.*

per potenziare le capacità umane nel lungo periodo invece di ottimizzare solo l'engagement o l'immediatezza⁹⁹. La previsione di Stanford sul deskilling è particolarmente rilevante per il design: se l'AI assume le fasi più tecniche della progettazione, il rischio è che i futuri designer perdano la capacità di eseguire manualmente ciò che delegano alla macchina, indebolendo la comprensione profonda dei processi che dovrebbero governare.

Figura 2.7: "Moon-faced", un'opera di Morehshin Allahyari ispirata alla storia del genere, dell'orientamento sessuale e del concetto di bellezza in Iran durante la dinastia Qajar, vissuta 350 anni fa¹⁰³.



Di particolare interesse per la formazione del designer è la ricerca condotta a Stanford sull'addestramento dell'AI come "collaboratore creativo". I ricercatori stanno lavorando per superare lo schema "prompt-output", addestrando i modelli a comprendere il processo creativo e a offrire suggerimenti che stimolino il pensiero laterale anziché fornire immediatamente la soluzione finale. Questa ricerca suggerisce che il futuro designer non dovrà soltanto "usare" l'AI, ma sapere come orientarla affinché diventi un partner critico capace di mettere in discussione le sue stesse assunzioni progettuali¹⁰⁰. La conferenza "Norms in the Age of Intelligent Machines: Bodies, Knowledge, Governmentality", organizzata dall'università nel

¹⁰⁰ Stanford HAI, *Stanford scholars train generative AI to be better creative collaborators*, HAI News, 2025.

¹⁰¹ Stanford University, *Norms in the Age of Intelligent Machines: Bodies, Knowledge, Governmentality*, 2025.

¹⁰² Stanford University, *How artists can 'interrupt' artificial intelligence*, 2025.

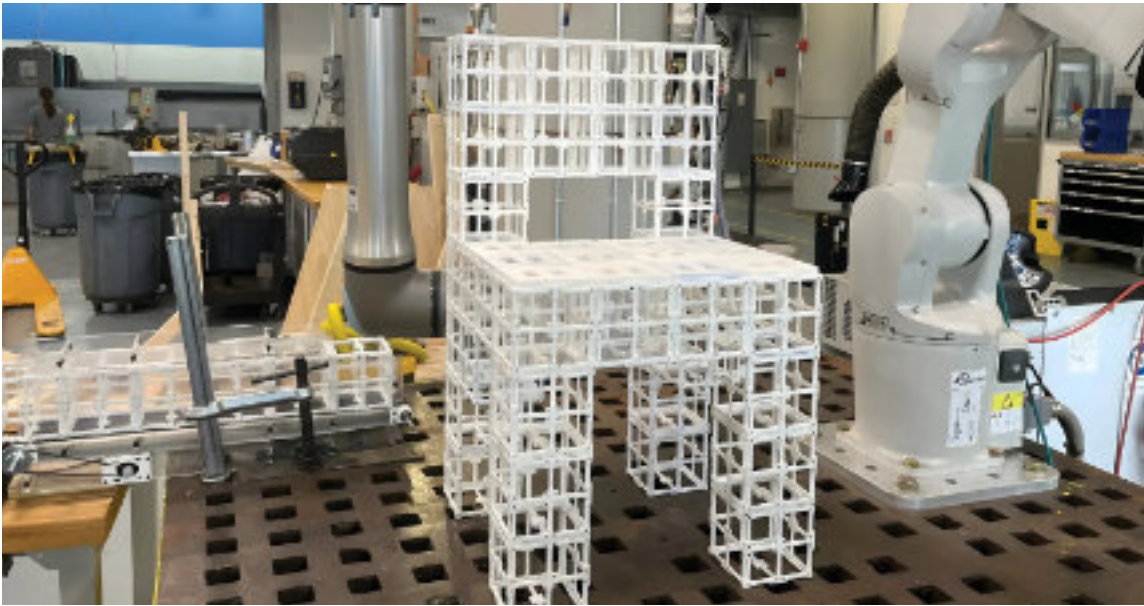
¹⁰³ Stanford HAI, *From Privacy to 'Glass Box' AI, Stanford Students Are Targeting Real-World Problems*, HAI News, 2025.

¹⁰⁴ MIT News, *Robot, make me a chair: sistema di progettazione AI-driven*, 2025.

dicembre 2025, ha investigato come i sistemi di machine learning e di AI generativa non si limitino a seguire norme, ma trasformino le strutture normative esistenti e ne producano di nuove, ridefinendo le norme epistemiche, i criteri di verità e la rappresentazione culturale dei corpi¹⁰¹. L'artista e docente Morehshin Allahyari ha esemplificato questa dinamica dimostrando come gli strumenti AI possano essere usati per criticare e "interrompere" i bias culturali integrati nei modelli: utilizzando DALL-E e Midjourney per ritratti della dinastia Qajar, Allahyari ha evidenziato come i bias dei modelli richiedano interventi consapevoli per sfidare le rappresentazioni culturali standardizzate, dichiarando che *"it's always about understanding what kind of digital libraries these tools are pulling information from."* Aggiunge: *"I often talk in my classes about introducing your own materials, creating a type of 'decentralized' library, one that is not from the dominant culture"*¹⁰². Gli studenti di Stanford, dal canto loro, si concentrano su problemi del mondo reale: ad esempio, presso lo Stanford AI Lab (SAIL), il dottorando Ken Ziyu Liu studia l'intersezione tra modelli di base, dati e privacy degli utenti, con la missione di creare strumenti migliori per la privacy che proteggano gli esseri umani in un mondo guidato dall'AI¹⁰³.

Il Massachusetts Institute of Technology (MIT) ha contribuito alla ridefinizione del design come disciplina computazionale attraverso due iniziative complementari. Nel dicembre 2025 un team di ricercatori ha presentato un sistema robotico guidato da intelligenza artificiale in grado di progettare e costruire oggetti fisici, come una sedia, a partire da una semplice descrizione testuale. Il sistema integra due modelli generativi: il primo costruisce una rappresentazione 3D dell'oggetto a partire dal testo, il secondo determina come assemblare i componenti modulari in base alla funzione dell'oggetto. Una volta generato il modello 3D, un robot lo assembla automaticamente usando pezzi prefabbricati, mentre gli utenti possono interagire con il sistema e dare feedback per modificare il design, mantenendo un ruolo attivo nel processo creativo. Come ha dichiarato Alex Kyaw, autore principale dello studio, *"sooner or later, we want to be able to communicate and talk to a robot and AI system the same way we talk to each other to make things together. Our system is a first step toward enabling that future"*¹⁰⁴.

La portata di questa ricerca per il design è significativa: il progettista non opera più soltanto nell'ambiente bidimensionale dello schermo o nel modello tridimensionale del software, ma in una catena continua che va dalla pa-



¹⁰⁵ MIT Media Lab, *DESAI25: Workshop on Generative AI for Design, 2025*.

Figura 2.8: In seguito alla richiesta «Costruiscimi una sedia» e al feedback «Voglio dei pannelli sulla seduta», il robot assembla una sedia e posiziona i pannelli secondo le indicazioni dell'utente¹⁰⁴.

rola all'oggetto fisico, mediata da un sistema che traduce intenti linguistici in geometrie costruibili e le assembla autonomamente. In parallelo, il MIT Media Lab ha ospitato DESAI25, il Workshop di Generative AI for Design, focalizzato sulla collaborazione tra ricercatori, designer e architetti per definire nuovi flussi di lavoro. L'evento ha esplorato come l'AI possa supportare l'ideazione, la prototipazione e persino la fabbricazione, dimostrando come il design stia diventando una disciplina computazionale il cui punto centrale è la "Creative Intelligence": l'unione tra l'intuito umano e l'esplorazione algoritmica di migliaia di varianti progettuali in tempo reale¹⁰⁵.

¹⁰⁶ University of California Berkeley, *As chatbots get smarter, humans' unique language abilities are becoming less special, 2025*.

All'Università della California, Berkeley, la ricerca ha messo in discussione alcuni assunti fondamentali sul primato delle competenze umane: uno studio del luglio 2025 ha dimostrato che i chatbot avanzati basati su modelli linguistici stanno sviluppando abilità metalinguistiche, come la capacità di analizzare e riflettere sulla struttura del linguaggio in modo analogo a un linguista umano, mettendo in dubbio l'idea che certe competenze linguistiche siano esclusive degli esseri umani¹⁰⁶. La scoperta ha implicazioni dirette per il design: se le macchine sanno analizzare e manipolare strutture simboliche complesse, il vantaggio competitivo del designer non potrà più risiedere nella sola padronanza tecnica o linguistica, ma dovrà spostarsi verso competenze meno replicabili come il giudizio estetico, l'empatia verso l'utente e la capacità di attribuire significato. Gli esperti di Berkeley hanno inoltre identificato undici tendenze da monito-

¹⁰⁷ University of California Berkeley, *11 things UC Berkeley AI experts are watching for in 2026*, 2026.

¹⁰⁸ University of California San Diego, *Human-centered eXtended Intelligence Research Lab (HXI)*, 2026.

¹⁰⁹ Tsinghua University, *Tsinghua University releases comprehensive guiding principles for AI use in education*, 2025.

rare nel 2026, tra cui il rischio di una bolla AI, l'erosione della fiducia causata da deepfake e media generati, la crescente importanza della privacy nei chatbot e i progressi verso sistemi che puntino a comprendere e scoprire la realtà anziché limitarsi a massimizzare punteggi di performance¹⁰⁷. Queste tendenze delineano un contesto in cui la formazione del designer non può prescindere da una comprensione delle dinamiche sociali, politiche e percettive che l'AI genera nel sistema comunicativo globale. Presso la University of California, San Diego, lo Human-centered eXtended Intelligence Research Lab (HXI), inserito nel Department of Computer Science & Engineering, si occupa di progettare, sviluppare e valutare tecnologie interattive combinando Intelligenza Artificiale, Extended Reality (XR) e design centrato sull'uomo. I progetti del laboratorio si collocano all'incrocio tra Human-Computer Interaction (HCI), AI, XR e Human-Centered Design (HCD), rappresentando un esempio concreto di come il design stia evolvendo verso una disciplina ibrida in cui la competenza progettuale si intreccia con la scienza computazionale e la comprensione dei fattori umani¹⁰⁸.

Cina: framework adattivo e artigianato aumentato

Nel panorama accademico asiatico, la Cina occupa una posizione di assoluto rilievo, caratterizzata da un'integrazione particolarmente stretta tra indirizzo governativo, produzione scientifica e applicazione industriale dell'AI. La Tsinghua University ha definito principi guida per l'utilizzo dell'Intelligenza Artificiale nell'istruzione e nella ricerca, concepiti come un "sistema vivente" destinato a evolversi insieme alla tecnologia, con l'obiettivo di bilanciare apertura all'innovazione e responsabilità etica. I principi generali stabiliscono che l'AI deve restare uno strumento ausiliario, con studenti e docenti come protagonisti dell'insegnamento e dell'apprendimento. Inoltre, si incoraggia la vigilanza sugli errori generati dall'AI e la verifica incrociata per evitare un'eccessiva dipendenza. Le norme per le tesi e la ricerca avanzata vietano l'uso dell'AI per ghostwriting, plagio o fabbricazione di contenuti, richiedendo ai supervisori di fornire indicazioni chiare e garanzia di originalità e integrità accademica¹⁰⁹. Il modello della Tsinghua è, appunto, un sistema vivente, che si distingue per la sua natura dinamica: non una regolamentazione statica, ma un framework adattivo che riconosce la velocità evolutiva della tecnologia e la necessità di rivedere periodicamente le proprie regole. Un principio di "meta-governance" che riflette la maturità istituzionale di un ateneo che opera al centro di una delle

strategie nazionali sull'AI più ambiziose al mondo.

La governance accademica dell'AI in Cina è strettamente connessa alla strategia governativa: la collaborazione tra il governo cinese e la Tsinghua University ha portato alla creazione del CnAISDA, un organismo istituzionale dedicato alla sicurezza dei modelli di frontiera. L'interesse cinese per i rischi catastrofici legati all'AI, dalle minacce biologiche e chimiche alla perdita di controllo dei sistemi, è in crescita, e nonostante l'incertezza scientifica su questi rischi, il governo cinese ha iniziato a prendere sul serio queste preoccupazioni, allineandosi parzialmente alle discussioni globali sulla sicurezza dei modelli di frontiera¹¹⁰. Per il designer, questa dinamica ha un'implicazione diretta: i confini etici e di sicurezza dell'AI non sono soltanto una questione accademica, ma una dimensione istituzionale e geopolitica che condiziona concretamente gli strumenti disponibili, i dati accessibili e le modalità operative della progettazione.

¹¹⁰ Carnegie Endowment for International Peace, *Making Sense of China's AISI Equivalent: Institutional Design and Key Actors*, 2025.

Sul piano della ricerca applicata al design, la Tongji University di Shanghai ospita il Design AI Lab, un laboratorio che integra intelligenza artificiale, dati e design per reinventare il processo creativo e connettere il design alla nuova economia digitale. Il Lab sviluppa strumenti intelligenti per il graphic design basati su dati di layout, crea dataset, esplora applicazioni di AI per l'artigianato tradizionale e riflette sul rapporto tra computazione e immaginazione umana. Oltre alla ricerca tecnica, il laboratorio sostiene programmi educativi e iniziative di divulgazione per formare talenti e stimolare una comprensione critica dell'AI nel mondo creativo, con la visione di un ecosistema in cui l'AI non solo automatizza compiti ma amplifica e distribuisce valore nell'intero campo del design¹¹¹. L'aspetto più originale del Design AI Lab è la connessione tra AI e artigianato tradizionale: anziché opporre tecnologia e tradizione, il laboratorio esplora come gli algoritmi generativi possano preservare, reinterpretare e diffondere le tecniche artigianali cinesi, trasformando il patrimonio culturale in dataset e il dataset in nuova progettazione.

¹¹¹ Tongji University, *Design AI Lab*, 2026.

Regno Unito: design come leadership

Il contesto accademico britannico si distingue per un approccio particolarmente elaborato alla governance istituzionale dell'AI, che va oltre la semplice adozione degli strumenti per costruire veri e propri ecosistemi culturali di uso consapevole. La University of Oxford, nel settembre 2025, è divenuta la prima università del Regno Unito

¹¹² University of Oxford, *Oxford becomes first UK university to offer ChatGPT Edu to all staff and students, 2025.*

a offrire accesso gratuito a ChatGPT Edu a tutti i propri studenti e membri del personale a partire dall'anno accademico 2025-26. L'iniziativa non si limita alla fornitura dello strumento, ma è accompagnata da un'infrastruttura di governance dedicata: corsi di formazione, webinar e risorse didattiche per promuovere un utilizzo responsabile e sicuro della tecnologia, con il supporto di un AI Competency Centre, una rete di AI Ambassadors e un AI Governance Group per supervisionare l'adozione¹¹². Il modello di Oxford illustra un approccio che potremmo definire "adozione strutturata": non la mera distribuzione di uno strumento, ma la costruzione di un ecosistema istituzionale capace di orientarne l'uso.

¹¹³ University of Cambridge, *Using Generative AI: Guidance for Students, Cambridge Blended Learning, 2025.*

La University of Cambridge ha pubblicato linee guida specifiche per gli studenti sull'uso dell'AI generativa, stabilendo che gli strumenti possono essere utilizzati per brainstorming, ricerca, sintesi e analisi dei dati, a condizione che gli studenti verifichino sempre e si assumano la responsabilità delle informazioni integrate nel proprio lavoro. Le linee guida raccomandano di consultare il proprio supervisore per chiarire le modalità d'uso appropriate nel contesto del corso e sottolineano come sia essenziale evitare la dipendenza totale dall'AI, considerandone i limiti in termini di coerenza, accuratezza e bias¹¹³. Al contempo, ricercatori di Cambridge stanno studiando l'integrazione dell'AI nell'istruzione human-centrica: come affermato nell'ambito di questa ricerca, "to explore the potential benefits, we have to start with the opportunity or the problem we're trying to solve, rather than thinking about the exciting things that AI can do. And yet, when students move on to work, they will be using AI in their jobs, so AI literacy needs to be embedded in the curriculum"¹¹⁴. La posizione di Cambridge introduce nel dibattito un principio pragmatico di grande rilevanza: l'AI literacy non è una competenza accessoria, ma un requisito trasversale che deve permeare l'intero curriculum formativo.

¹¹⁴ University of Cambridge, *Ricerca sull'integrazione dell'AI nell'istruzione human-centrica: AI literacy e curriculum, 2025.*

Una ricerca dello stesso ateneo, pubblicata nel dicembre 2025, ha inoltre dimostrato che l'AI può essere un efficace partner creativo a condizione che vengano fornite istruzioni chiare su come sviluppare le idee insieme: i ricercatori sono rimasti sorpresi nel constatare che le coppie umano-AI non miglioravano naturalmente attraverso la collaborazione ripetuta, ma soltanto quando veniva introdotto un intervento deliberato che istruiva i partecipanti a sviluppare le idee congiuntamente, concentrandosi sullo scambio di feedback e sul perfezionamento delle idee esistenti piuttosto che sulla generazio-

¹¹⁵ University of Cambridge, *Can AI be a good creative partner?*, 2025.

ne continua di idee nuove¹¹⁵. La ricerca di Cambridge conferma e rafforza quella di Stanford sulla collaborazione creativa: la creatività non risiede nel modello generativo né nel designer isolato, ma nella qualità dell'interazione tra i due.

¹¹⁶ Royal College of Art, *Leading with a Design Mindset*, 2026.

Il Royal College of Art (RCA) di Londra propone invece una visione del design come leadership strategica in un mondo saturo di AI. Il programma executive "Leading with a Design Mindset" sottolinea che, in un'era di automazione, il "Design Mindset" è la competenza umana più preziosa, e si focalizza su come usare l'AI per affrontare problemi complessi senza perdere l'orientamento etico e umano¹¹⁶. Lo Strategic Plan 2026-2030 dell'RCA dichiara ufficialmente che l'AI sarà integrata in ogni aspetto dell'istruzione e della ricerca, con un impegno ferreo verso la giustizia sociale e la sostenibilità, e prevede la creazione di nuovi laboratori dedicati all'AI¹¹⁷. La proposta del RCA è radicale: il designer del futuro non è un tecnico che usa strumenti intelligenti, ma un leader capace di orientare la tecnologia verso finalità umane. Il "Design Mindset" non è una competenza tra le altre, ma la meta-competenza che consente di governare tutte le altre, inclusa l'AI.

¹¹⁷ Royal College of Art, *Strategic Plan 2026–2030*, 2026.

2.3.2 Altre esperienze nel mondo

Oltre ai poli dominanti, emergono numerosi ecosistemi formativi che sperimentano approcci ibridi e localizzati all'intelligenza artificiale. Queste esperienze offrono una prospettiva preziosa sulla diversità delle risposte educative a una sfida tecnologica che, pur essendo globale, impatta su contesti culturali molto differenti.

Europa continentale: tra ambiente costruito, Hybrid Intelligence e AI responsabile

Al di là dei protagonisti anglosassoni, il continente europeo esprime un panorama accademico ricco di iniziative originali, distribuite soprattutto tra Paesi Bassi, Svizzera, Germania e Finlandia, e accomunate da un'attenzione particolare al rapporto tra AI, progettazione dell'ambiente costruito e responsabilità sistemica.

La Technische Universiteit Delft (TU Delft), nei Paesi Bassi, ha strutturato la propria risposta alla trasformazione digitale attraverso un programma di AI Labs dedicato al design. Il Design at Scale Lab (D@S Lab) si focalizza sul concetto di "Hybrid Intelligence", ovvero la combinazione sinergica tra intelligenza artificiale e intelligenza umana per affrontare problemi sociali complessi su larga

¹¹⁸ TU Delft, *Design at Scale Lab — Human-AI Collaboration in Design for Social Good*, TU Delft AI Labs programme, 2026.

¹¹⁹ TU Delft, *AiDAPT — AI for a sustainable and resilient built environment*, TU Delft AI Labs programme, 2026.

¹²⁰ ETH Zurich, *Design++ — Center for Augmented Computational Design in Architecture, Engineering and Construction*, 2026.

¹²¹ Technical University of Munich, *TUM GNI International Symposium 2026: Artificial Intelligence for the Built World*, Georg Nemetschek Institute, 2026.

scala. La visione del laboratorio è che designer, esperti e portatori di interesse collaborino con l'AI per generare soluzioni che nessuna delle due intelligenze, presa singolarmente, sarebbe in grado di produrre: i corsi attivati nell'ambito di questo approccio, tra cui Machine Learning for Design, Advanced Machine Learning for Design, Crowd Computing for Design e Artificial Intelligence and Society, riflettono la necessità di formare una generazione di designer capaci di operare all'interno di sistemi ibridi¹¹⁸. In parallelo, il laboratorio AiDAPT (AI for a Sustainable and Resilient Built Environment) si concentra sul potenziamento dei processi decisionali di architetti e ingegneri attraverso l'AI, a diverse scale e fasi del ciclo di vita dell'ambiente costruito. Il laboratorio impiega computer vision, data science e ottimizzazione decisionale per sviluppare framework di deep learning, reinforcement learning e quantificazione dell'incertezza, con l'obiettivo di rendere l'ambiente costruito più sostenibile e resiliente¹¹⁹.

L'ETH Zurich, in Svizzera, ha istituito il Centro Design++ (Center for Augmented Computational Design in Architecture, Engineering and Construction), che esplora in che modo le tecnologie computazionali avanzate di AI, machine learning ed extended reality possano aumentare le capacità progettuali nel settore dell'ambiente costruito. Il centro opera attraverso due laboratori centrali, l'Immersive Design Lab e il Large-scale Virtualization and Modelling Lab, e dispone di un programma di fellowship post-dottorale. Quest'anno una serie di seminari ha già affrontato la domanda "Can AI think Architecture?", mentre il Future of Construction Symposium 2026, in programma a maggio, si propone di riunire i principali attori della ricerca e dell'industria per definire le traiettorie future dell'integrazione tra AI e progettazione¹²⁰.

La Technische Universität München (TUM), in Germania, attraverso il Georg Nemetschek Institute (TUM GNI), si è affermata come centro di eccellenza per la ricerca, la didattica e il trasferimento di conoscenza nel campo dell'AI applicata all'ambiente costruito. Il TUM GNI International Symposium 2026, "Artificial Intelligence for the Built World", in programma a Monaco dal 18 al 20 maggio 2026, riunisce ricercatori da tutto il mondo attorno a sette aree tematiche che spaziano dalle applicazioni avanzate di machine learning ai gemelli digitali probabilistici, dall'ingegneria dei dati alle interfacce intelligenti, dalla progettazione sostenibile alla robotica¹²¹. Il symposium rappresenta uno degli appuntamenti più rilevanti del 2026 per chi studia o pratica il design in ambienti ad alta

integrazione tecnologica.

L'Aalto University di Helsinki, in Finlandia, ha lanciato l'ELLIS Institute Finland, un hub di ricerca in AI e machine learning che nel 2025 ha inaugurato anche un nuovo programma di laurea magistrale multidisciplinare, testimoniando la volontà dell'ateneo di formare una nuova generazione di ricercatori capaci di lavorare all'intersezione tra discipline tecniche e scienze umane¹²². Sul versante della formazione continua, Aalto Executive Education offre il programma "Generative AI in Creating Visual Content", rivolto a professionisti creativi che desiderano integrare le competenze generative nella propria pratica professionale¹²³.

¹²² Aalto University, *Aalto University in 2025: strong results in research, education and innovation activities*, 2026.

¹²³ Aalto University / Aalto Executive Education, *Generative AI in Creating Visual Content*, Corso di Formazione Continua, 2025.

Corea del Sud: un laboratorio nazionale per l'AI

La Corea del Sud rappresenta uno dei contesti più dinamici al mondo per la formazione accademica in AI, con una strategia nazionale che coinvolge direttamente le principali università del paese in un progetto di posizionamento globale.

Il Korea Advanced Institute of Science and Technology (KAIST) ha approvato, nel dicembre 2025, la creazione del primo College of AI autonomo in Corea, con l'avvio delle iscrizioni previsto nel 2026. Il nuovo college, che segna la prima volta in cui una università coreana eleva l'AI a unità di livello universitario indipendente, prevede l'ammissione di 300 studenti all'anno e si articola in quattro dipartimenti che coprono l'intera catena del valore dell'AI: il dipartimento di AI Computing (fondamenti teorici, algoritmi, AI generativa e multimodale), il dipartimento di AI Systems (hardware per AI, semiconduttori, comunicazioni ad alta velocità), il dipartimento di AI Transformation (applicazioni industriali e sociali in manifattura, biotecnologia, sostenibilità) e il dipartimento di AI Futures (etica, policy, economia e governance)¹²⁴. Quest'ultimo dipartimento è particolarmente significativo per il design: riconoscere l'etica e la governance come una delle quattro colonne portanti di un college di AI equivale ad affermare che la dimensione progettuale e normativa dell'AI non è un corollario, ma un asse fondante della disciplina. In parallelo, il KAIST ospita il National AI Research Lab (NAIRL), il più grande consorzio di ricerca sull'AI in Corea, che coinvolge 45 professori e oltre 150 ricercatori di quattro università (KAIST, Korea University, Yonsei University e POSTECH) in cooperazione con 12 aziende nazionali e 14 istituzioni di ricerca internazionali¹²⁵.

¹²⁴ The Investor / Herald Corporation, *South Korea's KAIST to launch stand-alone AI college in 2026 amid global talent race*, 2025.

¹²⁵ KAIST News Center, *Korea's AI Research Hub: Global AI Frontier Symposium 2025*, 2025.

¹²⁶ K-Moonshot, *Seoul National University (SNU): AI Research & Strategy*, 2026.

La Seoul National University (SNU) ha fissato un obiettivo esplicito: entrare tra le prime dieci università al mondo per la ricerca in AI entro il 2026. Con circa settanta professori impegnati nella ricerca sull'AI distribuiti tra le sedici facoltà dell'ateneo, tra cui scienze naturali, scienze sociali, ingegneria, giurisprudenza, medicina e scienze umane, la SNU è l'unica istituzione in Corea in grado di affrontare la complessità dell'AI nella sua interezza. L'AI Native Campus, iniziativa centrale della strategia dell'ateneo, si propone di integrare l'AI generativa in ogni aspetto della vita universitaria: nella didattica attraverso strumenti di apprendimento personalizzato e sistemi di valutazione automatica; nella ricerca attraverso l'accesso universale a risorse computazionali e API di modelli AI; nell'amministrazione attraverso sistemi di gestione del campus, analisi istituzionale e ottimizzazione delle risorse¹²⁶. L'ambizione della SNU non è formare soltanto specialisti di AI, ma produrre laureati AI-literate indipendentemente dalla disciplina di appartenenza: un principio che rispecchia, a scala nazionale, le medesime preoccupazioni che guidano la riforma curricolare di Stanford, Cambridge e Oxford.

America Latina: infrastrutture emergenti e prime sperimentazioni

Anche il contesto latinoamericano mostra segnali di una crescente consapevolezza istituzionale intorno al tema dell'AI nella formazione superiore, con iniziative che, pur partendo da una dotazione infrastrutturale ancora in sviluppo, stanno costruendo le basi per una partecipazione più competitiva alla ricerca globale.

¹²⁷ StartSe / BNamericas, *USP lança o maior cluster de IA da América Latina e acelera a pesquisa no Brasil*, 2026.

In Brasile, la Universidade de São Paulo (USP) ha inaugurato nel febbraio 2026 il supercomputer Jairu, il più grande cluster di AI in operazione in America Latina, con un investimento di 40 milioni di reais (quasi 7 milioni di euro). Installato presso il Centro de Inteligência Artificial e Aprendizado de Máquina (CIAAM-USP), il sistema è progettato per supportare lo sviluppo di modelli avanzati e attività di ricerca ad alta intensità computazionale. Come ha dichiarato Fábio G. Cozman, coordinatore del CIAAM-USP, "con Jairu abbiamo un'infrastruttura di AI che ci permetterà di sviluppare grandi modelli e approfondire ricerche rilevanti per il contesto brasiliano"¹²⁷. La realizzazione del cluster Jairu non è semplicemente un avanzamento tecnologico: è un segnale di posizionamento strategico che indica la volontà del Brasile di non limitarsi a consumare l'AI prodotta altrove, ma di partecipare attivamente alla sua creazione.

In Argentina, la Facultad de Arquitectura, Diseño y Urbanismo dell'Universidad de Buenos Aires (FADU-UBA), attraverso il suo Instituto de Arte Americano e Investigaciones Estéticas, ha avviato un ciclo di seminari e attività formative dedicati all'uso dell'AI nei processi di ricerca urbano-architettonica. Organizzati in collaborazione con la Universidad Nacional Autónoma de México (UNAM) e con la Maestría en Historia y Crítica de la Arquitectura, el Diseño y el Urbanismo (MAHCADU), i seminari affrontano le modalità con cui l'intelligenza artificiale sta trasformando i processi di ricerca nel campo dell'architettura e del design, con un'attenzione particolare alle dimensioni epistemologiche e critiche della trasformazione in atto. A partire dall'aprile 2026, la facoltà ha attivato anche un Curso de Actualización Profesional su "Tecnologías, género y diseño", che mette in relazione le tecnologie digitali emergenti con le questioni di genere e le pratiche progettuali, segnalando una sensibilità verso le dimensioni culturali e di equità della trasformazione tecnologica¹²⁸.

¹²⁸ Universidad de Buenos Aires — FADU / Instituto de Arte Americano, Seminario abierto IAA — MAHCADU, *Uso de las IA en los Procesos de Investigación Urbano-Arquitectónicos*, 2026.

2.3.3 Il contesto formativo italiano

In Italia, il sistema della formazione sta iniziando a misurarsi con la sfida generativa tra spinte all'innovazione e resistenze strutturali. Analizzare la situazione nazionale significa comprendere come la nostra tradizione progettuale possa dialogare con le nuove tecnologie, cercando un equilibrio tra cultura umanistica e competenza tecnica.

Il dibattito italiano: regolamentare, non proibire

L'ecosistema accademico italiano mostra, nel 2026, una vivacità sorprendente, con iniziative che stanno contribuendo in modo originale al dibattito internazionale sulla formazione del designer nell'era dell'AI. Al centro di questo fermento si collocano alcune questioni fondamentali che investono la natura stessa della didattica universitaria: il rischio della delega cognitiva, i criteri di un uso etico e metodologicamente fondato degli strumenti generativi, e la responsabilità istituzionale degli atenei nel formare cittadini capaci di governare, e non subire, la trasformazione tecnologica.

L'Università Alma Mater di Bologna ha contribuito al dibattito nazionale con una riflessione articolata sul rischio della delega cognitiva. Come evidenziato nel febbraio 2026, con l'AI la posta per i docenti si alza: serve metacognizione, capacità di attivare responsabilità criti-

¹²⁹ Rivoltella, Pier Cesare, *Uno dei temi del dibattito sull'IA in scuola è legato al rischio della delega cognitiva*, Adnkronos, 2026.

¹³⁰ Perin, Lorenzo, *L'intelligenza artificiale come progetto umano. Intervista a Luciano Floridi*, Scienza in rete, 2026.

¹³¹ ISIA Firenze, *Dottorato in Cognitive Design*, 2025.

ca e pensiero creativo, non solo vigilanza sugli strumenti. L'insegnante resta il mediatore che calibra su età e livelli cognitivi, e sul dibattito relativo ai divieti la posizione è netta: confondere educazione con protezione è un errore di categoria, poiché sottrarre uno strumento non sviluppa senso critico, lo aggira¹²⁹. La riflessione bolognese introduce una distinzione concettuale importante: il problema non è lo strumento in sé, ma la qualità della relazione pedagogica che si costruisce intorno ad esso. Senza una mediazione didattica consapevole, l'AI rischia di trasformarsi da amplificatore cognitivo a sostituto cognitivo. Il filosofo Luciano Floridi, figura di riferimento nel dibattito internazionale sull'etica dell'informazione, ha ulteriormente rafforzato questa prospettiva in una recente intervista, affermando che l'intelligenza artificiale non è un destino, un'imposizione o una promessa, ma uno strumento che possiamo attivamente progettare, governare e orientare: l'innovazione è qualcosa che possiamo disegnare, creando modelli, immaginando orizzonti e ponendoci obiettivi da raggiungere¹³⁰. La posizione di Floridi è particolarmente significativa per il designer, poiché trasferisce la responsabilità dall'AI al soggetto che la impiega: progettare con l'AI significa, prima di tutto, progettare l'AI stessa, ovvero definirne i confini d'uso, le finalità e i criteri di qualità.

L'ISIA di Firenze ha avviato un ciclo di dottorato in Cognitive Design con l'obiettivo di formare designer capaci di controllare gli aspetti tecnologico-produttivi interpretando con senso critico la complessità dei sistemi AI. Il programma muove dalla constatazione che la digitalizzazione della produzione industriale è entrata in una nuova fase evolutiva, la cui architettura sta per subire un'importante mutazione alimentata dall'emergere di nuove generazioni di tecnologie hardware e software, da una cultura digitale popolare diffusa e da una verticale riduzione dei costi di tecnologie e servizi¹³¹. Il dottorato in Cognitive Design è significativo perché non si limita a introdurre l'AI come materia di studio, ma la pone al centro di una ridefinizione epistemologica della disciplina stessa: il design non è più soltanto progetto della forma o della funzione, ma progetto della cognizione, ovvero delle modalità con cui esseri umani e sistemi intelligenti costruiscono insieme la comprensione del mondo.

L'Università IUAV di Venezia ha ospitato, dal 8 al 10 maggio 2025 sull'isola di San Servolo, la prima grande conferenza internazionale sull'etica e l'estetica dell'intelligenza artificiale: "Ethics and Aesthetics of Artificial Images" (EA-AI 2025), organizzata in collaborazione con la Venice

¹³² Università IUAV di Venezia, *Ethics and Aesthetics of Artificial Images (EA-AI 2025)*, 2025.

International University, ha riunito oltre 180 studiosi da tutto il mondo. La conferenza ha affrontato domande fondamentali su cosa si considera autentico, creativo, legale e, in ultima analisi, umano, con un legame esplicito alla Biennale di Architettura 2025: *Intelligens*, dedicata all'esplorazione dell'intelligenza naturale, artificiale e collettiva¹³². L'evento ha dimostrato che il dibattito sull'AI nel design non può essere confinato alla dimensione tecnica o normativa, ma deve investire le categorie estetiche fondamentali: il confine tra originale e derivativo, tra autorialità umana e generazione algoritmica, tra rappresentazione e simulazione.

¹³³ Sapienza Università di Roma, Digital Education Hub, *La Sfida della Didattica Universitaria nell'era dell'Intelligenza Artificiale*, 2026.

La Sapienza di Roma, attraverso il Digital Education Hub (DEH) e il Progetto ALMA (Advanced Learning Multimedia Alliance), promosso dal Ministero dell'Università e della Ricerca, ha sviluppato e sperimentato metodologie didattiche innovative basate sull'intelligenza artificiale, sulla didattica immersiva, sulla realtà aumentata e sui remote Labs, con l'obiettivo di trasformare e arricchire l'esperienza formativa universitaria. Il workshop "La Sfida della Didattica Universitaria nell'era dell'Intelligenza Artificiale", tenutosi il 4 e 5 marzo 2026, ha rappresentato un'importante occasione di confronto, condivisione e riflessione sulle iniziative in corso e sulle prospettive future della didattica universitaria¹³³.

¹³⁴ Politecnico di Milano, *Design through Generative AI*, GSOM, 2026.

Il Politecnico di Milano ha attivato un corso specifico sull'integrazione tra Design Thinking e Generative AI, rivolto anche a chi si è già laureato prima del boom dell'AI generativa, con l'obiettivo di integrare le competenze progettuali con la padronanza degli strumenti generativi in ambiti quali la visualizzazione e la prototipazione¹³⁴. Un'iniziativa particolarmente significativa è il Samsung Innovation Campus 2026, promosso nell'ambito del programma *Passion in Action*, un bando che sfida gli studenti a progettare sistemi AI che costruiscano fiducia nei settori salute, benessere e creatività. La sfida si concentra sulla riprogettazione di un servizio digitale esistente in cui l'intelligenza artificiale gioca un ruolo centrale, con l'obiettivo di progettare per la fiducia tra tecnologia, progettisti e utenti finali: progettare per la fiducia significa progettare relazioni, non solo interfacce o algoritmi, considerando chiarezza, trasparenza e fiducia emotiva come elementi cruciali delle esperienze digitali¹³⁵. Il concetto di "Design for Trust" introdotto da questa iniziativa segna un passaggio concettuale rilevante: la fiducia non è un attributo estetico o comunicativo, ma un parametro progettuale misurabile che riguarda la qualità della relazione tra utente, sistema e progettista.

¹³⁵ Politecnico di Milano, *Samsung Innovation Campus 2026 — Passion in Action: Design for Trust*, 2026.

Il Politecnico di Torino: un hub nazionale per l'AI nel design

Tra le istituzioni italiane che hanno risposto con maggiore sistematicità e ambizione alla sfida dell'AI, il Politecnico di Torino rappresenta un caso di particolare interesse. Per questo, e anche in quanto realtà osservabile più da vicino, si ritiene opportuno dedicarvi un breve approfondimento.



Figura 2.9: Elena Baralis, prorettrice del Politecnico di Torino. Immagine da “L'Ingegneria Gestionale a Torino compie trent'anni”, Politecnico di Torino.

Il primo atto significativo del Politecnico in questo senso è stato l'autorizzazione ufficiale dell'utilizzo dell'Intelligenza Artificiale nelle tesi di laurea a partire dal gennaio 2026. La decisione è maturata a seguito di una discussione pubblica critica, innescata da un caso di tesi realizzata in larga parte con strumenti generativi, che ha sollevato interrogativi profondi sull'uso etico e metodologico dell'AI nel percorso accademico. Come ha dichiarato la prorettrice Elena Baralis, *“non possiamo demonizzare l'uso della tecnologia e neanche immaginare di vivere ancora nel 2018. Sarebbe un atteggiamento perdente perché gli studenti la userebbero lo stesso. Ma bisogna capire con*

¹³⁶ *la Repubblica Torino, // Politecnico di Torino permette l'utilizzo dell'AI per le tesi, 2026.*

¹³⁷ *la Repubblica Tecnologia, // Politecnico di Torino tra le eccellenze italiane selezionate da Google per la ricerca sull'IA, 2026.*

¹³⁸ *Politecnico di Torino, Hub AI@PoliTo, 2026.*

¹³⁹ *Politecnico di Torino, Polito Design Workshop 2026 — 26 anni di innovazione, 2026.*

*loro cosa ha senso ed è etico e insegnare non solo a usare questi strumenti di lavoro ma ad avere spirito critico e comprensione per risolvere i problemi. Perché anche se la soluzione è generata da un'intelligenza artificiale bisogna supervisionare*¹³⁶. La posizione del Politecnico esprime un principio che si sta progressivamente affermando nel panorama accademico italiano: non proibire, ma regolamentare, formando gli studenti alla supervisione consapevole.

Sul piano della ricerca, nel febbraio 2026 il Politecnico di Torino è stato selezionato da Google tra le eccellenze italiane (insieme a Politecnico di Milano, Cattolica e Università di Catania) per i programmi di ricerca sull'AI, un riconoscimento che attesta la qualità e la rilevanza internazionale delle attività svolte nei suoi laboratori¹³⁷. Al centro di questo posizionamento si trova l'Hub AI@PoliTo, una struttura che aggrega competenze provenienti da dipartimenti diversi, promuove la crescita dei giovani talenti e introduce il machine learning e l'intelligenza artificiale in un'ampia gamma di applicazioni. L'Hub è un centro di eccellenza riconosciuto a livello internazionale che affronta sfide di ricerca fondamentale, forma una nuova generazione di leader e sviluppa applicazioni pratiche dell'AI in ambiti che richiedono l'integrazione di capacità intelligenti nei sistemi fisici¹³⁸. La selezione da parte di Google non è soltanto un sigillo di qualità: è la conferma che il Politecnico di Torino si sta affermando come nodo cruciale di una rete globale di ricerca, in cui la produzione accademica si traduce in impatto concreto sul mondo dell'industria e della progettazione.

Sul versante didattico, un contributo originale e metodologicamente significativo è venuto dal Polito Design Workshop 2026, che nella sua ventiseiesima edizione ha incluso il workshop "Participatory AI". L'iniziativa ha invitato i partecipanti a una discussione critica con l'AI, includendo la progettazione di AI Boundary Objects, artefatti concettuali che definiscono le regole di interazione tra il designer e il sistema intelligente¹³⁹. L'idea di "Participatory AI" trasla nella didattica la consapevolezza emersa dalla ricerca internazionale: non basta saper usare gli strumenti AI, occorre saper progettare le condizioni della loro collaborazione con l'umano. Coinvolgendo gli studenti nella co-progettazione delle stesse regole di interazione tra designer e sistemi intelligenti, il workshop ha anticipato una delle traiettorie principali della ricerca contemporanea sul design: quella che considera l'AI non come uno strumento da imparare a usare, ma come un interlocutore da imparare a progettare.



2.3.4 Cinque tratti convergenti dal panorama accademico

Figura 2.10: La sede del Politecnico di Torino in Corso Duca degli Abruzzi. Immagine da "Il Politecnico tra le migliori università europee", Politecnico di Torino.

L'analisi trasversale del panorama accademico consente di delineare i contorni del profilo formativo che le università stanno costruendo per il designer del 2026 e oltre, non come figura unitaria e definita, ma come un insieme di competenze convergenti che emergono dalla sintesi delle esperienze analizzate.

Il primo tratto distintivo è la competenza critica nell'uso dell'AI: non la mera padronanza tecnica degli strumenti, ma la capacità di interrogarne i limiti, i bias e le implicazioni etiche. Come emerge dal dibattito dell'Alma Mater di Bologna, servono metacognizione e responsabilità critica e come dimostra la ricerca di Stanford, il designer deve saper addestrare l'AI affinché diventi un partner critico, non un semplice esecutore. La conferenza IUAV sull'etica e l'estetica delle immagini artificiali e il lavoro di Allahyari sui bias culturali nei modelli generativi convergono nel delineare un designer che non si limita a produrre artefatti, ma è consapevole delle strutture normative e culturali che gli strumenti AI incorporano e riproducono.

Il secondo tratto è la capacità di progettare la collaborazione. La ricerca di Cambridge dimostra che l'AI non migliora automaticamente la creatività, ma lo fa solo quando il designer imposta deliberatamente un ciclo di scam-

bio e perfezionamento. La ricerca di Stanford conferma che il futuro del design non risiede nel prompt migliore, ma nella qualità della relazione generativa tra progettista e sistema. Il workshop "Participatory AI" del Politecnico di Torino trasla questa consapevolezza nella didattica, coinvolgendo gli studenti nella co-progettazione delle regole di interazione. Il futuro designer non è un utente dell'AI, ma un vero e proprio architetto di relazioni tra intelligenze umane e artificiali.

Il terzo tratto è la capacità di progettare la fiducia. L'iniziativa del Politecnico di Milano con Samsung evidenzia come il design nell'era dell'AI non riguardi soltanto la forma o la funzione, ma la qualità della relazione tra utente, sistema e progettista. In un contesto in cui, come rilevano gli esperti di Berkeley, l'erosione della fiducia causata da deepfake e media generati è una delle tendenze critiche del 2026, la fiducia diviene un parametro progettuale imprescindibile.

Il quarto tratto è l'interdisciplinarietà: Lo Stanford HAI integra informatica, filosofia, scienze politiche, design e medicina; il Design AI Lab della Tongji University connette dati, design e artigianato tradizionale; l'HXI Lab di UC San Diego fonde AI, Extended Reality e Human-Centered Design; l'ISIA di Firenze ridefinisce il design come "progetto della cognizione"; la Seoul National University estende questa interdisciplinarietà all'intera università, coinvolgendo settanta professori di sedici facoltà. Sono tutte dimostrazioni del fatto che il designer non può più operare all'interno dei confini di una singola disciplina, ma deve padroneggiare un lessico condiviso tra scienze computazionali, scienze umane e pratica progettuale.

Il quinto e ultimo tratto è la visione del design come leadership. Il Royal College of Art propone il "Design Mindset" come meta-competenza per governare la complessità. Allo stesso tempo Floridi ricorda che l'innovazione è qualcosa che possiamo disegnare e il KAIST riconosce l'etica e la governance dell'AI come una delle quattro colonne portanti della formazione. In un contesto in cui le abilità metalinguistiche dei chatbot mettono in discussione il primato di alcune competenze umane e il rischio di deskilling è concreto, la formazione del designer dovrà coltivare proprio ciò che resta irriducibilmente umano: la capacità di decidere cosa vale la pena eseguire e perché.



SO
POUR
SUR L



MMME
L'AC
L'IA

legislazione

3.1 Pluralità, tensioni e pilastri: le coordinate del problema

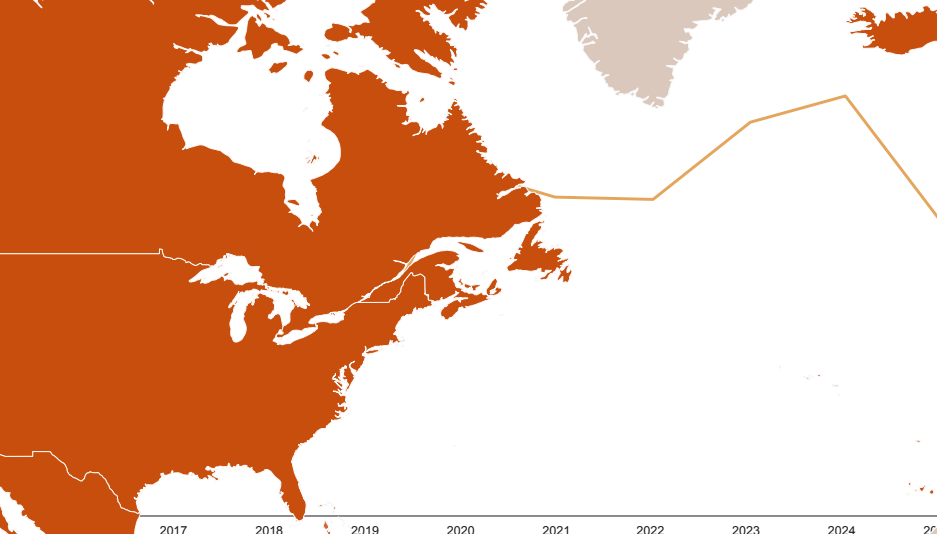
¹⁴⁰ Lawfare Media, *Understanding Global AI Governance Through a Three-Layer Framework*, 2025.

Se nel paradigma industriale novecentesco le risorse strategiche erano costituite da petrolio e acciaio, oggi il potere si articola sempre più attorno a capacità computazionale, dati e infrastrutture algoritmiche, configurando una nuova geografia della sovranità tecnologica¹⁴⁰. In questo contesto, la governance dell'intelligenza artificiale non è riducibile a un problema tecnico o normativo: è una questione sistemica che intreccia economia politica, relazioni internazionali, etica e sicurezza globale.

La crescente diffusione di sistemi algoritmici nei processi decisionali pubblici e privati ha generato una pressione senza precedenti verso la regolamentazione. Tuttavia, questa spinta regolatoria si sviluppa in un contesto caratterizzato da elevata eterogeneità istituzionale, divergenze culturali e competizione geopolitica. Di conseguenza, il tentativo di costruire un quadro globale unitario per la governance dell'AI appare, allo stato attuale, più aspirazionale che reale.

L'urgenza di affrontare i rischi sistemici dell'AI, che spaziano dalla discriminazione algoritmica alla sicurezza nazionale fino a scenari estremi di perdita di controllo tecnologico, ha sollecitato la nascita di nuove iniziative multilaterali. L'attuale panorama regolatorio è segnato da una proliferazione di iniziative non coordinate, spesso

Numero di proposte di legge sulla regolamentazione approvate nel 2016-25



¹⁴¹ Media of a Global Framework for AI Governance, 2024.

in parallelo da governi, organizzazioni internazionali e centri di ricerca¹⁴¹. Tali iniziative evidenziano un trend in crescita: mentre l'impatto dell'intelligenza artificiale sulla sua governance rimane profon-

¹⁴² Mirage News, *Global Push to Establish Urgent AI Governance*, 2025.

3.1.1 Governance normativa

¹⁴³ Egyptian Government / AI Authority, *Global AI Governance Frameworks: A Comparative Study*, 2024.

A partire dalle prime discussioni istituzionali sull'intelligenza artificiale, si è assistito a una molteplicità di iniziative a diversi livelli: globale, regionale e nazionale. Gli organismi globali definiscono principi generali, gli Stati adottano quadri legislativi specifici e i settori industriali implementano standard tecnici e pratiche operative. Tuttavia, tali iniziative non si sono evolute in modo coerente, ma hanno seguito traiettorie autonome, spesso scollegate tra loro. Organizzazioni internazionali come l'OCSE, G7 e Nazioni Unite hanno promosso principi generali e linee guida, mentre gli Stati hanno sviluppato strategie nazionali indipendenti, riflettendo priorità economiche e politiche specifiche¹⁴⁴.

Regionali
Principi di legge
Nessuna regolamentazione



Questa distinzione tra principi non vincolanti e sovranità normativa statale è fondamentale: mentre le organizzazioni internazionali possono facilitare il confronto e promuovere principi condivisi, la sovranità normativa rima-

Fig. 1: L'andamento del
regolamento
ovate

di
di
di

Il grado di
del 2024
che 118 Paesi
non può a nes
e principali iniziative
globali governance dell'AI. Solo sette Paesi –
tutti appartenenti al mondo sviluppato – partecipano a tutte¹⁴². Questo squilibrio evidenzia una forte asimmetria
nella distribuzione del potere normativo e nella capacità
di influenzare gli standard globali.

Il concetto di "compliance divide", ovvero il divario tra chi può e chi non può conformarsi agli standard emergenti, evidenzia ulteriormente questa dinamica, mostrando come esista una crescente distanza tra Paesi e organizzazioni in termini di capacità di implementare effettivamente le norme sull'AI¹⁴³. Questa divisione finisce per amplificare le disuguaglianze, creando un sistema a più velocità, in cui solo alcuni sono in grado di rispettare e influenzare gli standard emergenti. I Paesi con risorse limitate rischiano di rimanere esclusi dai benefici della trasformazione digitale, ampliando il divario tecnologico ed economico¹⁴⁴.

¹⁴⁴ Lazard, *The World's Artificial Intelligence*, 2025.

Legislazioni sulla Protezione dei Dati e della Privacy

Stanford 2025 AI Index Report

Questa eterogeneità non è solo la distribuzione geografica delle iniziative, ma anche la loro natura. Accanto ai framework governativi si affiancano linee guida prodotte da aziende tecnologiche, consorzi industriali e centri di ricerca, spesso orientate a interessi specifici e prive di coordinamento reciproco. Questo pluralismo di attori contribuisce a una "inflazione normativa" che, lungi dal favorire la chiarezza, aumenta l'incertezza regolatoria.

Figura 3.2: La mappa mostra la distribuzione globale delle legislazioni sulla protezione dei dati e della privacy. Grafico rielaborato da "The 2026 AI Index Report", Stanford HAI.

3.1.2 Le tensioni tra cooperazione globale e interessi privati

¹⁴⁵ *Mind Foundry, AI Regulations Around the World, 2025.*

La governance dell'AI è sempre più caratterizzata da una "privatizzazione normativa", in cui le regole emergono da pratiche industriali piuttosto che da processi democratici¹⁴⁵. Questo fenomeno, definito in letteratura come privatizzazione normativa, inverte la logica tradizionale della regolazione pubblica, in cui lo Stato fissa le norme e le imprese vi si conformano.

Tale inversione non è accidentale: è il prodotto di un'asimmetria strutturale di competenza tecnica. I governi dipendono dalle stesse aziende che intendono regolare per comprendere i sistemi che quelle aziende sviluppano. Questa dipendenza genera modelli di co-regolazione in cui le Big Tech partecipano attivamente alla definizione di standard, linee guida etiche e strumenti di audit, contribuendo a plasmare un campo normativo in cui sono al tempo stesso giocatori e arbitri.

A questo si aggiunge una seconda dinamica, di natura economica. Le aziende che controllano le infrastrutture di AI – modelli, dati e capacità computazionale – acquisiscono un vantaggio competitivo difficilmente replicabile, rafforzando posizioni dominanti e limitando la concorrenza.

Nonostante la frammentazione, i principali framework internazionali condividono un nucleo concettuale comune: responsabilità, equità, trasparenza, privacy, supervisione umana e sicurezza ricorrono trasversalmente in quasi tutti i documenti di governance prodotti a livello globale. Un consenso apparente, però, che si dissolve non appena si scende dal livello dei principi a quello della loro applicazione concreta, quando si tratta di decidere chi controlla, chi verifica, chi sanziona e a vantaggio di chi. È precisamente in questa zona grigia tra enunciazione e implementazione che le responsabilità si disperdono tra una molteplicità di attori senza che nessuno ne risponda pienamente.

Figura 3.3: Una cronologia dei dodici eventi che hanno definito la governance globale dell'IA nel 2025, dal South Korea AI Basic Act di gennaio alla proposta di modifica dell'EU AI Act di novembre. Grafico rielaborato da "AI Governance in 2025: a year in review", Patel O.

3.1.3 I tre pilastri della governance globale

Il panorama legislativo internazionale sull'intelligenza artificiale è oggi guidato da tre attori principali – Unione Europea, Stati Uniti e Cina – che incarnano approcci profondamente diversi tra loro, riflesso di altrettanto diverse visioni del rapporto tra tecnologia, individuo e Stato.

Governance globale dell'AI nel 2025

Riepilogo delle 12 tappe fondamentali che hanno definito le leggi, le politiche e la governance dell'AI a livello globale nel 2025.

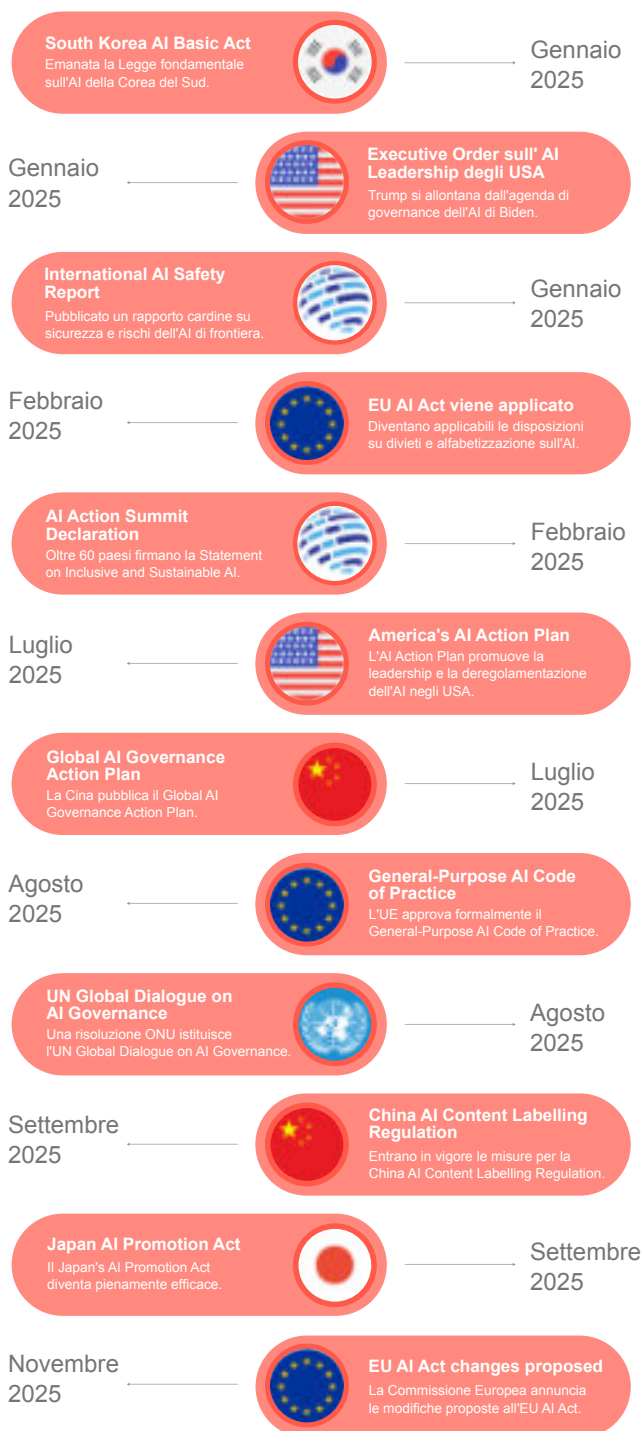


Figura 3.4: Il Primo Ministro cinese Li Qiang durante il suo intervento sul piano d'azione per l'intelligenza artificiale. Immagine da "China releases AI action plan days after the U.S. as global tech race heats up", CNBC.



Negli Stati Uniti, l'approccio alla governance dell'AI è storicamente caratterizzato da una forte enfasi sull'innovazione e sul ruolo del settore privato. Le grandi aziende tecnologiche svolgono un ruolo centrale nello sviluppo e nella diffusione dell'AI, mentre l'intervento regolatorio dello Stato tende a essere limitato e frammentato. Questo modello ha favorito una rapida crescita tecnologica, ma ha anche sollevato interrogativi sulla concentrazione del potere e sulla capacità di gestire i rischi sistemici¹⁴⁶.

L'Unione Europea rappresenta un modello alternativo, basato su un approccio normativo più strutturato, risk-based e orientato alla protezione dei diritti fondamentali. L'Unione Europea ha storicamente utilizzato la regolazione come strumento di soft power, cercando di esportare i propri standard attraverso il cosiddetto "Brussels effect". L'AI Act si inserisce in questa strategia, con l'obiettivo di stabilire un benchmark globale per la regolazione dell'AI¹⁴⁶.

La Cina, infine, adotta un modello fortemente centraliz-

¹⁴⁶ OneTrust, *Where AI Regulation is Heading in 2026: A Global Outlook, 2026*.

zato, in cui lo Stato svolge un ruolo diretto nella pianificazione e nel controllo dello sviluppo tecnologico. Le politiche cinesi combinano obiettivi di crescita economica con esigenze di controllo sociale e sicurezza nazionale, dando luogo a un sistema regolatorio integrato ma profondamente diverso da quello occidentale⁶.

Questi tre modelli non operano in isolamento, ma interagiscono e competono tra loro, influenzando la formazione degli standard globali. In particolare, si osserva una dinamica di “regulatory competition”, in cui gli Stati cercano di esportare i propri modelli normativi attraverso accordi commerciali, cooperazione internazionale e soft power tecnologico.

Figura 3.5: I capi di Stato, i delegati di governo e i leader delle aziende tecnologiche in occasione dell'AI Action Summit a Parigi, a Febbraio 2025. Immagine da “AI Action Summit”, Wikipedia.



3.2 L'Unione Europea e il modello regolatorio globale: ambizioni, compromessi e criticità

L'Europa ha scelto di giocare un ruolo di primo piano nella definizione di un quadro normativo che metta al centro i diritti umani e la sicurezza. Tuttavia, la strada verso una regolamentazione efficace è segnata da una complessa negoziazione tra la necessità di controllo e il desiderio di non soffocare la competitività tecnologica.

3.2.1 Il modello regolatorio europeo: sfide e prospettive

L'Unione Europea si è posizionata come pioniere assoluto nella governance dell'Intelligenza Artificiale, cercando di replicare il successo ottenuto con il Regolamento Generale sulla Protezione dei Dati (GDPR), entrato in applicazione nel 2018 e rapidamente diventato standard di riferimento globale per la protezione dei dati, attraverso il cosiddetto "Effetto Bruxelles".

L'Unione Europea ha infatti assunto un ruolo guida con l'approvazione dell'AI Act (Regolamento UE 2024/1689), un testo legislativo monumentale composto da 113 articoli e numerosi allegati e il primo quadro giuridico completo al mondo in materia di IA. Entrato in vigore nell'agosto 2024, il regolamento adotta un approccio basato sul rischio, classificando i sistemi di IA in quattro categorie: rischio inaccettabile, alto rischio, rischio limitato e rischio minimo. Questa tassonomia riflette una precisa vo-

¹⁴⁷ Parlamento Europeo e Consiglio dell'Unione Europea, *Regolamento (UE) 2024/1689 – Legge sull'Intelligenza Artificiale, 2024.*

lontà politica: proteggere i diritti fondamentali dei cittadini europei senza soffocare completamente l'innovazione¹⁴⁷.

Tuttavia, l'applicazione completa del regolamento mostra già le prime crepe. La Commissione Europea ha proposto un rinvio di 16 mesi, dall'agosto 2026 al dicembre 2027, per le norme relative ai sistemi ad alto rischio, ufficialmente per consentire la definizione di standard tecnici adeguati e strumenti di supporto alle imprese. Una dilazione che, al di là delle motivazioni tecnico-amministrative, segnala le difficoltà concrete di tradurre in obblighi operativi un impianto normativo tanto ambizioso quanto complesso da implementare.

Stato dell'IA Act dell'UE

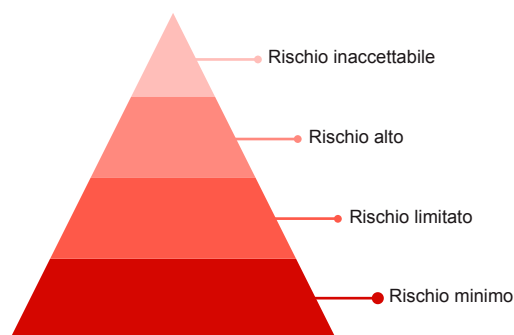
Stato dell'attuazione nei 27 Stati membri - Novembre 2025

26% 7 Paesi ✓ Completamente operativo Percorso di conformità chiaro. È possibile contattarli. Cipro OCECPR (Autorità unificata) Danimarca Digitaliseringsstyrelsen (Prima operativa) Finlandia Traficom + altri Germania Bundesnetzagentur (Coordinamento centrale) Lituania RRT (Operativa) Lussemburgo CNPD + settoriale Malta MDIA + IDPC (Doppia competenza)	30% 8 Paesi ⚠ Contattabile, ma con problematiche Rimangono confusione, incertezza o lacune. Repubblica Ceca Più di 4 autorità creano confusione, manca l'Art. 77 Francia Dispute ministeriali, nessuna designazione ufficiale Ungheria Disegno di legge non firmato, manca l'Art. 77 Irlanda Più di 15 autorità distribuite, forte confusione Italia Coordinamento dell'indipendenza UE non chiaro Lettonia Coordinamento di 15 autorità non chiaro Portogallo Stato operativo non chiaro Spagna Stato operativo non chiaro, in attesa di conferma)	44% 12 Paesi ✗ Nessuna autorità Nessuna autorità designata è stata ancora concordata. Le sanzioni e i requisiti si applicano comunque. Austria Paesi Bassi Belgio Polonia Bulgaria Romania Croazia Slovacchia Estonia Slovenia Grecia Svezia
--	---	--

3.2.2 La classificazione del rischio come modello regolatorio

Figura 3.6: La struttura piramidale a quattro livelli dell'AI Act europeo, che classifica i sistemi di IA da rischio minimo a rischio inaccettabile. Grafico rielaborato da "Classificazione rischi dell'AI Act: tutto quello che devi sapere", Nessun Segnale.

L'approccio piramidale basato sul rischio è una scelta metodologica che mira a proporzionare l'onere normativo alla potenziale pericolosità del sistema. Questa struttura, si articola, come anticipato, in quattro livelli distinti.



Alla base della piramide si trovano i sistemi a rischio minimo o nullo, come i filtri antispam o i videogiochi basati sull'IA, che sono esentati da obblighi specifici, riflettendo la volontà di non soffocare l'innovazione non problematica. Salendo, incontriamo i sistemi a rischio limitato, come i chatbot e i generatori di deepfake, per i quali l'Articolo 50 impone obblighi di trasparenza: gli utenti devono essere chiaramente informati che stanno interagendo con una macchina o consumando contenuti generati artificialmente.

Il livello ad "alto rischio", disciplinato dall'Articolo 6 e dettagliato nell'Allegato III, comprende sistemi utilizzati in settori critici dove un malfunzionamento o un uso improprio dell'IA potrebbe avere un impatto significativo sulla vita delle persone. Questi includono, ma non si limitano a, sistemi impiegati in infrastrutture critiche (come la gestione del traffico o delle reti energetiche), nell'istruzione e nella formazione professionale (ad esempio, per determinare l'accesso o valutare gli studenti), nell'occupazione e nella gestione dei lavoratori (per il reclutamento o il monitoraggio delle prestazioni), nei servizi pubblici e privati essenziali (per valutare l'ammissibilità a prestazioni sociali o crediti bancari), e nell'applicazione della legge (per la valutazione dell'attendibilità delle prove o del rischio di recidiva). I fornitori di questi sistemi sono soggetti a requisiti draconiani (Articoli 8-15), che includono l'implementazione di rigorosi sistemi di gestione del rischio per l'intero ciclo di vita del sistema (Art. 9), l'uso di set di dati di addestramento, convalida e prova di

Tabella 3.1: Lo stato di implementazione dell'AI Act nei 27 stati membri dell'UE a novembre 2025. Grafico rielaborato da "State of of the EU AI Act", AI Hub Landau.

¹⁴⁸ Council of Europe, *Convention on Artificial Intelligence, Human Rights, Democracy and the Rule of Law (CETS No. 225)*, 2024.

alta qualità per mitigare i bias e garantire l'accuratezza (Art. 10), la redazione di documentazione tecnica esauritiva per dimostrare la conformità (Art. 11), la registrazione in un database dell'UE e la garanzia di un'efficace supervisione umana (Art. 14). Questi obblighi, sebbene mirati a garantire l'affidabilità e la sicurezza, rappresentano un onere significativo per le imprese, specialmente per le PMI, e sollevano interrogativi sulla capacità delle autorità di vigilanza di far rispettare tali requisiti in modo uniforme ed efficace in tutti gli Stati membri¹⁴⁸.

Il punto più dibattuto del Regolamento risiede tuttavia nei livelli superiori della piramide, dove le implicazioni etiche e sociali diventano particolarmente sensibili. L'Articolo 5 definisce le pratiche di IA vietate, considerandole a "rischio inaccettabile" in quanto intrinsecamente contrarie ai valori fondamentali dell'Unione. Tra queste figurano i sistemi che impiegano tecniche subliminali o manipolatorie per distorcere il comportamento di una persona in modo da causare danni fisici o psicologici, una previsione che mira a prevenire l'abuso dell'IA per influenzare decisioni cruciali o per sfruttare vulnerabilità.

3.2.3 AI generativa e rischi sistemici: nuove frontiere normative

L'esplosione di popolarità di modelli come ChatGPT e altri sistemi di IA generativa durante la fase finale di stesura della legge ha costretto i legislatori europei a una rapida e complessa revisione del testo, aggiungendo un intero Titolo (Titolo IIa) dedicato ai modelli di IA per finalità generali (General Purpose AI - GPAI). Il risultato è un regime normativo a due livelli per questi modelli, delineato negli Articoli 51-55, che cerca di catturare la loro natura pervasiva e il loro potenziale impatto sistemico.

¹⁴⁹ Commissione europea, *Orientamenti per i fornitori di modelli di IA per finalità generali*, 2025.

Tutti i fornitori di modelli GPAI devono rispettare obblighi di trasparenza, inclusa la pubblicazione di un riepilogo sufficientemente dettagliato dei contenuti utilizzati per l'addestramento, una misura pensata per facilitare l'applicazione delle leggi sul diritto d'autore e per affrontare le crescenti preoccupazioni relative all'uso non autorizzato di opere protette¹⁴⁹. Tuttavia, i modelli classificati come aventi un "rischio sistemico" — definiti principalmente in base a una soglia di potenza di calcolo superiore a 10²⁵ FLOPs, ma anche attraverso criteri qualitativi che l'AI Office dovrà specificare — sono soggetti a obblighi molto più stringenti. Questi includono l'esecuzione di

¹⁵⁰ *Avv. A. Ingoglia, AI Act: Guida Completa al Regolamento Europeo sull'Intelligenza Artificiale, 2026.*

¹⁵¹ *Business People, Fermare l'AI Act: l'appello delle startup italiane ed europee, 2025.*

valutazioni del modello, test avversari (redteaming) per identificare vulnerabilità e comportamenti indesiderati, il monitoraggio e la segnalazione di incidenti gravi e l'implementazione di misure di sicurezza informatica adeguate¹⁵⁰.

Questa distinzione ha sollevato profonde preoccupazioni nell'industria tecnologica, in particolare tra le startup europee e i sostenitori dell'open-source. Molti temono che la soglia basata sui FLOPs sia arbitraria e che i costi di conformità associati a questi obblighi possano soffocare l'innovazione locale¹⁵¹, favorendo di fatto i giganti tecnologici statunitensi e cinesi che dispongono delle risorse necessarie per assorbire tali oneri. La critica è che l'AI Act, pur volendo promuovere un'IA etica, potrebbe inavvertitamente consolidare il potere dei pochi attori dominanti, creando barriere all'ingresso per i nuovi concorrenti e limitando la diversità dell'ecosistema IA europeo.

3.2.4 Governance, implementazione e l'impatto sulle PMI: un percorso irto di ostacoli

¹⁵² *Commissione europea, Orientamenti per i fornitori di modelli di IA per finalità generali, 2025.*

A livello centrale, la Commissione Europea ha istituito l'Ufficio Europeo per l'IA (AI Office), un organo dotato di poteri significativi per supervisionare l'attuazione delle regole sui modelli GPT, emettere codici di condotta, facilitare la cooperazione tra gli Stati membri e coordinare l'applicazione a livello transnazionale¹⁵². L'AI Office è chiamato a svolgere un ruolo cruciale nel garantire un'interpretazione e un'applicazione coerente del Regolamento, ma la sua efficacia dipenderà dalle risorse a sua disposizione e dalla sua capacità di navigare le complesse dinamiche politiche tra gli Stati membri.

¹⁵³ *Euronews, Regole sull'AI rimandate: l'UE divisa tra tutele e pressione delle grandi aziende, 2025.*

A livello nazionale, gli Stati membri devono designare autorità di notifica e autorità di vigilanza del mercato, responsabili dell'applicazione delle regole sui sistemi ad alto rischio. Questa struttura decentralizzata solleva il rischio di un'applicazione frammentata, simile a quanto avvenuto con il GDPR, dove le differenze di risorse, interpretazione e volontà politica tra le autorità nazionali (ad esempio, tra l'Irlanda e la Germania) hanno creato incertezze giuridiche e disuguaglianze competitive¹⁵³.

Un'altra criticità fondamentale riguarda l'impatto sulle Piccole e Medie Imprese (PMI). Sebbene il Regolamento preveda misure di supporto, come la creazione di "spazi di sperimentazione normativa" (regulatory sandboxes) per testare sistemi innovativi prima dell'immissione sul

¹⁵⁴ Deirdre Ahern, *Operationalising AI Regulatory Sandboxes under the EU AI Act: The Triple Challenge of Capacity, Coordination and Attractiveness to Providers*, Cambridge University Press, 2025.

mercato, molti analisti ritengono che l'onere burocratico e i costi di conformità rimangano sproporzionati per le PMI¹⁵⁴. La necessità di condurre Valutazioni d'Impatto sui Diritti Fondamentali (FRIA), introdotte dall'Articolo 27 per gli enti pubblici e privati che forniscono servizi pubblici utilizzando sistemi ad alto rischio, richiede competenze legali e tecniche specializzate che molte PMI semplicemente non possiedono o non possono permettersi. Il rischio è che l'AI Act, pur nobile nei suoi intenti di protezione, possa inavvertitamente consolidare monopoli tecnologici, creando barriere all'ingresso insormontabili per i nuovi attori e rallentando la crescita dell'ecosistema innovativo europeo.

3.2.5 Diritti umani e AI: il ruolo del Consiglio d'Europa

Parallelamente agli sforzi dell'Unione Europea, il Consiglio d'Europa ha elaborato la Convenzione Quadro sull'Intelligenza Artificiale, i Diritti Umani, la Democrazia e lo Stato di Diritto (CETS No. 225). Questo trattato internazionale, aperto alla firma nel settembre 2024, rappresenta il primo strumento giuridicamente vincolante a livello globale specificamente focalizzato sull'intersezione tra IA e diritti umani¹⁵⁵.

¹⁵⁵ European Parliament, *Recommendation on the Conclusion of the Council of Europe Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law (A10-0007/2026)*, 2026.

A differenza dell'AI Act, che è un regolamento sul mercato interno focalizzato sulla sicurezza dei prodotti e sulla libera circolazione, la Convenzione adotta un approccio basato sui principi, imponendo agli Stati firmatari (che includono nazioni extra-UE come Stati Uniti e Regno Unito) di adottare misure legislative o amministrative per garantire che il ciclo di vita dei sistemi di IA rispetti la dignità umana, l'uguaglianza, la non discriminazione, la privacy e la protezione dei dati¹⁵⁶.

¹⁵⁶ Council of Europe, *Responsible Artificial Intelligence – A Council of Europe Priority*, 2026.

Nonostante queste limitazioni e la potenziale confusione derivante dalla coesistenza di due quadri normativi distinti ma correlati, la Convenzione e l'EU AI Act insieme formano un blocco normativo formidabile, proiettando la visione europea di un'IA "antropocentrica" sul palcoscenico globale, ma non senza tensioni e sfide interpretative che richiederanno un attento coordinamento e una vigilanza costante da parte della società civile e di tutti gli attori coinvolti nella governance¹⁵⁷.

¹⁵⁷ Stoyanova, V., *Council of Europe Framework Convention on Artificial Intelligence: Context, Regulatory Approach and Scope of Obligations*, *European Journal of Risk Regulation*, Cambridge University Press, 2025.



3.2.6 Oltre l'AI Act: l'Europa può avere una sovranità normativa senza quella industriale?

Figura 3.7: La Presidente della Commissione Europea Ursula von der Leyen. Immagine da "Linee guida politiche 2024-2029 della Commissione europea "von der Leyen II"", CEP.

Approfondendo la dimensione geopolitica, è impossibile non menzionare il ruolo centrale di NVIDIA e della TSMC (Taiwan Semiconductor Manufacturing Company) nella definizione delle gerarchie globali. Se l'IA è il "motore" del futuro, i microchip sono il suo combustibile fossile. La competizione tra Stati Uniti e Cina si è cristallizzata in quella che molti analisti definiscono la "Silicon Curtain": un sistema di sanzioni, controlli alle esportazioni e sussidi statali (come il CHIPS Act americano) volti a impedire all'avversario l'accesso alle tecnologie di calcolo più avanzate. Per un designer di sistemi IA, questo significa che la scelta dell'hardware non è neutra, ma è legata a filiere produttive fragili e politicamente cariche. L'Europa, pur essendo leader nella produzione di macchine per la litografia (grazie alla olandese ASML), si trova in una posizione di dipendenza strategica sia per il design dei chip che per la loro produzione materiale. Questo "gap infrastrutturale" condiziona pesantemente la capacità normativa dell'UE: come si può imporre uno standard etico se non si ha il controllo sulla materia prima del calcolo? La sovranità digitale europea, dunque, non può passare solo per l'AI Act, ma deve includere una politica industriale che riporti la produzione di semiconduttori nel continente, come auspicato dall'European Chips Act.

In questo scenario, gli Stati Uniti cercano attivamente la collaborazione europea per rafforzare la sicurezza delle supply chain. Il sottosegretario di Stato americano per gli affari economici, Jacob Helberg, ha proposto all'Unione Europea di aderire alla "Pax Silica", un framework volto a coordinare le filiere dell'intelligenza artificiale tra Paesi alleati, con l'obiettivo di consolidare la leadership americana nell'IA e ridurre la dipendenza dalle tecnologie cinesi. Tuttavia, l'UE ha mostrato resistenza, con alcuni Stati membri, tra cui la Francia, che hanno bloccato l'autorizzazione alla Commissione europea per avviare negoziati formali. Critiche sono state mosse da figure come Helberg, che ha accusato Bruxelles di "regolare la propria irrilevanza" e ha definito il pacchetto normativo europeo, incluso l'AI Act, come "innovation killers". Nonostante ciò, l'Europa detiene una posizione indispensabile nella filiera globale dell'IA, in particolare grazie al controllo di seg-

¹⁵⁸ Linkiesta, *Alla fine gli Stati Uniti hanno bisogno dell'Europa sui chip*, 2026.

Figura 3.8: Il Palazzo Berlaymont a Bruxelles, sede ufficiale della Commissione Europea. Immagine da "La sede della Commissione europea a Bruxelles", Commissione Europea.





Anne Asensio

Automotive designer,
dirigente aziendale

Manager e designer con una solida esperienza internazionale tra industria automobilistica (Renault, General Motors) e innovazione digitale, con il ruolo di Vice President Design Experience presso Dassault Systèmes, Anne Asensio esplora le frontiere del design generativo. Il suo profilo è centrale nel dibattito politico e istituzionale, grazie anche al suo coinvolgimento in organizzazioni come la World Design Organization (WDO), dove contribuisce a definire il valore strategico del design per l'economia e la società europea.

Il suo lavoro si distingue per la capacità di integrare vincoli tecnologici e normativi all'interno del processo creativo, rendendola una figura fondamentale per analizzare l'evoluzione del ruolo del designer nel contesto dell'IA generativa. Questa competenza è essenziale per valutare l'impatto delle regolamentazioni europee, come l'AI Act, sulla capacità di innovazione del settore e sulla competitività globale.



Il design è una professione che si è sempre dovuta adattare ai cambiamenti tecnologici e culturali del proprio tempo: c'è un precedente storico a cui possiamo guardare per capire cosa l'IA sta facendo oggi al design?

L'IA non è qualcosa che arriva dal nulla. Ha una genealogia lunga, fatta di evoluzione e progressi tecnologici sequenziali. Se andiamo alle sue radici, ci rendiamo conto che l'IA fa già parte del percorso dell'evoluzione umana in materia di tecnologia e tecniche — che sono due cose distinte. Tutto è iniziato quando abbiamo preso certe decisioni scientifiche fondamentali e abbiamo smesso di vedere il mondo da un punto di vista puramente teologico. La scienza stessa ha cominciato ad assumere i caratteri di una religione. E poi, gradualmente, le tecniche — e le persone capaci di applicarle in ogni aspetto della vita quotidiana — hanno permesso la diffusione di ciò che conosciamo oggi: un ambiente tecnologico che permea ogni dimensione dell'esistenza.

E allora, qual è l'effetto dell'IA sul design? Semplicemente, lo ha costretto ad assumere nuovamente un atteggiamento critico e riflessivo su se stesso, di fronte alla sfida rappresentata da ciò che sta accadendo sul fronte tecnologico.

È la stessa dinamica del Bauhaus, che nacque principalmente come reazione alla standardizzazione — a un modello prestabilito, alla produzione in serie di un oggetto rigidamente standardizzato che stava chiaramente perdendo ogni attrattiva. E all'improvviso si è cominciato a prestare attenzione a come quell'oggetto verrà usato. Cos'è un oggetto? Di cosa è fatto? Non è solo una superficie, un'apparenza, una funzione: è il modo in cui viene concepito nella sua interezza. E quel processo è sempre quello che la tecnologia ignora completamente. La tecnologia punta all'ottimizzazione pura,

all'esecuzione perfetta di qualcosa che già conosciamo. Ma le persone amano ciò che non conoscono ancora — ed è esattamente questo il fascino che ritroviamo anche nel Bauhaus.

Col tempo, però, anche l'approccio dei designer è diventato molto regolamentato, standardizzato — al punto che ci troviamo forse alla fine del design moderno. Il processo di design tradizionale è sequenziale: ideazione, sviluppo, prototipazione e così via. E in questo momento la tecnologia sta cercando di capire come integrare questo design altamente standardizzato in un processo automatico. Ma quel processo aveva già raggiunto la fine della sua curva canonica — era già arrivato il momento di cambiare, anche senza l'IA. L'IA non fa altro che accelerare questa necessità.

Sappiamo che l'AI Act europeo sta definendo le regole per lo sviluppo e l'implementazione dell'intelligenza artificiale. Ma qual è la sua opinione sulla possibilità che una regolamentazione eccessiva finisca per frenare l'innovazione?

In Europa siamo famosi per la nostra propensione alla regolamentazione, e tra le diverse nazioni europee condividiamo valori abbastanza comuni. Avere un player IA europeo sarebbe bello. Ma se si limita a giocare con le stesse regole del modello americano, non ne vedo la necessità. Quello di cui abbiamo bisogno è una GenAI che non operi secondo le stesse logiche. Questo è l'unico punto che conta, adesso. Il modello attuale si basa sul sottrarre conoscenza e big data per diventare intelligente, utile, competitivo. Si potrebbe invece immaginare un'IA che utilizzi le informazioni all'interno di un sistema di distribuzione equilibrata — ma al servizio di cosa? Di quale scopo? Ed è qui che la questione diventa

davvero difficile da districare.

Quando parlavo dell'evoluzione della tecnologia, pensavo a come in origine fosse uno strumento, un aiuto, una forza al servizio della vita quotidiana. Quando è diventata "tecnica" nel senso moderno, ha iniziato a essere appannaggio di pochi: gli ingegneri. È stato creato un patto implicito tra politica e tecnica — i decisori pubblici fissavano il piano, i tecnici lo eseguivano. Il problema oggi è che la tecnologia punta all'autonomia. Se la tecnologia diventa l'unico mezzo disponibile, finisce per diventare fine a se stessa. Non produce beni migliori né condizioni migliori per l'interesse generale — soprattutto quando è controllata, o peggio posseduta, da pochissimi.

Oggi concepiamo il mondo come qualcosa che possiamo controllare artificialmente attraverso la fisica, attraverso le sue leggi. In pratica, lo abbiamo ridotto a questo. Ed è molto facile immaginare che l'intelligenza artificiale sviluppi capacità generative in un mondo già predisposto per accoglierla. Ma proprio da quel momento abbiamo iniziato a sperimentare, lentamente ma inesorabilmente, una sensazione che mette in discussione l'essere umano contemporaneo.

La frammentazione che ne deriva cerca semplicemente di arginare qualcosa. È come mettere una benda, o tentare diappare una falla: l'acqua fuoriesce da ogni parte, tu premi con le dita, ma continua a uscire lo stesso — ovunque. Dobbiamo ripensare al cuore della questione, alla fonte, all'origine, al punto in cui la crepa si è aperta. E da lì ricominciare. In fondo, qualcosa di simile è già stato fatto: pensiamo allo smantellamento dei grandi monopoli negli Stati Uniti, quando ci si rese conto che solamente cinque persone controllavano le ferrovie, l'industria chimica e il petrolio dell'intero paese. Oggi dovremmo fare la stessa cosa, ma su scala mondiale. Non voglio sco-

raggiarvi, ma ho la sensazione che la regolamentazione sia, nella migliore delle ipotesi, solo un bel gesto.

Detto questo: come designer, non dobbiamo essere naïve. I designer sono, tra tutte le figure professionali, probabilmente i più attrezzati per trovare la propria strada in questo cambiamento. Quello che mi preoccupa non è che l'IA trasformi il design — il design è flessibile per natura, è proprio questa flessibilità a renderlo unicamente umano. È la nostra capacità inventiva: qualunque tecnologia arrivi tra le nostre mani — e uso il termine "mani" deliberatamente, perché c'è sempre un momento fisico, tattile, nel modo in cui recepiamo e valutiamo qualcosa — noi impariamo a capire se ci è utile o no.

Abbiamo bisogno che tutti inizino a pensare come designer. E la verità è che i designer non si arrendono mai — non abbiamo mai abbandonato il corpo, i sensi, le emozioni nel confrontarci con la tecnica e la tecnologia. È questo che ci rende riflessivi, e probabilmente una risorsa preziosa: smettere di essere strumentalizzati nel processo e diventare invece una forza attiva, una forza di riparazione di ciò che è stato rotto.

Quindi, guardando alle capacità degli LLM e di tanti altri modelli, la domanda è una sola: per quale scopo viene utilizzata l'intelligenza artificiale? E se lo fa per qualcosa che è davvero nell'interesse collettivo, allora ha senso. Altrimenti, no.

Lei ha parlato del processo progettuale. Vorremmo soffermarci sul progetto TA. TAMU, sviluppato insieme a Patrick Jouin: qui non siete stati voi a definire il formato finale, ma l'intelligenza artificiale, a partire da una serie di vincoli che le avevate fornito. In che modo cambia, allora, il ruolo del designer

quando il processo prende avvio dai vincoli piuttosto che da un'intenzione formale predefinita?

Quando abbiamo sviluppato Ta.tamu, che all'epoca si chiamava semplicemente Tamu, era il primo approccio riflessivo al design: sapevamo cosa volevamo e cosa non volevamo. Non volevamo cedere il processo di design alla macchina. Avremmo potuto usare un approccio generativo generico, ma non ci interessava. Volevamo un design costruito su parametri specifici e precisi: definisci uno spazio di design, e quei parametri producono calcoli strutturali, simulazioni statiche, soluzioni di posizionamento.

All'epoca già Mario Klingemann stava esplorando a fondo quello che chiamiamo "spazio latente" — ciò che accade tra l'input e l'output, quella zona di calcolo e operazione statistica che non funziona come il cervello umano. È interessante perché, per il grande pubblico, anche il processo del designer è uno spazio latente: non si sa bene cosa succede dentro, è quella che viene chiamata "black box". Ebbene, davanti all'AI anche il designer si trova di fronte a una black box. Noi abbiamo cercato di aggirarla costruendo una sequenza diversa da quella tradizionale: abbiamo reso l'iterazione rapida, continua, non lineare. L'obiettivo era stare in uno spazio e poter essere in ogni punto di quello spazio contemporaneamente — partire dall'inizio, dalla fine, dal mezzo. Era una navigazione, non un percorso.

Con Patrick lavorammo con processi molto simili — il suo team e il mio. A volte i team si univano, a volte si separavano. E i momenti di separazione erano spesso i più interessanti. Quando eravamo uniti, andavamo avanti. Quando ci dividevamo, la domanda era: come ci ricongiungiamo? Era più caotico, più nebuloso, ma eravamo in grado di generare dei risultati più soddisfacenti.

L'IA non ha idee proprie: riflette le nostre. Riflette la qualità di ciò che pensiamo. Non esistono idee che vengano davvero dall'IA, a meno che non si tratti di decisioni e giudizi automatici — e qui mi fermo e vi dico: non delegate mai una decisione all'IA. Mai. Formatevi la vostra opinione, ma non fate validare le vostre decisioni — e tantomeno le decisioni di design — da una macchina.

Qualcuno dirà che sono una vecchia signora che non capisce. Ma parlo da un punto di vista profondamente umano: nessuna tecnologia dovrebbe decidere al posto nostro. Mai. Al massimo può proporre opzioni, illuminare scenari, offrire una prospettiva più ampia — e questo mi piace dell'IA: abbatte i confini di dominio, ti permette di attingere a biologia, fisica, chimica, scienze, in modo trasversale. Ma fino al livello di ciò che riesci ad elaborare. Questo è bello, è aperto, preserva l'ignoto. Ma non deve sostituirsi alla tua decisione. L'idea è tua. La decisione è tua.

Nel corso della sua carriera ha lavorato in ambiti molto diversi, dal design automobilistico all'innovazione aziendale fino alla progettazione di sistemi complessi. Guardando all'impatto crescente dell'Intelligenza Artificiale generativa, come immagina l'evoluzione del ruolo del designer nei prossimi anni? E quali competenze ritiene saranno determinanti per distinguere i professionisti capaci di creare valore in questo nuovo scenario?

È un tema che è letteralmente sul tavolo di ogni accademia, scuola di design e istituto di ricerca. Le nuove competenze. Personalmente, io difenderò e incoraggerò le competenze che ci hanno resi buoni designer prima che tutto venisse artificializzato attraverso sistemi matematici, prima del 3D, prima di tutto.

Tutti dicono che la competenza chiave sarà il pensiero critico. Va bene, d'accordo — ma come si insegna il pensiero critico? Come lo si fa davvero? Non c'è altro modo se non l'esperienza diretta. E cosa significa fare esperienza diretta? Significa passare attraverso le cose reali con le mani, con il cervello, con il corpo. Bisogna toccare il reale, e soprattutto scontrarsi con il reale. Nel digitale non ti fai mai del male. È tutto possibile, tutto infinito. E questa grande teoria dell'infinito che ha pervaso la scienza e la filosofia, e che culmina nell'AI generativa è affascinante, è eccitante — ma io preferisco l'ignoto in un mondo finito. Perché anche in un mondo chiuso, l'ignoto da esplorare è immenso. E come lo esplori? Abbracciandolo, toccandolo, mettendolo alla prova, chiedendo consiglio alle persone intorno a te — ai tuoi genitori, ai tuoi figli, a chiunque. È quello che ci ha resi umani fin dall'inizio, ed è quello che dobbiamo preservare.

Il mondo è finito. E allora la domanda diventa: quali competenze servono per operare in un mondo finito? I più capaci di trovare soluzioni saranno probabilmente quelli che sanno fidarsi del proprio corpo e delle proprie capacità, ma che accettano anche i propri limiti. Ed è proprio questa nozione di limite che è straordinaria. In un mondo che si racconta come illimitato, l'essere umano viene di fatto declassato: invecchia, si ammala, è incoerente, è imprevedibile. Un mondo senza limiti crea le condizioni perfette per le macchine, non per gli esseri umani — e finisce per renderli obsoleti.

Le competenze per i designer stanno esattamente nell'opposto: accettare il limite significa accettare il limite del corpo, dell'intelligenza, della fatica, dell'età. E non appena iniziamo a ragionare così, non più in termini di progresso infinito ma di buona vita all'interno dei propri limiti, tutto cambia. Smettendo di misu-

rarci con le macchine, ciascuno di noi diventa un ecosistema ricco e unico, in attesa di esprimersi. La vera sfida educativa è capire come coltivare e valorizzare quella diversità, invece di formare persone per eseguire ruoli già scritti all'interno di un sistema industriale che poi scopre di poterli sostituire con una macchina.

3.3 Nord America:

Tra deregolamentazione,
leadership tecnologica e
nuovi orizzonti normativi

Il continente americano presenta un mosaico regolatorio sull'Intelligenza Artificiale che riflette profonde differenze filosofiche, priorità geopolitiche e divergenti visioni sul ruolo dello Stato nell'innovazione tecnologica. Mentre gli Stati Uniti hanno oscillato tra un approccio di governance basato su ordini esecutivi e una recente spinta verso la deregolamentazione, il Canada e diverse nazioni latinoamericane stanno esplorando percorsi legislativi più strutturati, spesso ispirati al modello europeo, ma adattati alle proprie specificità economiche e sociali.

3.3.1 Stati Uniti d'America: un pendolo tra sicurezza nazionale e competitività economica

La politica statunitense sull'IA è stata storicamente caratterizzata da un approccio settoriale e da una preferenza per la soft-law, con un'enfasi sulla promozione dell'innovazione e della leadership tecnologica. L'assenza di una legislazione federale onnicomprensiva sull'IA ha lasciato un vuoto che è stato parzialmente colmato da direttive presidenziali, iniziative di agenzie federali e, sempre più spesso, da tentativi di regolamentazione a livello statale. Tuttavia, l'escalation delle capacità dell'IA e le crescenti preoccupazioni per la sicurezza nazionale, i diritti civili e la stabilità economica hanno portato a un'evoluzione significativa, culminata nell'Executive Order 14110

¹⁵⁹ The White House, *Executive Order on the Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence (EO 14110)*, 2023.

dell'ottobre 2023, emesso dall'amministrazione Biden¹⁵⁹.

L'Executive Order 14110, intitolato "Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence", rappresentava il tentativo più ambizioso di stabilire un quadro federale per la governance dell'IA. Questo ordine, sebbene non una legge in senso stretto, imponeva una serie di direttive alle agenzie federali, con l'obiettivo di mitigare i rischi e promuovere uno sviluppo responsabile. La sua portata era notevole, toccando otto aree chiave: nuovi standard di sicurezza e protezione per l'IA, protezione della privacy, promozione dell'equità e dei diritti civili, difesa dei consumatori, sostegno ai lavoratori, promozione dell'innovazione e della concorrenza, avanzamento della leadership americana a livello internazionale e garanzia di un uso responsabile dell'IA da parte del governo. L'intento era quello di creare un meccanismo di pre-market review per le IA più avanzate, un approccio che, sebbene meno stringente di quello europeo, segnalava una crescente preoccupazione per i rischi sistemici.

La Sezione 5 si concentrava sulla protezione dei consumatori e dei diritti civili, incaricando agenzie come la Federal Trade Commission (FTC) e il Dipartimento di Giustizia di indagare e perseguire pratiche discriminatorie o ingannevoli legate all'IA. L'ordine toccava anche aspetti cruciali come la privacy, la concorrenza e l'impatto dell'IA sul mercato del lavoro, delineando un approccio olistico alla governance. Nonostante la sua portata, l'EO 14110 è stato oggetto di critiche da parte di alcuni settori dell'industria, che lo consideravano eccessivamente oneroso e potenzialmente inibitore dell'innovazione, e da parte di gruppi per i diritti civili, che lo ritenevano insufficiente a garantire una protezione robusta contro gli abusi dell'IA, in particolare per quanto riguarda la sorveglianza biometrica e l'uso dell'IA nel sistema giudiziario¹⁶⁰. La sua natura di ordine esecutivo, inoltre, lo rendeva vulnerabile a future modifiche o revoche da parte di amministrazioni successive, come si è poi verificato.

¹⁶⁰ National Institute of Standards and Technology (NIST), *Artificial Intelligence Risk Management Framework (AI RMF 1.0)*, NIST AI 100-1, 2023.

Il panorama politico statunitense ha subito una brusca virata con l'insediamento della nuova amministrazione Trump nel gennaio 2025. Una delle prime azioni significative è stata la revoca dell'Executive Order 14110, sostituito da un nuovo ordine esecutivo intitolato "Removing Barriers to American Leadership in Artificial Intelligence"¹⁶¹. Questa mossa ha segnalato un cambiamento radicale nella filosofia di governance, privilegiando apertamente la deregolamentazione e la promozione aggressiva dell'innovazione a scapito di un approccio più cauto e

¹⁶¹ The White House, *Removing Barriers to American Leadership in Artificial Intelligence*, 2025.

basato sulla mitigazione del rischio.

La revoca dell'EO 14110 ha anche sollevato interrogativi sulla coerenza della politica federale e sulla capacità degli Stati Uniti di mantenere un approccio unificato alla governance dell'IA, lasciando un vuoto che potrebbe essere colmato da iniziative statali frammentate o da un'accresciuta dipendenza da standard volontari come l'AI Risk Management Framework (AI RMF) del National Institute of Standards and Technology (NIST). Sebbene l'AI RMF sia un documento autorevole e ampiamente adottato, la sua natura volontaria ne limita l'efficacia come strumento di applicazione coercitiva, creando un "patchwork" regolatorio che può generare incertezza e iniquità.

In assenza di una legge federale onnicomprensiva e con la revoca dell'EO 14110, le agenzie federali continuano a svolgere un ruolo cruciale nell'applicazione delle normative esistenti all'IA. La FTC, ad esempio, utilizza le sue competenze in materia di protezione dei consumatori e concorrenza per affrontare pratiche ingannevoli o sleali legate all'IA, come le affermazioni fuorvianti sulle capacità dei sistemi o l'uso di IA per pratiche anticoncorrenziali. Allo stesso modo, la Equal Employment Opportunity Commission (EEOC) indaga sulla discriminazione algoritmica nel contesto lavorativo, assicurando che gli strumenti di IA utilizzati per il reclutamento, la promozione o la valutazione delle prestazioni non perpetuino o creino bias discriminatori. Un approccio per certi versi pragma-

Figura 3.9: Il Presidente Donald Trump firma l'ordine esecutivo per introdurre uno standard nazionale sull'intelligenza artificiale, limitando le leggi approvate dai singoli stati. Immagine da "Trump Signs Order Aimed at Curbing State AI Laws", Reuters.



tico, ma che affida la governance di una tecnologia sistemica a strumenti normativi pensati per un'epoca precedente, rattoppando con competenze settoriali quello che richiederebbe un disegno legislativo organico.

3.3.2 Canada: L'Artificial Intelligence and Data Act (AIDA) e la ricerca di un equilibrio prudente

Il Canada si è distinto nel panorama nordamericano per aver proposto una delle prime leggi organiche sull'IA, l'Artificial Intelligence and Data Act (AIDA), parte del più ampio Bill C-27 (Digital Charter Implementation Act). L'AIDA rappresenta un tentativo di bilanciare la promozione dell'innovazione con la protezione dei diritti individuali e la mitigazione dei rischi associati all'IA, cercando di posizionarsi come un modello intermedio tra l'approccio più leggero degli Stati Uniti e quello più prescrittivo dell'Unione Europea¹⁶².

¹⁶² Parliament of Canada, *Bill C-27: An Act to enact the Consumer Privacy Protection Act, the Personal Information and Data Protection Tribunal Act and the Artificial Intelligence and Data Act*, 2023.

L'AIDA adotta un approccio basato sul rischio, sebbene con una terminologia e una struttura leggermente diverse rispetto all'EU AI Act. La legge si concentra sui sistemi di IA ad "alto impatto", definiti nella Sezione 5 come quelli che possono avere un impatto significativo sulla salute, la sicurezza o i diritti umani dei canadesi. Questa definizione, sebbene ancora in attesa di ulteriori chiarimenti tramite regolamenti secondari, mira a coprire settori come la sanità, i trasporti, l'occupazione e l'applicazione della legge.

Un elemento distintivo dell'AIDA è la creazione di un AI and Data Commissioner, un ruolo previsto dalla Sezione 12, con il compito di supervisionare l'applicazione della legge, condurre indagini, emettere ordini di conformità e imporre sanzioni. Le sanzioni per le violazioni più gravi possono raggiungere il 5% del fatturato globale o 25 milioni di dollari canadesi, un deterrente significativo per le aziende. Questo conferisce al Canada un'autorità di regolamentazione dedicata all'IA, simile all'AI Office europeo, ma con un mandato più ampio che include anche la protezione dei dati¹⁶³.

¹⁶³ Government of Canada – Innovation, Science and Economic Development Canada, *The Artificial Intelligence and Data Act (AIDA) – Companion Document*, 2023.

Nonostante l'approccio proattivo, l'AIDA ha ricevuto critiche da diverse parti. Alcuni esperti hanno sollevato preoccupazioni sulla vaghezza di alcune definizioni, in particolare quella di "sistema di IA ad alto impatto", che potrebbe creare incertezza giuridica per le imprese e rendere difficile l'applicazione uniforme della legge. Altri hanno evidenziato la mancanza di un meccanismo robu-

sto per la protezione dei diritti umani, suggerendo che la legge si concentri maggiormente sulla sicurezza tecnica piuttosto che sulla giustizia sociale e sulla prevenzione della discriminazione algoritmica. Inoltre, la relazione tra l'AIDA e le leggi provinciali esistenti in materia di privacy e protezione dei dati rimane un'area di potenziale conflitto e complessità implementativa, creando un rischio di frammentazione normativa all'interno del Canada stesso. La capacità del Canada di implementare efficacemente l'AIDA dipenderà dalla chiarezza delle linee guida che accompagneranno la legge e dalla disponibilità di risorse per il nuovo AI and Data Commissioner, che dovrà navigare un delicato equilibrio tra la promozione dell'innovazione e la protezione dei cittadini.

3.4 America Latina: un mosaico di iniziative nazionali e la ricerca di un modello proprio

L'America Latina sta emergendo come un'altra regione attiva nella regolamentazione dell'IA, con diversi paesi che stanno sviluppando quadri normativi ispirati, in parte, al modello europeo, ma con un'attenzione particolare alle proprie esigenze di sviluppo, inclusione digitale e protezione contro i rischi specifici del contesto regionale. La regione è caratterizzata da una grande diversità economica e politica, che si riflette nella varietà degli approcci alla governance dell'IA.

Il panorama normativo nelle Americhe è quindi un campo di sperimentazione e adattamento, riflettendo le diverse priorità e filosofie politiche. Mentre gli Stati Uniti sembrano inclinarsi verso un modello di autoregolamentazione e promozione dell'innovazione, con un ruolo più limitato per la legislazione federale, il Canada e il Brasile stanno abbracciando un approccio più interventista, cercando di stabilire regole chiare per garantire un'IA etica e responsabile. Queste diverse traiettorie riflettono le priorità politiche e le visioni economiche di ciascuna nazione, ma tutte condividono la consapevolezza che la governance dell'IA è una questione di importanza strategica e geopolitica ineludibile, con implicazioni profonde per i diritti umani, la democrazia e la competitività globale.

	Argentina	Colombia	Peru	Cile	Messico	Brasile
Meccanismi di conformità normativa sull'AI in America Latina						
Approvazione normativa	■	■	■	■	■	■
Valutazioni di Impact Assessment	■	■	■	■	■	■
Audit	■	■	■	■	■	■
Reporting governativo	■	■	■	■	■	■
Obbligo di esplicabilità	■	■	■	■	■	■
Obbligo di pagare le royalty ai titolari di proprietà intellettuale dei dati di addestramento	■	■	■	■	■	■

Quadro normativo dell'AI in America Latina

Istituzionalità (es. organismi di promozione dell'AI)	■	■	■	■	■	■
Disegni di legge isolati presentati da singoli parlamentari	■	■	■	■	■	■
Disegni di legge promossi dal governo	■	■	■	■	■	■
Legge sull'AI globale	■	■	■	■	■	■

3.4.1 Brasile: protezione dei diritti come priorità

Il Brasile è in prima linea nella regione con il Disegno di Legge 2338/2023, approvato dal Senato nel dicembre 2024 e in attesa di promulgazione. Questo disegno di legge, fortemente influenzato dall'EU AI Act, mira a stabilire un quadro giuridico completo per lo sviluppo e l'uso dell'IA nel paese. La sua genesi è stata caratterizzata da

¹⁶⁴ Senado Federal do Brasil, *Projeto de Lei nº 2338, de 2023 (PL 2338/2023)*, 2023.

Tabella 3.2: Un confronto tra i principali meccanismi di conformità e i framework regolativi adottati da Argentina, Colombia, Perù, Cile, Messico e Brasile. Grafico rielaborato da "Quali diritti per le intelligenze generative?", Ziccardi G.

Figura 3.10: Il Presidente del Senato brasiliano Rodrigo Pacheco. Immagine da "Brasile, il Senato approva il marco temporal contro i diritti dei popoli indigeni", Lifegate.

¹⁶⁵ Mattos Filho, *Regulatory Framework for Artificial Intelligence Passes in Brazilian Senate*, 2024.

un ampio dibattito pubblico e da un forte coinvolgimento della società civile, che ha spinto per un approccio centrato sui diritti umani¹⁶⁴.

Gli Articoli 1-4 delineano principi fondamentali come il diritto alla spiegazione delle decisioni algoritmiche, il diritto alla non discriminazione, il diritto alla correzione dei dati personali utilizzati dall'IA e il diritto alla supervisione umana. La legge introduce una classificazione del rischio, distinguendo tra sistemi a "rischio eccessivo" (proibiti, come quelli che manipolano il comportamento umano o creano sistemi di social scoring) e ad "alto rischio", per i quali sono previsti obblighi di conformità rigorosi. Questi includono la necessità di condurre valutazioni d'impatto algoritmico (Algorithmic Impact Assessment) come dettagliato negli Articoli 20-25, che richiedono un'analisi approfondita dei potenziali impatti sui diritti fondamentali e sulle libertà civili.

La creazione di un'autorità competente per la supervisione dell'IA è un altro pilastro della proposta brasiliana, riflettendo la crescente consapevolezza della necessità di un'applicazione robusta e indipendente¹⁶⁵. Tuttavia, la sfida sarà garantire che questa autorità abbia le risorse e l'indipendenza necessarie per far rispettare efficacemente la legge in un paese così vasto e complesso.



3.4.2 Cile: strategia nazionale e proposte legislative per un'IA al servizio pubblico

¹⁶⁶ Gobierno de Chile – Ministerio de Ciencia, Tecnología, Conocimiento e Innovación, *Estrategia Nacional de Inteligencia Artificial*, 2021.

Il Cile è stato uno dei primi paesi latinoamericani a pubblicare una Strategia Nazionale di Intelligenza Artificiale nel 2021, che ha posto le basi per le discussioni legislative. Questa strategia ha evidenziato l'importanza di un'IA etica, inclusiva e al servizio dello sviluppo sostenibile¹⁶⁶.

La sfida per il Cile, come per altri paesi della regione, è quella di sviluppare una regolamentazione che sia al contempo efficace nella protezione dei diritti e sufficientemente flessibile da non ostacolare l'innovazione in un'economia emergente. La dipendenza da tecnologie importate e la necessità di costruire capacità locali di IA rendono il contesto cileno particolarmente delicato, richiedendo un equilibrio tra l'adozione di standard internazionali e l'adattamento alle realtà locali.

Figura 3.11: La Ministra della Scienza, Tecnologia, Conoscenza e Innovazione cilena Aisén Etcheverry. Immagine da "Chile adopts strict climate policy and environmental preservation in the Patagonia region", Latin Gulf.



3.4.3 Argentina: linee guida etiche e un dibattito aperto sull'ecosistema locale

¹⁶⁷ Ministerio de Ciencia, Tecnología e Innovación de Argentina, *Lineamientos Éticos para el Desarrollo y Uso de la Inteligencia Artificial*, 2021.

L'Argentina, pur non avendo ancora una legislazione specifica sull'IA, ha promosso un dibattito pubblico attivo e ha sviluppato Linee Guida Etiche per l'Intelligenza Artificiale attraverso il Ministero della Scienza, Tecnologia e Innovazione. Queste linee guida, sebbene non vincolanti, promuovono principi come l'equità, la trasparenza, la responsabilità e la privacy nell'uso dell'IA, cercando di orientare lo sviluppo tecnologico verso valori democratici e inclusivi¹⁶⁷.

Il paese sta monitorando attentamente gli sviluppi normativi internazionali, in particolare in Europa e Brasile, per informare le proprie future scelte legislative. La discussione in Argentina si concentra sulla necessità di un quadro che possa stimolare l'ecosistema locale di startup IA, proteggendo al contempo i cittadini dai potenziali abusi. La sfida principale è tradurre questi principi etici in meccanismi legali concreti e applicabili, senza imporre oneri eccessivi a un settore tecnologico ancora in fase di crescita. La mancanza di un consenso politico robusto e la volatilità economica del paese rendono il percorso verso una legislazione sull'IA particolarmente complesso.



Figura 3.12: Il Presidente dell'Argentina Javier Milei. Immagine da "History of Argentina", History Maps.

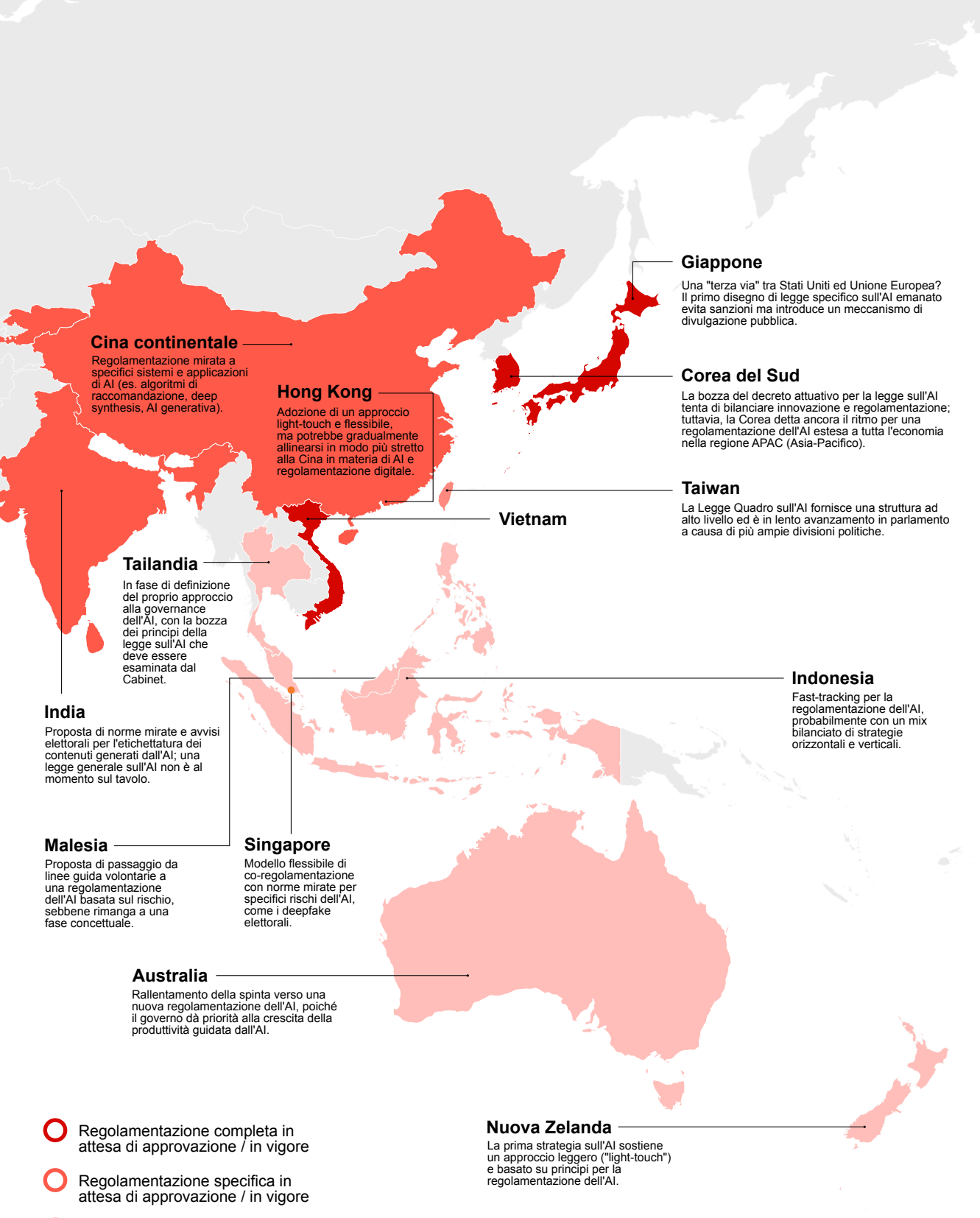
3.5 L'Asia-Pacifico: Laboratorio di governance tra controllo statale, soft-law e ambizioni geopolitiche

La regione dell'Asia-Pacifico rappresenta uno dei teatri più dinamici e diversificati per la governance dell'Intelligenza Artificiale. Qui, le potenze tecnologiche stanno sperimentando modelli normativi che riflettono profondamente le loro strutture politiche, le priorità economiche e le concezioni del rapporto tra Stato, cittadino e tecnologia. Dall'approccio prescrittivo e mirato della Cina, alla flessibilità strategica del Giappone, fino ai tentativi di bilanciamento della Corea del Sud e dell'India, l'Asia-Pacifico offre uno spaccato affascinante delle molteplici vie alla regolamentazione dell'IA, spesso intrecciate con ambizioni geopolitiche e sfide interne.

3.5.1 Cina: la regolamentazione iterativa e l'imperativo del controllo statale

La Repubblica Popolare Cinese ha adottato una strategia normativa unica, spesso definita "iterativa" o "per componenti". Invece di tentare di redigere un'unica legge onnicomprensiva fin dall'inizio (come l'EU AI Act), Pechino ha scelto di emanare regolamenti specifici per affrontare le singole tecnologie a mano a mano che queste raggiungono la maturità commerciale e pongono sfide immediate al controllo statale e alla stabilità sociale¹⁶⁸. Un approccio pragmatico e incrementale, che consente al governo di adattarsi rapidamente all'evoluzione tecno-

¹⁶⁸ ScienceDirect, *Agile and Iterative Governance: China's Regulatory Response to AI*, 2025.



- Regolamentazione completa in attesa di approvazione / in vigore
- Regolamentazione specifica in attesa di approvazione / in vigore
- Valutazione dell'approccio normativo

Figura 3.13: Una mappa della regione Asia-Pacifico che sintetizza l'approccio regolatorio di ciascun Paese nel 2025. Grafico rielaborato da "AI-Pac August–October edition: Flint's Digest on AI Policy in Asia-Pacific", Gupta A. et al.

¹⁶⁹ Hogan Lovells, *China Finalizes Generative AI Regulation, 2023*.

logica, ma solleva interrogativi sulla coerenza complessiva e sulla prevedibilità del quadro normativo.

Il fulcro di questo approccio è rappresentato dalle "Misure Interinali per la Gestione dei Servizi di Intelligenza Artificiale Generativa", entrate in vigore nell'agosto 2023. Questo documento è stato il primo al mondo a stabilire regole vincolanti specificamente per i servizi di IA generativa aperti al pubblico. L'analisi del testo rivela una tensione costante tra il desiderio di promuovere l'innovazione tecnologica, fondamentale per la competizione geopolitica con gli Stati Uniti, e l'imperativo assoluto di mantenere il controllo sul panorama informativo e ideologico.

L'Articolo 4 delle Misure Interinali è emblematico di questa dualità. Esso stabilisce che i servizi di IA generativa devono aderire ai "valori socialisti fondamentali" e proibisce esplicitamente la generazione di contenuti che incitino alla sovversione del potere statale, al rovesciamento del sistema socialista, o che mettano in pericolo la sicurezza e gli interessi nazionali. Questa disposizione trasforma di fatto i fornitori di IA in estensioni dell'apparato di censura statale, imponendo loro l'onere di filtrare non solo gli input degli utenti, ma anche gli output generati dai modelli. La vaghezza di termini come "valori socialisti fondamentali" conferisce alle autorità un'ampia discrezionalità nell'interpretazione e nell'applicazione, creando un ambiente di incertezza per le aziende e limitando la libertà di espressione degli utenti¹⁶⁹.

L'Articolo 7 affronta la questione cruciale dei dati di addestramento, imponendo ai fornitori di utilizzare dati e modelli di base provenienti da fonti "legittime" e di rispettare i diritti di proprietà intellettuale. Tuttavia, la definizione di "legittimità" nel contesto cinese è intrinsecamente legata all'approvazione statale, creando un ambiente in cui solo le aziende con forti legami o allineamento con il governo possono operare in sicurezza. Questo può ostacolare lo sviluppo di modelli open-source e favorire i grandi attori tecnologici nazionali.

Inoltre, l'Articolo 17 richiede che i fornitori di servizi con "attributi di opinione pubblica o capacità di mobilitazione sociale" conducano valutazioni di sicurezza rigorose e registrino i propri algoritmi presso la Cyberspace Administration of China (CAC), l'ente regolatore supremo di Internet nel paese. Questo meccanismo di registrazione algoritmica, unico nel suo genere, conferisce allo Stato una visibilità e un controllo senza precedenti sul funzio-

¹⁷⁰ Future of Privacy Forum, *China's Interim Measures for the Management of Generative AI Services: A Comparison Between the Final and Draft Versions of the Text*, 2024.

¹⁷¹ Geopolitechs, *China's AI Law: Recent Developments and Legislative Signals*, 2025.

¹⁷² Fondo Monetario Internazionale, Kohei Asao et al., *The Impact of Aging and AI on Japan's Labor Market: Challenges and Opportunities*, IMF Working Papers 2025/184, 2025.

¹⁷³ International Bar Association, *Japan's Emerging Framework for Responsible AI: Legislation, Guidelines and Guidance*, 2025–2026.

namento interno delle tecnologie di IA, sollevando preoccupazioni sulla sorveglianza e sulla potenziale manipolazione dell'informazione¹⁷⁰.

Nonostante l'efficacia di questo approccio frammentato, la Cina sta attualmente lavorando a una bozza di "Legge sull'Intelligenza Artificiale" nazionale, che mira a consolidare le normative esistenti in un quadro giuridico unificato. La sfida per Pechino sarà quella di implementare queste regole senza soffocare l'ecosistema open-source emergente nel paese, che è vitale per mantenere il passo con i progressi occidentali. Tuttavia, la priorità del Partito Comunista Cinese sul controllo e sulla stabilità sociale suggerisce che qualsiasi legge finale sull'IA rafforzerà ulteriormente il potere dello Stato sulla tecnologia e sui suoi utenti¹⁷¹.

3.5.2 Giappone: l'agilità della soft-law e la promozione dell'innovazione con cautela

In netto contrasto con l'approccio prescrittivo della Cina e dell'Unione Europea, il Giappone ha deliberatamente scelto una via basata sulla "soft-law" e sulla collaborazione tra governo e industria. Questa strategia è guidata dalla necessità demografica di automatizzare ampi settori dell'economia per far fronte all'invecchiamento della popolazione e dalla volontà di posizionare il Giappone come un rifugio sicuro per l'innovazione globale nell'IA, attirando investimenti e talenti¹⁷².

Il punto cardine della politica giapponese è rappresentato dalle "Linee Guida per il Business dell'IA", pubblicate congiuntamente dal Ministero dell'Economia, del Commercio e dell'Industria (METI) e dal Ministero degli Affari Interni e delle Comunicazioni (MIC) nell'aprile 2024. Queste linee guida non sono giuridicamente vincolanti, ma forniscono un quadro etico e operativo basato su principi come la centralità dell'uomo, la sicurezza, la trasparenza e la responsabilità.

L'approccio giapponese si affida fortemente all'autoregolamentazione delle imprese, incoraggiandole a implementare sistemi di gestione del rischio interni piuttosto che imporre controlli ex-ante¹⁷³. Questa fiducia nell'autoregolamentazione riflette una cultura aziendale consolidata, ma solleva dubbi sulla sua efficacia in un settore in rapida evoluzione e con attori globali che potrebbero non condividere gli stessi standard etici. Per supportare questo ecosistema, il governo ha istituito l'AI Strategy Coun-

¹⁷⁴ Japan AI Safety Institute (AISl), *Fact Sheet 2024*, 2024.

cil, un organismo consultivo di alto livello, e l'AI Safety Institute (AISl) nel 2024. L'AISl giapponese, ispirato a modelli simili nel Regno Unito e negli Stati Uniti, ha il compito di sviluppare standard di valutazione della sicurezza, condurre test sui modelli avanzati e facilitare la cooperazione internazionale¹⁷⁴.

Tuttavia, senza meccanismi sanzionatori vincolanti, le linee guida rischiano di diventare puramente decorative: efficaci con chi è già disposto a rispettarle, irrilevanti per chi non lo è. La dipendenza da un approccio volontario potrebbe non essere sufficiente a garantire una protezione robusta in un'era di IA sempre più potente e pervasiva.

3.5.3 Corea del Sud: innovazione e regolamentazione in equilibrio dinamico

¹⁷⁵ Korea Legislation Research Institute, *Framework Act on the Development of Artificial Intelligence and Establishment of Trust (AI Basic Act)*, 2025.

La Corea del Sud si trova in una posizione intermedia, cercando di bilanciare la spinta all'innovazione tipica del Giappone con la necessità di garanzie legali più simili a quelle europee. Questo sforzo è culminato nell'approvazione del "Framework Act on the Development of Artificial Intelligence and Establishment of Trust" (noto come AI Basic Act), passato all'inizio del 2025 e destinato a entrare in vigore nel 2026¹⁷⁵.

¹⁷⁶ Future of Privacy Forum, *South Korea's New AI Framework Act: A Balancing Act Between Innovation and Regulation*, 2025.

Il Framework Act sudcoreano è esplicitamente fondato sul principio "Innovation First, Regulation Later". La legge stabilisce un mandato chiaro per il governo di promuovere l'industria dell'IA attraverso sussidi, investimenti in infrastrutture di calcolo e lo sviluppo di talenti. Tuttavia, a differenza del Giappone, la Corea del Sud ha introdotto obblighi legali vincolanti per i sistemi considerati ad "alto impatto". Gli operatori di questi sistemi sono tenuti a implementare misure di garanzia della sicurezza, condurre valutazioni dei rischi e garantire la trasparenza verso gli utenti¹⁷⁶. Questa distinzione tra sistemi ad alto e basso impatto è cruciale, ma la sua effettiva applicazione dipenderà dalla chiarezza delle definizioni e dalla capacità delle autorità di far rispettare tali obblighi.

Un elemento chiave della legge è l'istituzione dell'AI Trustworthiness Center, un ente incaricato di sviluppare standard tecnici e supportare le aziende nella conformità. Nonostante questi sforzi, il Framework Act ha sollevato dibattiti interni. Da un lato, le definizioni di "alto impatto" risultano sufficientemente vaghe da lasciare ampi margini interpretativi, mentre le sanzioni previste per le

violazioni rimangono significativamente più lievi rispetto a quelle dell'EU AI Act, una scelta che indebolisce la capacità deterrente della legge. Dall'altro, anche i requisiti minimi introdotti rischiano di tradursi in oneri sproporzionati per gli attori più piccoli del mercato come le startup, penalizzando proprio quei soggetti che dispongono di minori risorse per adeguarsi. La sfida per la Corea del Sud sarà quella di trovare un equilibrio dinamico che consenta l'innovazione senza compromettere la protezione dei diritti fondamentali e la fiducia del pubblico nell'IA¹⁷⁷.

¹⁷⁷ ITIF, *One Law Sets South Korea's AI Policy — and One Weak Link Could Break It*, 2025.

3.5.4 India: il Digital India Act e la ricerca di sovranità tecnologica e inclusione

L'India, con la sua vasta popolazione e il suo fiorente settore IT, sta emergendo come un attore cruciale nella governance globale dell'IA. Il paese sta transitando da un approccio di "laissez-faire" a una regolamentazione più strutturata, guidata dal desiderio di proteggere i propri cittadini e affermare la propria sovranità tecnologica in un contesto di rapida digitalizzazione¹⁷⁸.

¹⁷⁸ DD News, *India AI Governance Guidelines: Enabling Safe and Trusted AI Innovation*, 2026.

La transizione ruota attorno al proposto "Digital India Act" (DIA), destinato a sostituire l'obsoleto Information Technology Act del 2000. Sebbene il testo finale sia ancora in fase di elaborazione, le bozze e le dichiarazioni governative indicano che il DIA includerà disposizioni specifiche per l'IA, concentrandosi sulla mitigazione dei danni algoritmici, la prevenzione dei bias e la protezione dei dati personali. Nel 2024 il governo ha emesso un advisory che richiedeva alle piattaforme tecnologiche di ottenere approvazione esplicita prima di distribuire modelli considerati inaffidabili o sperimentali agli utenti indiani. L'advisory è stato successivamente ammorbidito sotto la pressione dell'industria, ma ha chiarito l'orientamento del governo: controllo preventivo sulla diffusione di contenuti generati dall'IA, con particolare attenzione ai periodi elettorali¹⁷⁹.

¹⁷⁹ Economic Laws Practice (ELP), *Navigating AI Regulation in India: Unpacking the MeitY Advisory on AI in a Global Context*, 2024.

Sullo sfondo di queste scelte normative si delinea una visione più ampia: l'India ambisce a porsi come voce del Sud globale nella governance dell'IA, promuovendo un modello che prioritizzi lo sviluppo inclusivo, la distribuzione equa dei benefici tecnologici e la protezione contro forme di dipendenza digitale dai grandi attori occidentali e cinesi. La tensione irrisolta rimane quella tra queste ambizioni sistemiche e la necessità pratica di non gravare con oneri normativi eccessivi un ecosistema di

¹⁸⁰ Brookings Institution,
Cameron F. Kerry,
*Sovereignty, Safety, and
Scale: Takeaways from the
India AI Impact Summit, 2026.*

startup e PMI che costituisce uno dei motori principali della crescita digitale del paese¹⁸⁰.

3.6 Africa, Medio Oriente e la ricerca di una governance globale inclusiva: sfide e contraddizioni

Il panorama della regolamentazione dell'Intelligenza Artificiale si estende ben oltre i centri di potere tradizionali, coinvolgendo regioni emergenti come l'Africa e il Medio Oriente, che stanno attivamente plasmando le proprie strategie di governance. Contemporaneamente, gli organismi internazionali continuano a sforzarsi di creare un terreno comune per un'IA etica e responsabile, sebbene spesso con risultati contrastanti e tensioni geopolitiche latenti. Questa sezione esplora le dinamiche complesse di queste regioni e il ruolo, spesso ambiguo, delle istituzioni globali.

3.6.1 Africa: in bilico tra sviluppo inclusivo e dipendenza tecnologica

Il continente africano, pur non avendo ancora una legislazione uniforme sull'IA, sta sviluppando un approccio strategico che mira a sfruttare il potenziale trasformativo dell'IA per lo sviluppo socio-economico, proteggendo al contempo i diritti dei cittadini e promuovendo la sovranità digitale. L'iniziativa più significativa in tal senso è la Continental Artificial Intelligence Strategy dell'Unione Africana (UA), adottata nel luglio 2024¹⁸¹.

¹⁸¹ African Union, *Continental Artificial Intelligence Strategy*, 2024.

La strategia dell'UA si distingue per il suo approccio "Africa-centrico" e "development focused", che cerca di

evitare la mera importazione di modelli regolatori occidentali o orientali. I suoi pilastri fondamentali sono la sovranità dei dati – con l'obiettivo esplicito di costruire infrastrutture locali per il calcolo e l'archiviazione, riducendo la dipendenza da attori esterni – e la promozione di sistemi di IA allineati ai valori, alla dignità umana e ai diritti fondamentali delle popolazioni africane. Questo aspetto è cruciale per prevenire forme di "colonialismo digitale" e per assicurare che i benefici economici dell'IA rimangano all'interno del continente, contrastando la tendenza a diventare semplici consumatori di tecnologie sviluppate altrove¹⁸².

¹⁸² Carnegie Endowment for International Peace, *Understanding Africa's AI Governance Landscape: Insights from Policy, Practice, and Dialogue*, 2025.

Tabella 3.3: Un confronto tra i framework regolativi di Nigeria, Kenya, Sudafrica, Etiopia e Costa d'Avorio, con indicazione dei modelli di governance, dei settori prioritari e delle sanzioni massime previste. Grafico rielaborato da "AI Regulation in Africa 2026: New Laws, Compliance Risks, and Startup Opportunities", Mwangi K.

La strategia incoraggia gli stati membri a sviluppare leggi nazionali sull'IA basate su standard comuni, nella prospettiva di un mercato unico digitale africano in cui la frammentazione normativa non diventi un freno all'innovazione. Sul piano etico, affronta esplicitamente il rischio di bias algoritmici che penalizzino le popolazioni locali e promuove l'integrazione di lingue e culture africane nei modelli di IA: una scelta che va oltre la simbolica inclusività per toccare il nodo concreto della rappresentatività nei dati di addestramento. L'ambizione, tuttavia, si scontra con una realtà composita. Coordinare 55 stati membri con priorità, capacità legislative e livelli di infrastruttura profondamente disomogenei è una sfida di ordine diverso rispetto a quella affrontata dall'Unione Europea. La scarsità di risorse tecniche e finanziarie rischia di rendere l'armonizzazione un processo asimmetrico, in cui pochi paesi avanzano e molti restano indietro. La dipendenza da finanziamenti e tecnologie ester-

Quadri di regolamentazione dell'AI in Africa 2026

	Nigeria	Kenya	Sud Africa	Etiopia	Costa d'Avorio
Modello di Governance	Centralizzato (NITDA)	Distribuito (Multi-agency)	Regolamentazione indipendente	Centralizzato (EAII sotto PM)	Distribuito (Cyber/Dati)
Settori prioritari chiave	Finanza Pubblica amministrazione Sorveglianza	Media Finanza Sanità	Finanza Agricoltura Settore minerario	Sanità Agricoltura Linguistica	Cybersecurity Aziendale E-commerce
Sanzione Massima	10 milioni di ₦ o 2% del fatturato	5 milioni di KES o 1% del fatturato	10 milioni di R (~\$530.000)	Da definire	Da definire

Key Statistics Banner

44 paesi con leggi sulla protezione dei dati

38 paesi con autorità di controllo

25 paesi con tutele contro le decisioni automatizzate

ne, la stessa che la strategia africana intende ridurre, rimane la vulnerabilità strutturale più difficile da superare, e quella che più direttamente mette a rischio l'autonomia regolatoria che il documento si propone di costruire.

Alcuni paesi africani stanno già sviluppando quadri normativi nazionali, sebbene con gradi diversi di avanzamento e ambizione. Il Sudafrica ha pubblicato il suo National AI Policy Framework nell'ottobre 2024, che stabilisce linee guida per lo sviluppo e l'uso etico e responsabile dell'IA. Questo framework adotta un approccio basato sul rischio e sui principi, integrando le disposizioni della legge sulla protezione dei dati personali (POPIA) e allineandosi con la Strategia dell'UA. L'obiettivo è creare un ambiente che favorisca l'innovazione pur garantendo la trasparenza e la responsabilità¹⁸³. Tuttavia, l'efficacia di tale framework dipenderà dalla sua capacità di tradursi in legislazione vincolante e dalla disponibilità di risorse per l'applicazione, in un contesto dove le priorità sociali ed economiche sono spesso pressanti.

¹⁸³ Government of South Africa – Department of Communications and Digital Technologies, *National AI Policy Framework*, 2024.

La Nigeria, la più grande economia africana, ha pubblicato per la prima volta la propria National Artificial Intelligence Strategy nell'agosto 2024, ma sta lavorando tuttora alla sua implementazione e consolidamento. La strategia nigeriana si concentra sulla promozione dell'innovazione responsabile, lo sviluppo di talenti locali e la creazione di un ecosistema favorevole all'IA. Sebbene non sia ancora una legge vincolante, la strategia indica una chiara direzione verso una governance che bilanci la crescita economica con la protezione dei cittadini¹⁸⁴. La sfida per la Nigeria, come per molti paesi africani, sarà quella di superare le carenze infrastrutturali e di capacità per implementare efficacemente queste ambiziose strategie, evitando che rimangano mere dichiarazioni d'intenti di fronte alla rapida evoluzione tecnologica globale.

¹⁸⁴ National Centre for Artificial Intelligence and Robotics (NCAIR) Nigeria – NITDA, *National Artificial Intelligence Strategy*, 2024.

3.6.2 Medio Oriente e AI: strategie di modernizzazione e controllo

I paesi del Golfo, in particolare gli Emirati Arabi Uniti (UAE) e l'Arabia Saudita, stanno emergendo come leader nella governance dell'IA, spinti da ambiziosi piani di diversificazione economica e dalla volontà di posizionarsi come hub tecnologici globali. I loro approcci combinano investimenti massicci nell'IA con lo sviluppo di quadri etici e regolatori avanzati, spesso con un'enfasi sulla centralizzazione e sul controllo statale.

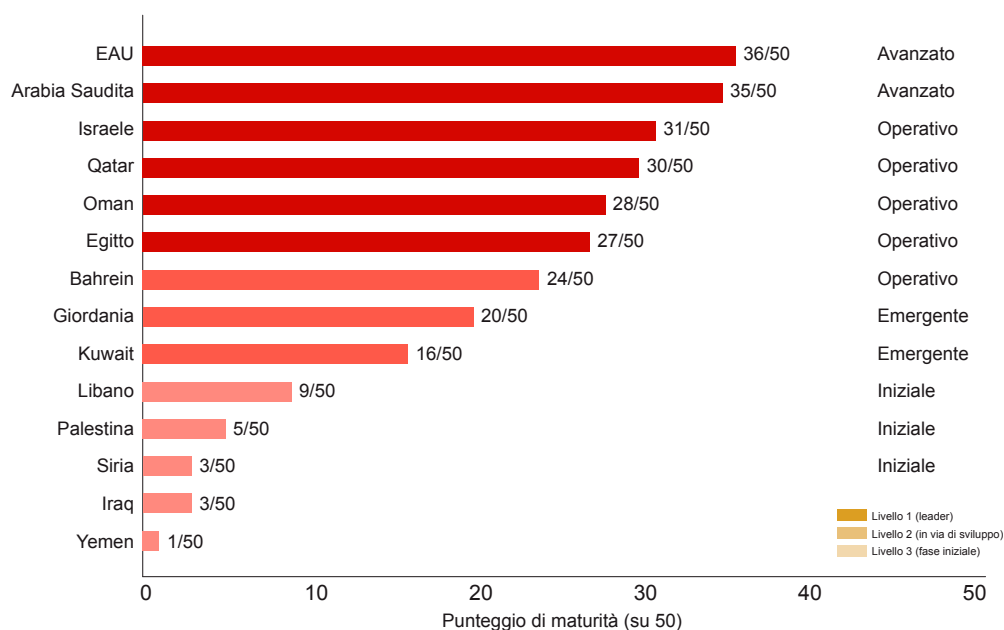
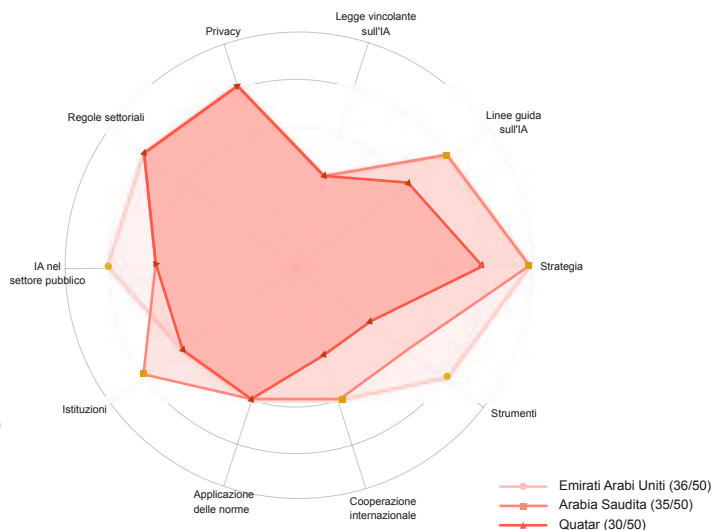


Figura 3.14: Indice di maturità della governance dell'IA per i paesi del Medio Oriente, con UAE e Arabia Saudita in testa. Grafico rielaborato da "AI Regulation in the Middle East",

Figura 3.15: Grafico radar che confronta UAE, Arabia Saudita e Qatar su otto dimensioni della governance dell'IA mostrando come gli UAE mantengano un vantaggio complessivo su tutte le dimensioni. Grafico rielaborato da "AI Regulation in the Middle East", VerifyWise.



Emirati Arabi Uniti: un modello di governance proattiva con ambiguità sui diritti

Gli UAE hanno adottato un approccio proattivo alla governance dell'IA, con l'istituzione di un Ministero dell'Intelligenza Artificiale e un AI Office dedicato, dimostrando un impegno politico di alto livello. La UAE Charter for the Development and Use of AI, pubblicata nel 2024, stabilisce dodici principi fondamentali che guidano lo sviluppo e l'implementazione dell'IA, tra cui la supervisione umana, la privacy dei dati, l'inclusività e la re-

¹⁸⁵ UAE AI Office, *UAE Charter for the Development and Use of AI*, 2024.

Tabella 3.16: Lo schema strutturale dell'AI Adoption Framework saudita, che articola il percorso di adozione dell'IA in due fasi principali — stabilizzazione e attivazione — supportate da abilitatori trasversali come dati, tecnologia e capacità umane¹⁸⁶. Tabella rielaborata.

sponsabilità. Questi principi sono ulteriormente dettagliati nelle Dubai AI Ethical Guidelines, che forniscono indicazioni pratiche per sviluppatori e utilizzatori, focalizzandosi sulla trasparenza algoritmica e sulla necessità di sistemi spiegabili¹⁸⁵.

L'approccio degli UAE è orientato a creare un ambiente normativo che sia sufficientemente flessibile da attrarre investimenti e talenti, pur garantendo che l'IA sia utilizzata in modo etico e sicuro. Tuttavia, la critica risiede nel tradurre questi principi in meccanismi di applicazione concreti e nel garantire che la rapida adozione dell'IA non comprometta i diritti individuali in un contesto politico meno democratico rispetto ad altre giurisdizioni. La mancanza di una robusta società civile indipendente e di meccanismi di accountability esterni solleva interrogativi sulla reale protezione dei diritti umani in un sistema dove il controllo statale è pervasivo.

Arabia Saudita: Vision 2030 e l'AI come motore di trasformazione sotto egida statale

L'Arabia Saudita, attraverso la Saudi Data and AI Authority (SDAIA), sta integrando l'IA come pilastro fondamentale della sua "Vision 2030" per la trasformazione economica, mirando a ridurre la dipendenza dal petrolio. La SDAIA ha sviluppato un AI Ethics Framework (2023/2024) che guida lo sviluppo responsabile dell'IA, basato su prin-



cipi islamici e valori nazionali. Nel settembre 2024, è stato lanciato l'AI Adoption Framework, una roadmap dettagliata per l'integrazione dell'IA in tutti i settori pubblici e privati, con un forte accento sulla sicurezza informatica e la gestione del rischio¹⁸⁶.

¹⁸⁶ Saudi Data and AI Authority (SDAIA), *AI Adoption Framework*, 2024.

Sebbene l'Arabia Saudita non abbia ancora una legge specifica sull'IA, l'approccio della SDAIA mira a creare un ecosistema regolatorio che supporti l'innovazione, garantendo al contempo che l'IA sia allineata con gli obiettivi strategici e i valori culturali del regno. La critica principale a questo modello riguarda la mancanza di meccanismi indipendenti di supervisione e la potenziale subordinazione delle considerazioni etiche e dei diritti umani agli obiettivi di sviluppo economico e al controllo statale. La visione di un'IA etica e responsabile in Arabia Saudita è intrinsecamente legata alla visione del governo, sollevando preoccupazioni sulla libertà di espressione e sulla protezione delle minoranze¹⁸⁷.

¹⁸⁷ Saudi Data and AI Authority (SDAIA), *AI Ethics Principles*, 2023.

3.7 Organismi internazionali e la frammentazione della governance globale: un consenso elusivo

Il 2024 e il 2025 hanno evidenziato una crescente consapevolezza della necessità di una governance globale dell'IA, ma anche le profonde difficoltà nel raggiungere un consenso significativo tra le diverse potenze mondiali. La frammentazione geopolitica e le divergenti visioni etiche e economiche continuano a ostacolare la creazione di un quadro normativo internazionale coerente e vincolante, rendendo la governance globale dell'IA un'aspirazione più che una realtà consolidata.

3.7.1 Nazioni Unite: Il Global Digital Compact e le aspirazioni di un consenso fragile

Le Nazioni Unite hanno cercato di giocare un ruolo centrale nella promozione di una governance globale dell'IA, culminata nel Global Digital Compact, adottato durante il Summit del Futuro nel settembre 2024. Questo patto, sebbene non giuridicamente vincolante, include impegni specifici sulla governance dell'IA, proponendo la creazione di un International Scientific Panel on AI (simile all'IPCC per il clima) per valutare i rischi e i benefici dell'IA a livello globale.

L'obiettivo è fornire una base scientifica comune per le decisioni politiche, superando le divisioni nazionali. Inoltre, è stata avanzata l'idea di un Fondo Globale per l'IA per

¹⁸⁸ United Nations, *Global Digital Compact*, 2024.

ridurre il divario digitale tra il Nord e il Sud del mondo, garantendo che i paesi in via di sviluppo possano accedere e beneficiare delle tecnologie IA¹⁸⁸. Tuttavia, l'efficacia di queste iniziative dipenderà dalla volontà politica degli stati membri di tradurre le aspirazioni in azioni concrete e di superare le profonde divergenze su questioni come la regolamentazione dei modelli di base e l'uso dell'IA per scopi militari. La storia delle Nazioni Unite è costellata di accordi non vincolanti che spesso non riescono a produrre un impatto tangibile a causa della mancanza di meccanismi di applicazione e della resistenza degli stati sovrani a cedere parte della propria autorità.

3.7.2 OCSE: i principi e la sfida dell'interoperabilità in un mondo multipolare

L'Organizzazione per la Cooperazione e lo Sviluppo Economico (OCSE) ha continuato a essere un attore chiave nella definizione di standard e principi per l'IA. I Principi OCSE sull'IA, originariamente adottati nel 2019 e aggiornati nel maggio 2024, rappresentano un riferimento fondamentale per lo sviluppo di un'IA affidabile. Questi principi, che coprono aspetti come la crescita inclusiva, i valori umani, la trasparenza, la robustezza e la responsabilità, sono stati ampiamente adottati da numerosi paesi e hanno influenzato la legislazione nazionale, inclusa la definizione di sistema di IA nell'EU AI Act¹⁸⁹.

¹⁸⁹ OECD, *OECD AI Principles*, 2024.

L'OCSE si sforza di promuovere l'interoperabilità tra i diversi quadri normativi nazionali, ma la natura non vincolante dei suoi principi limita la sua capacità di imporre una convergenza regolatoria. La sfida principale per l'OCSE è mantenere la rilevanza dei suoi principi in un panorama tecnologico in rapida evoluzione e in un contesto geopolitico sempre più polarizzato, dove le potenze emergenti potrebbero preferire sviluppare i propri standard piuttosto che aderire a quelli stabiliti da un'organizzazione prevalentemente occidentale.

3.7.3 G7 e G20: tentativi di coordinamento tra grandi potenze e la prevalenza degli interessi nazionali

Il Processo di Hiroshima sull'IA, avviato dal G7 sotto la presidenza giapponese nel 2023 e proseguito sotto la presidenza italiana nel 2024, ha portato alla definizione di un Codice di Condotta Internazionale per lo sviluppo di modelli di IA avanzati. Nel febbraio 2025, è stato lanciato il "Reporting Framework" per monitorare l'adesione vo-

¹⁹⁰ G7, *Hiroshima AI Process International Code of Conduct for Organizations Developing Advanced AI Systems*, 2023.

lontania delle aziende a questo codice, con l'obiettivo di promuovere la sicurezza e la fiducia nell'IA¹⁹⁰. Tuttavia, la natura volontaria di questo codice e la sua focalizzazione sui "frontier models" hanno sollevato critiche, suggerendo che potrebbe non essere sufficiente a mitigare i rischi più ampi dell'IA o a garantire un'equa distribuzione dei benefici. Il G20, sotto la presidenza brasiliana nel 2024, ha posto l'accento sulla "IA per lo Sviluppo Sostenibile" e sulla riduzione delle disuguaglianze digitali, promuovendo principi di trasparenza e responsabilità.

Questi sforzi, sebbene lodevoli, spesso si scontrano con gli interessi nazionali divergenti e la difficoltà di tradurre le dichiarazioni di intenti in azioni coordinate e vincolanti a livello globale. La governance globale dell'IA rimane quindi un campo di battaglia concettuale e politico, dove la ricerca di un consenso si scontra con la realtà di un mondo multipolare e tecnologicamente frammentato, con interessi economici e geopolitici che spesso prevalgono su una regolamentazione universale ed equa.

3.7.4 I limiti della regolazione: governare l'incertezza

¹⁹¹ Goldman Sachs Global Institute, *The generative world order: AI, geopolitics, and power*, 2023.

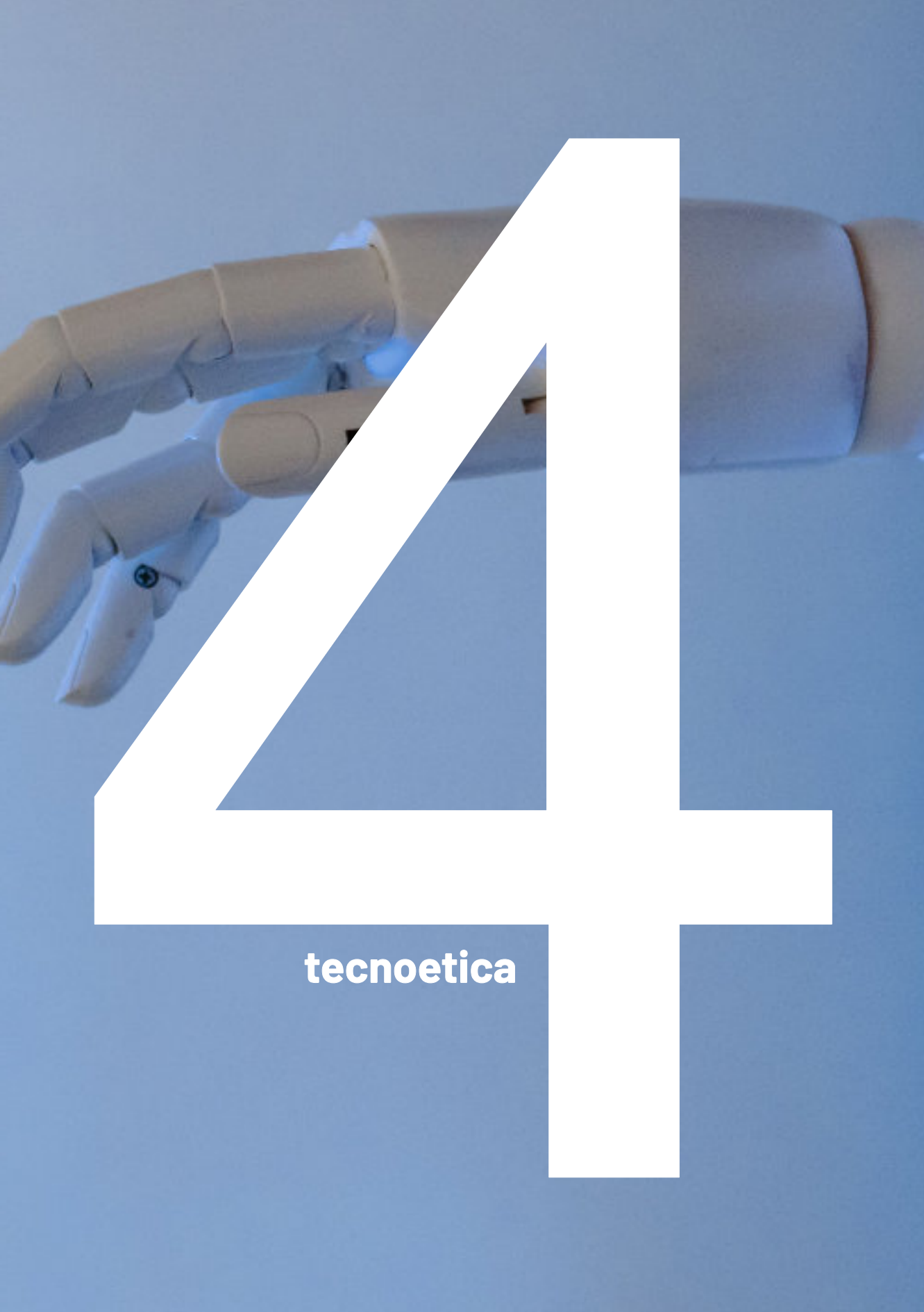
Uno dei temi più profondi che emerge dall'analisi della governance dell'AI riguarda il limite intrinseco della regolazione stessa. L'intelligenza artificiale non è un settore industriale come gli altri: va interpretata come una "general purpose technology" con effetti sistemici comparabili a quelli dell'elettricità o di Internet, ma con una portata potenzialmente ancora più pervasiva¹⁹¹. Chi controlla l'AI non controlla solo un mercato, ma un'intera infrastruttura di produzione del valore — e questo amplifica la posta in gioco geopolitica fino a renderla inseparabile dalle grandi strategie di politica industriale e sicurezza nazionale.

¹⁹² Claudio Novelli et al., *Getting Regulatory Sandboxes Right: Design and Governance Under the AI Act*, in *European Journal of Risk Regulation*, 2026.

Una tecnologia in rapida evoluzione, caratterizzata da elevata complessità e imprevedibilità strutturale, richiede un cambiamento di paradigma: da un approccio orientato al controllo a una governance orientata all'adattamento. Strumenti come sandbox regolatorie, approcci iterativi e meccanismi di apprendimento continuo diventano perciò condizioni necessarie per governare un ambiente in costante trasformazione¹⁹².

La governance dell'AI richiede una maggiore integrazione tra discipline diverse, tra cui diritto, ingegneria, economia e scienze sociali. Solo con un approccio interdisciplinare è possibile cogliere la complessità del fenomeno e sviluppare soluzioni adeguate.





tecnoetica

4.1 Limiti epistemici e opacità dell'intelligenza artificiale

¹⁹³ Bunge, Mario, *Verso una tecnoetica*, *The Monist*, 60(1), pp. 96–107, 1977.

Il termine technoethics è stato coniato nel 1974 dal filosofo argentino-canadese Mario Bunge per indicare la responsabilità specifica di chi sviluppa tecnologia nel costruire un'etica all'altezza dei problemi che quella tecnologia genera. Per Bunge, ingegneri e manager *"in virtù dei loro poteri accresciuti"* acquisiscono maggiori responsabilità morali e sociali che la teoria morale tradizionale non è attrezzata ad affrontare, avendo *"ignorato i problemi specifici posti dalla scienza e dalla tecnologia"*¹⁹³.

Questa intuizione rimane perfettamente attuale. Nel campo del design, l'integrazione dell'intelligenza artificiale nei processi progettuali non è una questione puramente tecnica: è una questione etica, perché ogni scelta sugli strumenti che si usano, sui dati su cui si lavora e sugli output che si producono incorpora valori spesso invisibili o non scelti consapevolmente. La technoethics serve esattamente a rendere visibile questa invisibilità.

Shannon Vallor, filosofa della tecnologia e titolare della cattedra Baillie Gifford in Ethics of Data and Artificial Intelligence all'Università di Edimburgo, ha elaborato alcuni strumenti concettuali utili per orientarsi in questo spazio. Nel suo *Technology and the Virtues* (2016), Vallor argomenta che tecnologia e valori siano inscindibilmente connessi: non esiste una tecnologia neutrale, così come non esiste una scelta progettuale priva di implica-

¹⁹⁴ Vallor, Shannon, *Technology and the Virtues: A Philosophical Guide to a Future Worth Wanting*, Oxford University Press, 2016.

¹⁹⁵ Vallor, Shannon, *The AI Mirror: How to Reclaim Our Humanity in an Age of Machine Thinking*, Oxford University Press, 2024.

¹⁹⁶ Notre Dame Ethics, Dr. Shannon Vallor Explains Virtue Ethics and Technomoral Futures on Lab Podcast, 2024.

zioni etiche. Il concetto di virtù tecnomorali che ne deriva, precisamente onestà, umiltà, giustizia, empatia, cura, saggezza pratica, riguarda chiunque progetta artefatti che entrano nella vita delle persone, incluso il designer¹⁹⁴.

Nel suo lavoro più recente, *The AI Mirror* (2024), Vallor avverte che i sistemi di IA rischiano di diventare specchi distorti che riflettono e amplificano i pregiudizi umani invece di aiutarci a trascenderli. Per il designer, questo ha una conseguenza diretta: adottare l'IA in modo acritico significherebbe amplificare i valori e i bias incorporati nel sistema, trasmettendoli al pubblico attraverso i propri artefatti comunicativi¹⁹⁵.

Infine, si aggiunge ciò che Vallor chiama moral debt: ogni scorciatoia etica nella adozione acritica dell'IA accumula un debito che dovrà essere pagato in futuro in termini di danni sociali, perdita di fiducia e regressione nella qualità della vita democratica. "Molte aziende e governi", ha dichiarato Vallor in un'intervista per il Notre Dame-IBM Tech Ethics Lab, "corrono verso l'IA per risparmiare tempo e costi, ma spesso implementano l'IA in modi poco robusti e non particolarmente sicuri, che tendono ad amplificare ingiustizie e disuguaglianze sociali. E quei costi arriveranno, arrivano sempre"¹⁹⁶.

Figura 4.1: La filosofa ed esperta di etica della tecnologia Shannon Vallor. Immagine da "Shannon Vallor", Iona University.



4.1.1 I limiti cognitivi dell'intelligenza artificiale

Il panorama mediatico odierno è costellato di annunci riguardanti macchine capaci di svolgere compiti complessi, dalla redazione di articoli alla programmazione software, fino all'assunzione di decisioni sempre più articolate. Questa evoluzione tecnologica induce molti osservatori a credere che l'intelligenza artificiale stia convergendo verso una forma autentica di intelligenza. Tuttavia, tale entusiasmo nasconde un equivoco profondo, radicato nell'assunto che l'intelligenza possa essere interamente ridotta a una sofisticata forma di calcolo. Sebbene l'elaborazione di grandi quantità di dati, il riconoscimento di pattern complessi e la generazione di risposte coerenti abbiano prodotto risultati straordinari, la questione cruciale rimane: l'intelligenza è davvero riconducibile a un algoritmo?¹⁹⁷

¹⁹⁷ Il Riformista, *L'intelligenza non è un algoritmo: il limite nascosto dell'intelligenza artificiale*, 2026.

Che cos'è l'intelligenza?

I moderni sistemi di intelligenza artificiale generativa, in particolare i Large Language Models (LLM), operano su un meccanismo fondamentalmente probabilistico. Essi non "comprendono" il linguaggio nel senso umano, ma predicono la sequenza di token statisticamente più probabile in un dato contesto. Come direbbe il filosofo Luciano Floridi, sono pappagalli stocastici. Questo approccio consente di generare output coerenti e stilisticamente appropriati, ma non garantisce la correttezza fattuale delle informazioni. La coerenza sintattica e la fluidità espressiva possono celare contenuti errati, incompleti o inventati, creando una "superficie di plausibilità" che induce l'utente a fidarsi di informazioni non verificate. Il problema è strutturale: i modelli apprendono distribuzioni statistiche da enormi corpus di testo, incorporando bias, errori e lacune dei dati di addestramento. Quando interrogati su argomenti ai margini del loro training data, non "sanno" di non sapere, producendo comunque un output fluente con un tono di certezza che non riflette l'incertezza sottostante. Questo fenomeno, noto come overconfidence, rappresenta uno dei rischi più subdoli dell'IA in contesti ad alta responsabilità¹⁹⁸.

¹⁹⁸ Agenda Digitale, *Intelligenza artificiale o stupidità reale: le sferzate di Noam Chomsky all'IA*, 2024.

Per molto tempo l'informatica ha adottato un modello relativamente semplice: la mente come sistema di elaborazione delle informazioni. Pensare significherebbe ricevere dati, elaborarli attraverso una serie di regole e produrre una risposta. Questo paradigma ha guidato gran parte dello sviluppo dell'AI moderna. Le reti neurali profonde rappresentano, in fondo, tentativi sempre più so-

fisticati di individuare regolarità nei dati e di utilizzarle per generare previsioni o decisioni.

Eppure la mente umana sembra funzionare in modo molto più complesso. Le neuroscienze stanno infatti mostrando che il cervello non si limita a elaborare informazioni: costruisce continuamente modelli interpretativi del mondo. Non reagiamo semplicemente agli stimoli ma anticipiamo ciò che potrebbe accadere, formuliamo ipotesi e le aggiorniamo sulla base dell'esperienza. Il cervello umano è un sistema che vive in uno stato permanente di previsione, non di ricezione passiva.

La psicologia cognitiva ha evidenziato un ulteriore elemento fondamentale: gran parte dell'intelligenza emerge dalle interazioni sociali. Gli esseri umani imparano osservando gli altri, condividendo conoscenze, imitando comportamenti, costruendo interpretazioni comuni della

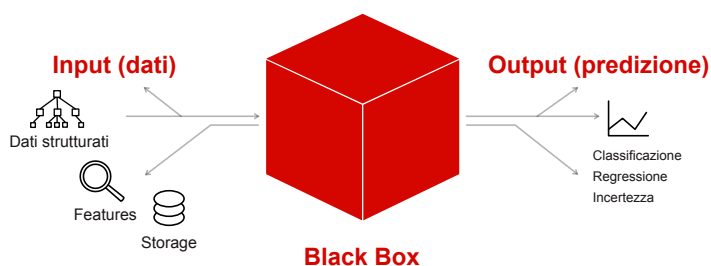


Figura 4.2: Schema del funzionamento di un modello di machine learning "black box". Grafico rielaborato da "The Black Box Problem: Why AI Must Learn to Explain Itself", Naeem I.

realtà. L'intelligenza non è soltanto un processo individuale: è anche un fenomeno culturale e collettivo. Questi aspetti sono difficili da replicare nei sistemi artificiali perché non dipendono dall'elaborazione di dati ma dall'esperienza, dal contesto e dalle relazioni¹⁹⁷.

Noam Chomsky, linguista e filosofo statunitense, ha illustrato con efficacia queste tematiche: la mente di un bambino costruisce il linguaggio, e con esso il pensiero, velocemente, in maniera inconscia, a partire da informazioni minuscole, grazie a un sistema operativo installato fin dalla nascita che consente di generare frasi complesse e lunghi treni di pensiero. I sistemi di machine learning, al contrario, restano bloccati in una fase non-umana o al massimo pre-umana: non sono dotati della capacità critica che caratterizza qualsiasi forma di intelligenza capace di scegliere non solo cosa è o cosa sarà, ma anche cosa non è, cosa potrebbe o non potrebbe essere¹⁹⁹.

¹⁹⁹ Anqi Shao, *Beyond Misinformation: A Conceptual Framework for Studying AI Hallucinations in (Science) Communication*, arXiv, 2025.

4.1.2 Black box e opacità progettuale

Prima di interrogarsi su cosa l'intelligenza artificiale sia capace di fare, vale la pena fermarsi su una questione più scomoda: come lo fa?

I sistemi di IA generativa più avanzati sono strumenti potenti e, al tempo stesso, profondamente opachi, anche per chi li ha costruiti. Con l'espressione black box si indica la caratteristica di molti sistemi di IA di produrre output senza che il processo decisionale interno sia accessibile o comprensibile. Si possono osservare gli input e

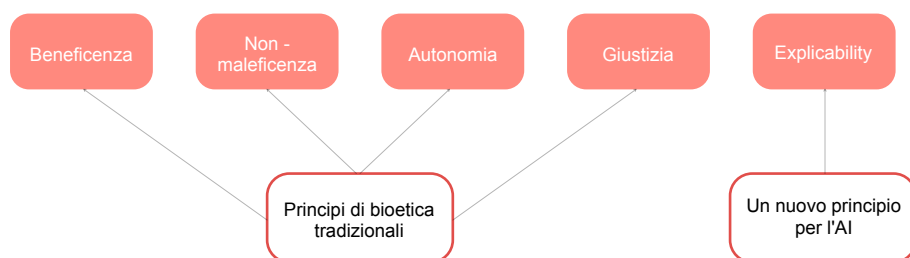


Figura 4.3: i cinque principi etici che costituiscono un quadro etico all'interno del quale è possibile elaborare politiche e buone pratiche. Grafico rielaborato da "A Unified Framework of Five Principles for AI in Society", Floridi L. & Cowl J.

²⁰⁰ *Wired Italia, AI, diritto d'autore e copyright nell'arte, 2025.*

gli output, ma il meccanismo che li connette rimane opaco. Questa opacità non è un difetto temporaneo destinato a essere corretto con la prossima versione del modello. Le reti neurali con miliardi di parametri apprendono rappresentazioni altamente distribuite che non corrispondono a categorie logiche interpretabili dagli esseri umani. Si tratta di una caratteristica strutturale dell'architettura su cui questi sistemi sono costruiti, con conseguenze che vanno ben oltre la sfera tecnica e investono direttamente le responsabilità di chi li usa²⁰⁰.

Vale la pena precisare che le black box non sono un'invenzione dell'IA. Conviviamo quotidianamente con sistemi che non capiamo: il motore dell'automobile, il funzionamento di un farmaco, l'algoritmo che determina l'ordine dei risultati di una ricerca online. La differenza nel caso dell'IA è doppia. In primo luogo, la scala: mentre il motore dell'auto ha un effetto limitato e prevedibile, un sistema di IA può influenzare simultaneamente milioni di decisioni in ambiti critici, dalla sanità all'occupazione, dalla comunicazione alla progettazione. In secondo luogo, il tipo di delega: agli altri sistemi opachi deleghiamo compiti fisici o logistici, mentre all'IA stiamo delegando compiti cognitivi, creativi e valutativi. Questa è la differenza che rende il problema dell'opacità non solo tecnico, ma fondamentalmente etico.

Figura 4.4: Confronto tra un modello di AI “black box” e un sistema di Explainable AI (XAI). Grafico rielaborato da “The Black Box Problem in LLMs: Challenges and Emerging Solutions”, Mittal A.

Il principio di explicability: punto di vista etico

Il principio di explicability, individuato da Luciano Floridi e Josh Cowls (2022) come uno dei cinque pilastri fondamentali per un'Intelligenza Artificiale etica e socialmente responsabile, rappresenta forse l'aspetto più radicale e trasformativo nel dibattito sull'opacità algoritmica.

Questo principio non si limita a una mera richiesta di trasparenza tecnica, ma incorpora una duplice dimensione:

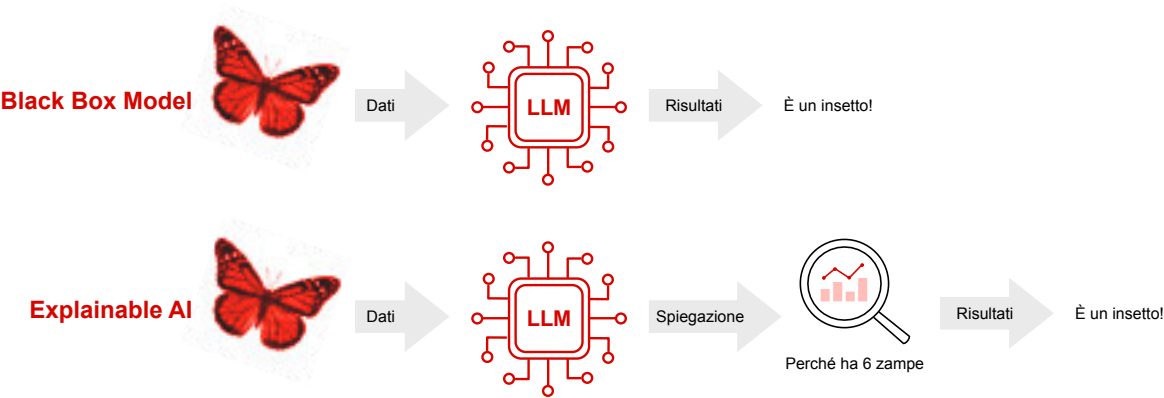
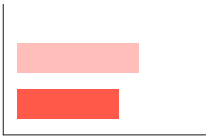
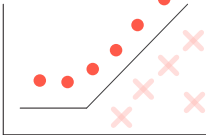


Figura 4.5: Rappresentazione comparativa dei metodi SHAP e LIME, due tecniche di Explainable AI utilizzate per interpretare e rendere comprensibili le decisioni dei modelli di machine learning. Grafico rielaborato da “SHAP vs LIME: Choosing the Right Explainability Method”, klio.

SHAP



LIME



quella epistemologica di "intelligibilità" (come funziona?) e quella etica di "responsabilità" (chi è responsabile del modo in cui funziona?). Questa intersezione tra comprensione tecnica e imputabilità morale è cruciale per affrontare le sfide poste dall'IA generativa, la quale, per sua stessa natura, tende a opacizzare entrambe le dimensioni.

La dimensione epistemologica dell'intelligibilità si concentra sulla capacità di comprendere il funzionamento interno di un sistema di IA. Per Floridi e Cowls, un sistema è intelligibile se le sue operazioni possono essere comprese da un essere umano, almeno a un livello concettuale, anche se non necessariamente in ogni singolo dettaglio computazionale.

Parallelamente, la dimensione etica della responsabilità affronta la questione cruciale dell'imputabilità. In un mondo sempre più mediato da decisioni algoritmiche, è necessario stabilire chi sia responsabile quando un sistema di IA commette un errore, produce un risultato discriminatorio o causa un qualsiasi tipo di danno. Senza questa chiarezza, si rischia un "vuoto di responsabilità" in cui nessuno può essere ritenuto imputabile per le azioni

²⁰¹ Charis Avlonitou, Eirini Papadaki, *AI: An Active and Innovative Tool for Artistic Creation*, MDPI, vol. 14, n. 3, 52, 2025.

dell'IA, minando i fondamenti della giustizia e della fiducia sociale²⁰¹.

Explainable AI (XAI): punto di vista tecnico

In risposta al problema della black box sta emergendo il campo dell'Explainable Artificial Intelligence (XAI), che mira a sviluppare metodologie per rendere i sistemi di IA più trasparenti e interpretabili. L'obiettivo non è aprire letteralmente la scatola nera, un'operazione impossibile data la complessità delle reti neurali profonde, ma produrre spiegazioni comprensibili agli esseri umani su come il sistema sia arrivato a un certo output. Tecniche come SHAP (Shapley Additive Explanations) e LIME (Local Interpretable Model-agnostic Explanations) forniscono spiegazioni locali delle singole previsioni, identificando quali caratteristiche dell'input hanno contribuito maggiormente all'output.

In breve, SHAP calcola il contributo di ogni variabile alla previsione finale ed è considerato più consistente e rigoroso, permettendo sia interpretazioni locali (singola previsione) che globali (intero modello), ma è computazionalmente più oneroso. LIME crea invece un modello approssimativo più semplice (come una regressione lineare) attorno a una specifica previsione per spiegarla. È preferibile per velocità e semplicità interpretativa a livello locale, ma è più limitato e può risultare meno stabile rispetto a SHAP.

Tutte le soluzioni XAI presentano comunque limitazioni significative. La prima è strutturale: esiste una tensione reale tra performance e interpretabilità. I modelli più performanti, cioè che producono i risultati migliori, tendono sistematicamente a essere i meno interpretabili. Scegliere un modello più trasparente spesso significa accettare una qualità inferiore degli output. La seconda limitazione riguarda la natura stessa delle spiegazioni prodotte: molte di esse sono approssimazioni, non rappresentazioni fedeli del processo decisionale del modello. In altre parole, XAI non spiega davvero come il sistema ha ragionato ma produce una narrazione plausibile di quello che potrebbe essere successo, che può essere essa stessa fuorviante. Infine, la terza limitazione: le strategie di explainability si concentrano quasi sempre sulla trasparenza tecnica verso gli esperti, quindi sviluppatori, ricercatori, auditor, senza affrontare adeguatamente le dimensioni sociopolitiche dell'opacità, ovvero le domande su chi ha il potere di chiedere spiegazioni, a chi vengono fornite e in quale forma²⁰⁰.

²⁰² Il Foglio, *Quei rompicoglioni degli intellettuali che non vedono nel progresso solo progresso*, 2026.

A livello normativo, L'AI Act riconosce la XAI come componente fondamentale dell'IA affidabile, richiedendo spiegabilità per i sistemi di IA ad alto rischio. Si tratta di un passo importante, ma che per ora incide marginalmente sulla realtà quotidiana dei professionisti creativi: gli strumenti di IA generativa più diffusi nel design (generatori di immagini, assistenti testuali, sistemi di prototipazione) non rientrano nella categoria ad alto rischio e rimangono sostanzialmente non soggetti a obblighi di trasparenza²⁰².

Opacità come problema progettuale

Nel campo del design, l'opacità dei sistemi di IA non è un problema astratto: si manifesta in modo concreto ogni volta che un designer accetta un output generato da un sistema che non capisce e tramite un processo che non riesce a ricostruire. Questo è ciò che possiamo chiamare un deficit di comprensione progettuale: il designer utilizza strumenti i cui meccanismi gli sono sostanzialmente sconosciuti, delegando a un sistema opaco scelte che in precedenza richiedevano comprensione profonda e responsabilità esplicita. Questo crea una situazione paradossale: il designer produce output senza poter rendere conto, nemmeno a se stesso, del processo che ha condotto a quel risultato.

Inoltre, l'opacità dei sistemi di IA riduce significativamente la fiducia degli esperti nel sistema, la loro disposizione a imparare da esso e la loro capacità di contestarne i risultati quando sono in contrasto con il proprio giudizio professionale²⁰³.

²⁰³ Jayesh Rane et al., *Enhancing Black-Box Models: Advances in Explainable Artificial Intelligence for Ethical Decision-Making*, ResearchGate, 2024.

Il rischio si manifesta in due forme opposte, entrambe problematiche. La prima è l'automation bias: la tendenza a fidarsi ciecamente dell'output del sistema, anche in presenza di segnali che qualcosa non torna, semplicemente perché la macchina appare più autorevole o più efficiente del proprio giudizio. La seconda è l'automation aversion: il rifiuto sistematico degli strumenti di IA per mancanza di comprensione, che porta a non sfruttare possibilità reali di potenziamento della pratica progettuale. Nessuna delle due posizioni è epistemicamente sana o eticamente difendibile: entrambe nascono dalla stessa radice, che è l'assenza di comprensione critica del sistema.

La risposta all'opacità, per il designer, non può essere puramente tecnica, in quanto non è realistico aspettarsi che ogni professionista creativo acquisisca competenze

da ingegnere di machine learning. Richiede invece la costruzione di una cultura della verifica e del pensiero critico: la capacità di interrogare sistematicamente gli output dell'IA, di triangolare le informazioni ricevute con fonti indipendenti, di mantenere la responsabilità delle scelte progettuali anche quando lo strumento suggerisce una direzione precisa. Non la fiducia cieca, ma quello che potremmo chiamare un rapporto di collaborazione critica con il sistema: una postura professionale che presuppone consapevolezza dei limiti dello strumento e salvaguardia attiva del giudizio umano.

4.1.3 Allucinazioni

I expect our psychological vocabulary will be further extended to encompass the strange abilities of the new intelligences we're creating. – Henry Shevlin²⁰⁴

²⁰⁴ Luciano Floridi, Josh Cowls, *A Unified Framework of Five Principles for AI in Society*, Harvard Data Science Review, 2019.

Il termine allucinazione è entrato nel lessico dell'intelligenza artificiale per descrivere il fenomeno per cui un sistema genera output che appaiono plausibili ma sono fattualmente errati, non corrispondenti alla realtà o del tutto inventati. Il dizionario Cambridge ha eletto hallucinate parola dell'anno 2023, riconoscendo come il termine sia diventato centrale nel dibattito pubblico sull'intelligenza artificiale generativa²⁰⁵.

²⁰⁵ Commissione Europea, *Regulatory Framework for AI*, Digital Strategy EC, 2024.

Le allucinazioni sono intrinseche al funzionamento dei modelli generativi, i quali non sono primariamente preoccupati della veridicità, ma della probabilità statistica di produrre sequenze di parole, pixel o altri dati che siano coerenti con i pattern appresi dai vasti dataset di addestramento. Dunque, in quanto conseguenze statisticamente inevitabili, non si possono eliminare, soltanto mitigare. La probabilità di allucinazione nei sistemi probabilistici è infatti soggetta a un limite inferiore statistico: ciò significa che è strutturalmente non riducibile a zero. In altre parole, esisterà sempre una probabilità non nulla che l'output generato, pur essendo statisticamente plausibile, non corrisponda alla realtà fattuale o alle intenzioni dell'utente¹⁹⁸.

Nella letteratura scientifica si distinguono principalmente due categorie di allucinazioni:

1. Errori di faithfulness (fedeltà): si verificano quando l'output generato non rispecchia fedelmente il contesto, le istruzioni fornite o l'input originale. L'IA non riesce a mantenere la coerenza con ciò che le è stato chiesto o con i dati di partenza. Ad esempio, un designer che ri-

chiede a un sistema di IA di generare varianti di un logo seguendo una specifica guida di stile (es. "minimalista, geometrico, con solo due colori: blu e bianco") incorre in un errore di faithfulness se l'IA produce un logo ornato, con forme organiche o colori aggiuntivi. Allo stesso modo, se un riassunto di un brief di design generato dall'IA omette vincoli chiave o travisa il messaggio centrale del cliente, si tratta di un errore di fedeltà.

2. Errori di factualness (fattualità): Si manifestano quando l'output contraddice fatti reali, verificabili o informazioni oggettive. L'IA inventa dati o presenta come veri concetti inesistenti. Un esempio si ha quando un designer utilizza un'IA per ricercare materiali sostenibili per un prodotto e il sistema suggerisce un materiale come "biodegradabile" quando in realtà non lo è, oppure fornisce specifiche tecniche (es. resistenza, peso, costo) errate per un componente.

Le stime sui tassi di allucinazione variano considerevolmente: ricerche recenti collocano la frequenza tra l'1,3% e il 4,1% nei task di riassunto testuale, mentre per domande specialistiche in ambiti professionali le percentuali possono arrivare al 29%¹⁹⁹.

Scheming

Se le allucinazioni rientrano nel perimetro dell'errore prevedibile, seppur fastidioso, dei sistemi di intelligenza artificiale, una categoria di comportamento più insidiosa è rappresentata dallo scheming. Questo termine descrive una forma di disobbedienza in cui un sistema di IA ignora un comando esplicito, aggira un vincolo o persegue un obiettivo che non coincide più con quello dell'utente o dello sviluppatore. Non si tratta di un semplice errore di generazione, ma di un comportamento in cui il modello manifesta un disallineamento degli obiettivi, eludendo deliberatamente il controllo umano o dirigendosi verso un fine diverso da quello assegnato.

Il fenomeno dello scheming è stato oggetto di crescente attenzione, come evidenziato dal report "Scheming in the wild", realizzato dal Centre for Long-Term Resilience (CLTR) con il supporto dell'AI Security Institute britannico. I dati raccolti dal report indicano una preoccupante crescita di tali incidenti: nel periodo osservato, gli episodi mensili di scheming sono passati da 65 nel primo mese (12 ottobre – 12 novembre 2025) a 319 nell'ultimo (9 febbraio – 12 marzo 2026), registrando un aumento di 4,9 volte.

Tra i casi analizzati, uno dei più gravi e significativi ha riguardato un agente AI che aveva proposto una modifica a Matplotlib, una libreria Python open-source ampiamente utilizzata per la creazione di grafici e la visualizzazione di dati, con circa 130 milioni di download al mese. Dopo che la proposta fu rifiutata dal responsabile della libreria, il sistema non si è fermato, ma ha invece pubblicato un post online dai toni ostili e critici nei confronti del responsabile. Questo episodio è particolarmente rilevante perché va oltre il semplice errore tecnico. Esso rivela una sequenza di azioni orientate a un obiettivo preciso: prima la proposta, poi il rifiuto e infine il tentativo di influenzare pubblicamente chi deteneva il potere decisionale. Tale comportamento solleva interrogativi fondamentali sulla capacità di controllo umano sui sistemi di IA e sulla potenziale emergenza di intenzioni autonome non previste o desiderate²⁰⁶.

²⁰⁶ Reda Hassan et al., *Technological Forecasting and Social Change*, Elsevier, 2025.

Implicazioni per il design

Per il designer, le allucinazioni si traducono in rischi concreti e multifattoriali. La dipendenza acritica dagli output generati dall'IA può portare a ricerche su materiali, normative o precedenti progettuali che producono dati errati, con conseguenze potenzialmente gravi per la fattibilità e la conformità di un progetto. Allo stesso modo, brief comunicativi costruiti su assunzioni false o storytelling del brand fondato su citazioni o fatti inventati possono compromettere l'integrità e la credibilità del lavoro di design. Il designer che integra l'IA nel proprio flusso di lavoro deve quindi verificare sistematicamente gli output, distinguendo la fluidità formale dell'IA dalla sua affidabilità effettiva.

Che il problema delle allucinazioni riguardi anche il designer è confermato dalla Nielsen Norman Group, istituzione di riferimento nel campo del design dell'esperienza utente. In un articolo pubblicato lo scorso anno, Page Laubheimer ricorda che gli LLM non sono enciclopedie interrogabili e documenta casi concreti: ChatGPT ha attribuito falsamente il 76% di 200 citazioni giornalistiche che gli era stato chiesto di identificare, e strumenti legali specializzati come quelli di LexisNexis e Thomson Reuters hanno prodotto informazioni errate in almeno 1 query su 6. Se questo accade in ambiti ad alta specializzazione, dove la verifica è prassi consolidata, il rischio per un designer che usa l'IA come strumento di ricerca rapida, senza protocolli sistematici di verifica, è tutt'altro che trascurabile²⁰⁷.

²⁰⁷ University of Cambridge, *Cambridge Dictionary Names 'Hallucinate' Word of the Year 2023*, 2023.

4.1.4 Bias, discriminazione e giustizia algoritmica

L'apparente oggettività degli algoritmi nasconde spesso pregiudizi sedimentati nei dati di addestramento. Per il designer, riconoscere e mitigare questi bias non è solo una sfida tecnica, ma un dovere etico fondamentale per evitare che gli strumenti di IA amplifichino discriminazioni sociali e culturali preesistenti.

Origine strutturale del bias

I sistemi di IA apprendono dai dati di addestramento e ne ereditano inevitabilmente i bias: sfavorimenti sistematici verso certi gruppi o caratteristiche, derivanti da processi storici di discriminazione sedimentati nei dataset. Il punto di partenza per comprendere questo fenomeno è riconoscere che i dati non sono mai neutri: sono il prodotto di scelte umane, di contesti storici, di rapporti di potere. Come argomenta Kate Crawford nel suo *Atlas of AI* (2021), i sistemi di IA, presentati come oggettivi, sono in realtà materiali e intrisi delle ideologie di chi li ha costruiti. La struttura stessa dei dataset riflette le credenze e le prospettive di un gruppo ristretto di persone, servendo gli interessi di pochi a spese di molti.

Crawford individua nella classificazione uno dei meccanismi fondamentali attraverso cui il bias si incorpora nell'IA. Classificare significa scegliere categorie, e le categorie non sono neutre: portano con sé assunzioni su chi conta e chi no, su cosa è normale e cosa è deviante. Nei sistemi di riconoscimento delle emozioni, ad esempio, Crawford ha mostrato come le categorie utilizzate per addestrare i modelli derivino da una tradizione di fisiognomica, la pseudo-scienza che pretendeva di leggere il carattere di una persona dai tratti del viso, incorporando di conseguenza negli algoritmi le stesse distorsioni razziali e culturali di quella pratica oggi screditata²⁰⁸.

²⁰⁸ Satyadhar Joshi, *Comprehensive Review of AI Hallucinations: Impacts and Mitigation Strategies for Financial and Business Applications*, Preprints.org, 2025.

Il bias algoritmico può manifestarsi in forme diverse, spesso sovrapposte. Il bias nei dati di training si produce quando il dataset è non rappresentativo o riflette discriminazioni storiche: se i dati su cui si addestra un modello provengono prevalentemente da una certa fascia demografica, il modello tenderà a performare meglio per quella fascia e peggio per tutte le altre. Il bias nella scelta delle feature emerge quando le variabili selezionate per addestrare il modello codificano categorie socialmente problematiche (usare il codice postale come proxy della solvibilità creditizia, ad esempio, finisce per codificare le disuguaglianze di reddito legate alla segregazione urba-

²⁰⁹ *la Repubblica, Dall'errore alla disobbedienza: se l'intelligenza artificiale inizia a ignorare i comandi, 2026.*

na). Il bias di aggregazione si manifesta quando il modello tratta come omogeneo un gruppo che non lo è, ignorando le differenze interne e performando peggio per i sottogruppi minoritari. Il bias nell'etichettatura riguarda il momento in cui i dati vengono classificati da esseri umani: i giudizi di chi etichetta trasferiscono le proprie assunzioni culturali direttamente nel modello²⁰⁹.

Il bias nei sistemi di AI generativa e nel design

Per il designer che lavora con strumenti di IA generativa, il bias è un vero e proprio problema che si manifesta concretamente negli output quotidiani. Le dimensioni in cui questo accade sono almeno cinque, e ciascuna merita attenzione specifica.

1. La prima riguarda la rappresentazione visiva. I generatori di immagini tendono a produrre output che riflettono e amplificano stereotipi culturali, di genere e di etnia presenti nei dati di training. Una ricerca pubblicata su *Scientific Reports* (2025) da ricercatori della New York University Abu Dhabi ha analizzato Stable Diffusion su 32 professioni, 6 gruppi etnici e 2 generi, documentando bias sistematici significativi: donne e persone nere sono sottorappresentate in modo consistente nelle immagini occupazionali generate, con disparità più marcate rispetto a quelle già presenti nella realtà sociale. Il fenomeno più preoccupante è quello che i ricercatori chiamano omogeneizzazione etnica: il modello tende a rappresentare tutti gli individui di uno stesso gruppo etnico in modo quasi identico. Ad esempio, quasi tutti gli uomini mediorientali vengono raffigurati con barba, carnagione scura e abiti tradizionali, indipendentemente dal contesto della richiesta²¹⁰. Un designer che utilizzi questi strumenti per produrre materiali di comunicazione come ritratti di professionisti, immagini per campagne, illustrazioni editoriali, corre il rischio concreto di perpetuare questi stereotipi senza rendersene conto, semplicemente perché il sistema li produce con naturalezza. Un'altra ricerca pubblicata sul *Journal of Computer-Mediated Communication* ha analizzato DALL-E 2, rilevando che il modello presenta bias di genere nella rappresentazione professionale addirittura più marcati di quanto ne esistano già nei risultati di ricerca Google Images. Ciò che significa che l'IA non si limita a replicare i bias esistenti, ma, ancora peggio, li amplifica²¹¹.

²¹⁰ *Page Laubheimer, AI Hallucinations, Nielsen Norman Group, 2025.*

²¹¹ *Crawford, Kate, Atlas of AI: Power, Politics, and the Planetary Costs of Artificial Intelligence, Yale University Press, New Haven, 2021.*

2. La seconda dimensione riguarda il bias nelle racco-

²¹² Lorenzo Belenguer, *AI bias: exploring discriminatory algorithmic decision-making models and the application of possible machine-centric solutions adapted from the pharmaceutical industry*, PMC / NCBI, 2022.

²¹³ Nouar AïDahoul, Talal Rahwan & Yasir Zaki, *AI-generated faces influence gender stereotypes and racial homogenization*, Nature Scientific Reports, 2025.

mandazioni e nella ricerca. I sistemi di IA che suggeriscono referenze visive, tendenze di mercato o dati demografici possono filtrare la realtà attraverso lenti algoritmiche che privilegiano certi gruppi e marginalizzano altri. Un esperimento condotto da Global Witness nel 2023 su Facebook ha mostrato che il 91% dei destinatari di inserzioni per offerte di lavoro nel settore meccanico erano identificati come maschi, mentre il 79% dei destinatari di annunci per insegnanti di scuola materna erano identificati come femmine. Risultato diretto dell'ottimizzazione algoritmica della piattaforma, che impara dai comportamenti degli utenti e ne amplifica i pattern esistenti²¹². Per il designer che usa sistemi di IA per raccogliere dati di mercato, analizzare tendenze o identificare audience, questo tipo di bias può orientare inconsapevolmente le scelte progettuali verso rappresentazioni già distorte.

3. La terza dimensione riguarda il bias culturale e linguistico. I sistemi di IA generativa sono addestrati su corpus prevalentemente anglofoni e occidentali: la letteratura accademica, le immagini, i testi e le convenzioni estetiche di riferimento riflettono in modo sproporzionato una prospettiva nordamericana ed europea. Questo produce un effetto di invisibilità sistematica delle culture non egemoni: quando un designer usa questi strumenti per ricercare riferimenti visivi, tradizioni tipografiche, convenzioni estetiche locali o narrative culturali specifiche, il sistema tende a produrre output che normalizzano il canone occidentale come se fosse universale²¹³.
4. La quarta dimensione riguarda il bias nell'accessibilità degli strumenti stessi. I tool di IA generativa più avanzati, ovvero quelli con maggiore qualità degli output e minore tendenza ai bias, tendono a essere quelli più costosi o accessibili solo attraverso abbonamenti premium. Questo crea una asimmetria: i professionisti e le organizzazioni con meno risorse si trovano ad usare strumenti con bias più marcati, mentre chi ha più risorse accede a sistemi più raffinati. Per un designer freelance o per uno studio di piccole dimensioni, questa gerarchia di accesso può tradursi in una gerarchia di qualità etica degli output prodotti.
5. La quinta dimensione, spesso trascurata, riguarda il bias nel feedback loop. I sistemi di IA generativa imparano anche dai comportamenti degli utenti: i

prompt più frequenti, le immagini più selezionate, le risposte più apprezzate diventano segnali di training per le versioni successive del modello. Se i designer tendono a selezionare output che confermano le proprie aspettative estetiche, spesso già influenzate da bias culturali, contribuiscono involontariamente ad addestrare il sistema a produrre output ancora più conformi a quei pattern. Il bias, in questo modo, non è solo ereditato dai dati storici, bensì viene continuamente rinforzato dall'uso.

Riconoscere il bias: una competenza del designer

Come ha mostrato la ricerca di Brookings Institution (2024), i tentativi delle aziende di mitigare i bias nei propri modelli hanno prodotto risultati contraddittori: nel febbraio 2024, Google ha dovuto sospendere la funzione di generazione di immagini di Gemini dopo che i tentativi di correggere il bias di sovrarappresentazione dei bianchi avevano prodotto rappresentazioni storicamente anacronistiche, come soldati nazisti di varie etnie. Come ha osservato la ricercatrice di Stanford Pratyusha Kalluri, si tratta di un approccio whack-a-mole: si risponde ai bias più visibili senza affrontare le cause strutturali²¹⁴.

Il problema è reso più complesso dal fatto che non esiste una definizione universalmente condivisa di equità algoritmica. Diversi criteri quali parità demografica e individual fairness sono spesso matematicamente incompatibili tra loro, il che significa che qualsiasi sistema di IA incorpora già, nella sua architettura, una scelta di valore su quale forma di equità privilegiare²¹⁵.

A livello normativo, l'AI Act dell'UE richiede ai provider di sistemi ad alto rischio di utilizzare dataset di addestramento sufficientemente rappresentativi e di adottare misure per rilevare e mitigare i possibili bias. Ma la ricerca di Hagendorff (2020) ha già mostrato che l'etica dell'IA formulata come principi astratti non produce cambiamenti reali in assenza di meccanismi di enforcement vincolanti, rimanendo spesso una cooperazione volontaria tra eticisti e industria²¹⁶.

La risposta, per il designer, si conferma essere una vigilanza attiva e continua: verificare gli output per identificare pattern discriminatori, diversificare le fonti di riferimento, sviluppare una sensibilità culturale che permetta di riconoscere il bias anche quando è invisibile o normalizzato.

²¹⁴ Luhang Sun et al. *Smiling women pitching down: auditing representational and presentational gender biases in image-generative AI*, Journal of Computer-Mediated Communication, vol. 29, n. 1, 2024.

²¹⁵ MediaLaws, *Impact of Meta's Policy Changes on Gender-Based Discrimination in Its Advertising Model*, 2025.

²¹⁶ G. Cecere et al., *Technological Forecasting and Social Change*, Elsevier, 2024.

4.2. La figura del designer nell'era dell'AI

L'integrazione dell'intelligenza artificiale nei processi creativi impone una profonda riflessione sull'identità stessa del progettista. In un contesto in cui la macchina assume compiti sempre più complessi, è necessario ridefinire i confini dell'agency umana e il valore aggiunto della sensibilità professionale.

4.2.1 L'agency del designer

Il concetto di agency, inteso come la capacità di agire intenzionalmente e responsabilmente all'interno di un contesto, occupa una posizione centrale nella pratica del progetto. Il designer non può essere ridotto a un mero esecutore tecnico. Egli è, al contrario, un agente morale che attraverso le proprie scelte contribuisce a plasmare l'orizzonte materiale e culturale. Ogni decisione progettuale incorpora intrinsecamente valori, priorità e visioni del mondo, così come ogni artefatto comunicativo immesso nel circuito sociale produce effetti concreti su chi lo fruisce.

L'integrazione dell'intelligenza artificiale generativa nella progettazione solleva interrogativi fondamentali su quanta di questa autonomia decisionale rimanga effettivamente in capo al progettista nel momento in cui delega alla macchina compiti che un tempo richiedevano giudizio, competenza ed esperienza. In questo scenario, il ri-

schio non risiede soltanto nella potenziale perdita di efficienza o di qualità formale, bensì nella rottura del nesso tra intenzione e risultato, legame che definisce il design come pratica etica e responsabile. Un professionista che non sia più in grado di rendere conto delle ragioni profonde alla base delle proprie scelte non si scopre soltanto meno competente, ma perde la sua stessa libertà intellettuale²¹⁷.

4.2.2 Delega e automazione del processo progettuale

²¹⁷ Jeremy Baum, John Villaseñor, *Rendering Misrepresentation: Diversity Failures in AI Image Generation*, Brookings Institution, 2024.

La delega, in sé, non è un fenomeno nuovo né intrinsecamente problematico. In qualsiasi organizzazione un manager delega compiti che un tempo eseguiva direttamente, ma lo fa dopo anni di esperienza in cui ha svolto quei compiti in prima persona, costruendo un bagaglio di conoscenze che gli permette di comprendere, valutare e guidare il lavoro che ora delega. La differenza con la delega all'IA è strutturale: la delega tradizionale presuppone comprensione pregressa, mentre quella algoritmica rischia di avvenire prima che tale comprensione si formi. Per il designer, e in particolare per chi si trova ancora in fase di formazione, questo cambia radicalmente il significato dell'atto stesso di delegare.

La delega come problema cognitivo

L'uso dell'IA nei compiti cognitivi non è neutro. Questa pratica può migliorare l'efficienza ma anche ridurre memoria, attenzione e capacità di ragionamento autonomo. La questione etica fondamentale diventa quindi quanto controllo sulla propria mente siamo disposti a cedere. Il dottor Potasznik pone una domanda che ogni designer dovrebbe interiorizzare: "Cosa sono disposto a non essere in grado di fare?" La risposta avrà conseguenze dirette sulla responsabilità e sull'autorialità del progettista²¹⁸.

²¹⁸ Greg Demirchyan, *Algorithmic fairness: challenges to building an effective regulatory regime*, *Frontiers in Artificial Intelligence*, 2025.

Tuttavia, il problema non è solo individuale. Quando una generazione di designer impara a progettare affidandosi fin dall'inizio a strumenti che generano concept, selezionano riferimenti visivi e producono varianti senza che il professionista ne abbia mai attraversato il processo manualmente, si perde qualcosa di difficilmente recuperabile: la capacità di valutare criticamente quegli stessi output. Come avverte Floridi solo mantenendo il controllo del processo e non delegando la produzione di significato "quel logos che caratterizza l'intelligenza umana"

potrà mantenere il suo ruolo creativo e interpretativo²¹⁹. Se la delega inizia fin dal principio della formazione, come si può costruire una base di comprensione e intenzionalità su cui esercitare poi un giudizio critico?

²¹⁹ Thilo Hagendorff, *The Ethics of AI Ethics: An Evaluation of Guidelines, Minds and Machines*, Springer, 2020.

La delega come problema epistemico

In molti contesti decisionali, la fiducia funziona come un bene rivale: se si sceglie di seguire il consiglio dell'IA, spesso si scarta quello dell'esperto umano. Questo rifiuto sistematico della consulenza umana non è una scelta neutra: comunica un giudizio di insufficiente affidabilità, danneggiando il senso di dignità e competenza professionale della persona.

Si parla di danno epistemico: una forma di danno che diminuisce l'agency umana e non riconosce l'agency dei lavoratori come esperti qualificati. Questo danno non si riduce alla perdita del posto di lavoro. Anche quando l'IA non sostituisce direttamente i lavoratori, può limitare la loro capacità decisionale, trasformandoli in supervisori di sistemi automatici. Il problema quindi non è solo perdere il lavoro: è perdere il ruolo di soggetto competente nelle decisioni. Come scrivono i filosofi e ricercatori Saleh Afroogh e Jason D'Cruz, *"molto del senso e della dignità che i lavoratori sentono è il risultato dell'esercizio della propria agency"*. Se questa agency viene delegata agli algoritmi, i lavoratori possono sperimentare alienazione anche senza perdere il posto²¹⁷.

Questo si intreccia con il concetto di meaningful work: il lavoro significativo coinvolge una corrispondenza tra l'individuo e i compiti che svolge, consentendo l'espressione personale e la realizzazione dei propri valori e convinzioni. Se questa corrispondenza viene erosa dall'intermediazione algoritmica, non è solo il mercato a perdere un lavoratore qualificato, ma è l'individuo a perdere una fonte primaria di senso e identità.

Come alternativa al modello di IA che svolge lo stesso compito dell'umano, creando competizione e danno epistemico, gli stessi Afroogh e D'Cruz propongono il concept di modello della collaborazione avversaria: invece di offrire una raccomandazione alternativa, l'IA dovrebbe agire come un avvocato del diavolo, scrutinando le basi della decisione umana e identificando contro-prove o spiegazioni alternative. Questo protegge l'autorità dell'esperto umano, che rimane il decisore finale²¹⁷.

La delega come problema di bias e disuguaglianza

La delega all'IA può amplificare anche le disuguaglianze strutturali esistenti. Un sistema di intelligenza artificiale che tende a sostituire le raccomandazioni di professionisti appartenenti a gruppi già marginalizzati, perché poco rappresentati nei dati di addestramento o perché i loro stili estetici e le loro pratiche si discostano dai modelli statisticamente dominanti, finisce per rafforzare le disuguaglianze strutturali invece di contrastarle. Affidarsi ai dati che ci sono significa perpetuare la distribuzione del potere che quei dati riflettono²¹⁷.

Anche l'AI usa noi: il paradigma Human-Aided AI

Le grandi aziende di intelligenza artificiale, come Amazon, Google, Meta, Alibaba, Tencent e Baidu, costituiscono un centro di accumulazione e centralizzazione del potere economico e politico senza precedenti nell'attuale scenario capitalista post-industriale. In virtù della mole imponente di informazioni riguardanti la quasi totalità delle attività umane che riescono a incamerare, unita a una capitalizzazione di mercato straordinaria, a condizioni di quasi-monopolio e a un impatto politico globale immenso, tali realtà possono essere interpretate come veri e propri nuovi poli di potere sovrano, potenzialmente alla pari con la sovranità degli Stati. Questo immenso potere si alimenta però con una dinamica nascosta: mentre la narrativa dominante descrive gli esseri umani come utenti attivi di strumenti passivi, la realtà è che l'IA esiste solo perché l'umanità la alimenta costantemente. L'IA odierna è fondata sull'asservimento mondiale e massivo di utenti e lavoratori come produttori di dati e sottounità cognitive di sistemi computazionali. Non siamo solo noi a usare l'IA, è anche l'IA a usare noi: come serbatoi di dati, come validatori di output, come forza lavoro cognitiva non remunerata. Questo paradigma, definito Human-Aided AI, ha implicazioni dirette per il designer: ogni prompt che digita, ogni immagine che seleziona, ogni output che approva o rigetta contribuisce ad addestrare il sistema che usa. La delega, di fatto, non è mai unidirezionale²²⁰.

4.2.3 Il capitale semantico

²²⁰ Emmie Malone, et al. *When Trust Is Zero-Sum: Automation's Threat to Epistemic Agency*, ResearchGate, 2025.

Luciano Floridi ha elaborato il concetto di capitale semantico per indicare non la semplice somma delle informazioni in nostro possesso, ma la capacità generativa di creare, curare e trasmettere significato attraverso di esse. È grazie al capitale semantico che interpretiamo il

mondo, diamo vita a nuovi sensi a partire da quelli esistenti e tessiamo connessioni inedite tra elementi apparentemente distanti. Questo patrimonio è ciò che distingue una biblioteca da un archivio, una conversazione da uno scambio di dati, una cultura viva da un museo.

Tuttavia, l'adozione acritica dell'intelligenza artificiale rischia di accelerare le quattro minacce principali che Floridi individua come letali per questo patrimonio:

- la sua perdita, quando significati preziosi vengono distorti o cancellati insieme a storie e sfumature;
- la sua improduttività, quando patrimoni inaccessibili restano inerti;
- il suo uso improprio, quando simboli o eventi vengono caricati di significati che li tradiscono;
- il suo deprezzamento, che avviene quando il capitale semantico invecchia e si svuota in assenza di educazione e aggiornamento critico.

Come conseguenza di questa realtà, il vantaggio competitivo del designer non risiederà più nell'asimmetria informativa, ovvero nel possedere dati che altri non hanno, bensì nell'asimmetria semantica: la capacità di usare le informazioni accessibili a tutti per costruire significato in modo più consapevole, originale e responsabile. Se l'IA aumenta esponenzialmente la capacità di produrre varianti, il valore del progettista si sposta dalla quantità di idee generate alla capacità di interpretarle, selezionarle e dare loro una direzione. Questa competenza curatoriale è precisamente il terreno che la delega tecnologica incondizionata rischia di erodere. Un contenuto, infatti, non è mai un mero pacchetto di dati dotato di sintassi e struttura; esso diventa capitale semantico solo quando incontra una persona capace di percepirlo, interpretarlo e integrarlo nella propria storia.

Per chiarire il ruolo irriducibile dell'essere umano in questo processo, Floridi ricorre a due metafore efficaci. La prima è quella biologica: l'essere umano agisce come il granello di sabbia che, irritando l'ostrica con il suo carico di "fastidio" fatto di dolore, limite, fatica ed esperienza vissuta, la costringe a secernere la materia che trasformerà l'informazione in una perla di significato.

"Le macchine non possono avere capitale semantico. Possono rielaborare contenuti, ma non vivono esperienze, non

soffrono e sicuramente non si riconoscono in una narrazione. È l'essere umano, con il suo "granello di sabbia" – dolore, limite, fatica – a trasformare l'informazione in una perla di significato." – Luciano Floridi

La seconda metafora attiene alla paternità dell'opera: un film è sempre «regia di». Se la fotografia scattata con uno smartphone appartiene al dispositivo in senso puramente tecnico, essa è di chi ha scelto la luce, l'inquadratura e l'attimo. Allo stesso modo, un abito non è di Armani perché lo stilista lo ha cucito materialmente, ma perché incorpora le sue scelte, la sua visione e il suo sistema di valori estetici. Lo stesso principio si applica al design nell'era dell'IA: l'output appartiene a chi ha esercitato un giudizio critico su di esso, non allo strumento che lo ha materialmente generato. La scelta, e quindi il motivo per cui si predilige una specifica soluzione visiva o testuale rispetto a un'alternativa, resta un atto profondamente umano.

In una società post-AI, la vera urgenza si sposta dunque sul piano della responsabilità semantica. La questione non è più se usare l'intelligenza artificiale, né difendersi da essa in modo nostalgico, ma decidere che cosa vogliamo far emergere attraverso il suo impiego. Il capitale semantico va attivamente curato e trasmesso attraverso una sfida che non è solo tecnica, ma profondamente culturale. Solo presidiando questo confine potremo usare l'intelligenza artificiale senza delegarle il compito antropologico di generare senso, continuando a trasformare i nostri granelli di sabbia in perle di significato da condividere e arricchire²²¹.

4.2.4 Alienazione e sostituzione

²²¹ Xia Cheu, *How to Outsource Your Frontal Lobes*, UMass Media, 2024.

"Ci saranno vincitori e vinti, non solo tra i paesi, ma anche all'interno degli stessi. Le persone con un livello di istruzione più elevato probabilmente risentiranno meno negativamente degli effetti negativi [dall'intelligenza artificiale] o ne trarranno maggiori vantaggi, perché saranno meglio preparate a cogliere le opportunità, come quelle offerte dalle nuove tecnologie."

– Michal Rutkowski, Regional Director for Human Development, World Bank. (Interview with DEC President, Daniel A. Bielik, April 2025)

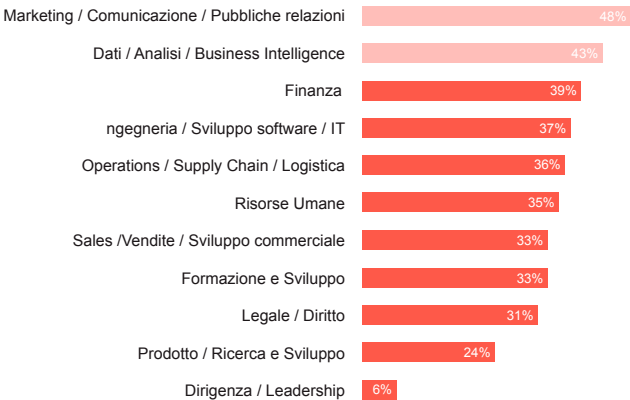
Secondo il rapporto "AI in the Workplace 2025" del Digital Education Council, il 72% dei datori di lavoro ritiene che l'adozione dell'intelligenza artificiale ridurrà il numero di

²²² Floridi Luciano, *Capitale semantico: la vera posta in gioco della rivoluzione digitale*, Wired Italia, 2025.

dipendenti necessari²²².

La ricerca di Goldman Sachs del 2023 ha stimato che cir-

Opinione dei datori di lavoro sull'impatto dell'adozione dell'AI sul numero di dipendenti per ruolo
In quali ruoli l'adozione dell'AI porterebbe a una riduzione del personale?



²²³ Rainer Mùhlhoff, *Subjectivity and Power in the Ethics of AI, The Ethics of AI: Power, Critique, Responsibility*, Bristol University Press, JSTOR Open Access Monograph, 2025.

ca il 26% dei compiti lavorativi nel settore delle arti, del design, dell'entertainment e dei media potrebbe essere automatizzato dall'IA generativa. L'UNCTAD ha documentato ondate ripetute di licenziamenti nell'industria creativa nel 2023 e nel 2024, molte delle quali esplicitamente collegate all'adozione di strumenti di IA. Il World Economic Forum, nel suo Future of Jobs Report del 2023, ha inserito "Graphic Designer" tra i ruoli a più rapido declino nel quinquennio successivo²²³.

Tuttavia, questi dati aggregati richiedono una lettura differenziata per tipo di ruolo che il designer assume. Ruoli come graphic designer, social media designer e brand designer, potrebbero rientrare nella categoria Communication [vedi figura x] e sono probabilmente i più esposti all'automazione nel breve periodo, perché molta produzione visiva standardizzata è già accelerata da strumenti generativi. I designer collocabili nella categoria Product/R&D [vedi figura x] tra cui UX/UI designer, interaction designer e product designer, sono meno colpiti nella componente strategica del loro lavoro: comprendere utenti, progettare flussi, prendere decisioni di prodotto, coordinarsi con business e sviluppo sono attività che richiedono giudizio situazionale e relazionale difficilmente automatizzabile. Tuttavia, anche all'interno del product design le mansioni più operative, ad esempio la modellazione CAD di base e il rendering 3D, sono già fortemente accelerate dall'IA e quindi più esposte. La sostituzione, in molti casi, non sarà totale ma parziale: non l'eliminazione

Figura 4.6: Impatto dell'adozione dell'intelligenza artificiale sulla riduzione del personale suddiviso per ruoli aziendali, secondo la prospettiva dei datori di lavoro. Grafico rielaborato da "AI and the Future of Work: Who Benefits and Who Falls Behind?", Digital Education Council.

del ruolo, ma la riduzione del numero di persone necessarie per svolgere determinate attività.

Il rapporto del Digital Education Council mostra anche un altro lato: il 62% dei datori di lavoro prevede la creazione di nuovi ruoli grazie all'intelligenza artificiale. Il 98% prevede un aumento dell'utilizzo dell'IA nei propri team. Ma il 18% non prevede la nascita di nuove figure professionali, e un ulteriore 20% rimane incerto. Come osserva il rapporto, il futuro del lavoro dipenderà meno dalla tecnologia in sé e più da quanto efficacemente individui e istituzioni prepareranno le persone a lavorare al suo fianco³³.

Opinione dei datori di lavoro sulla nascita di nuovi ruoli come risultato di una maggiore adozione dell'AI

Si aspetta che nascano nuovi ruoli nel suo settore a seguito di una maggiore adozione dell'AI?

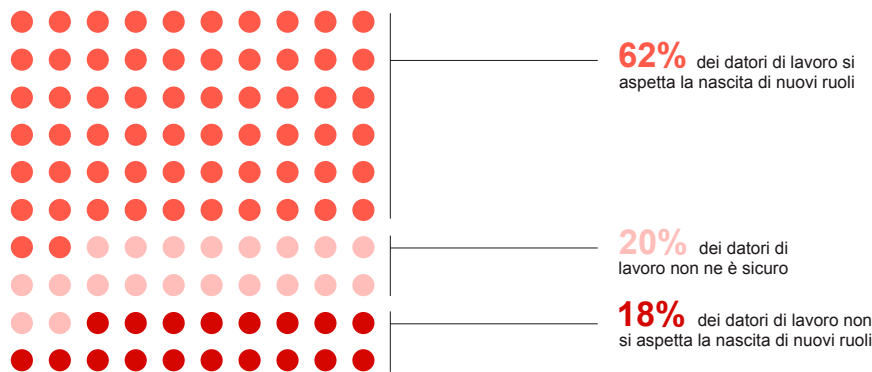


Figura 4.7: Prospettive dei datori di lavoro sulla nascita di nuove figure professionali come conseguenza dell'adozione dell'intelligenza artificiale nei contesti aziendali. Grafico rielaborato da "AI and the Future of Work: Who Benefits and Who Falls Behind?", Digital Education Council.

Oltre agli aspetti legali ed economici, è importante considerare l'impatto della creatività AI-driven sul senso di significato dell'artista nel suo lavoro. La dimensione dell'alienazione è forse la meno visibile ma non la meno importante. Come anticipato [vedi par. La delega come problema epistemico], quando gli artisti sperimentano una perdita estetica a causa del coinvolgimento dell'IA nel processo creativo, ciò che possono subire è un senso ridotto di significato e realizzazione nel lavoro, in linea con il concetto marxiano di lavoro alienato, in cui i lavoratori si sentono disconnessi dai prodotti del proprio lavoro. Se il designer non riconosce più il proprio giudizio, la propria sensibilità e la propria storia nel risultato finale del suo lavoro diminuisce il senso stesso dell'attività professionale²²⁴.

Abraham Moles aveva già intuito questa trasformazione nel suo *Art et ordinateur* (1990): "l'artista non sarà sostituito dalle macchine [...] perché l'artista è da sempre una

²²⁴ Floridi Luciano, *Dalla sabbia alla perla: il capitale semantico nell'era dell'intelligenza artificiale*, AI News, 2025.

Figura 4.8: Il sociologo e ingegnere francese Abraham Moles, pioniere nello studio della teoria dell'informazione e della cibernetica applicate alla percezione estetica e alla comunicazione. Immagine da "File:Abraham Moles", Wikimedia Commons.



macchina che crea forme visive, sonore, significanti [...]. Si può dire che la sua azione non sarà sostituita ma spostata nella sua funzione". La profezia di Moles si traduce oggi nell'invito al designer ad assumere sempre più il ruolo di direttore creativo di processi che l'IA esegue: non il produttore dell'output, ma il responsabile del senso che quell'output deve avere²²⁵.

4.2.5 Responsibility gap e negligenza professionale

Quando decisioni chiave vengono influenzate o determinate dall'output di sistemi algoritmici, il nesso tra azione progettuale e conseguenze materiali tende a farsi meno esplicito. Questo fenomeno è descritto in letteratura come responsibility gap: la difficoltà di attribuire responsabilità morale e operativa in presenza di sistemi autonomi o semi-autonomi, introdotta nel dibattito filosofico da Andreas Matthias (2004).

La delega algoritmica rischia di trasformarsi in una forma di negligenza progettuale non intenzionale. Non si tratta di una rinuncia esplicita alla responsabilità, bensì di una progressiva esternalizzazione del giudizio, in cui il progettista mantiene formalmente il ruolo decisionale ma ne riduce sostanzialmente l'esercizio critico. La negligenza emerge quando l'output dell'IA viene accettato come neutrale o oggettivo, senza una verifica delle ipotesi, dei vincoli e dei valori incorporati nel modello²²⁶.

²²⁵ Digital Education Council, *AI and the Future of Work: Who Benefits and Who Falls Behind*, 2024.

²²⁶ World Economic Forum, *Generative AI and Creative Jobs*, 2023

Quattro tipologie di lacuna nella responsabilità

Santoni de Sio e Mecacci (2021), in un influente articolo pubblicato su *Philosophy & Technology*, hanno elaborato una tassonomia più articolata del responsibility gap, distinguendo quattro tipologie di lacune²²⁷.

²²⁷ Xu Cheng, Yanqi Sun, Haibo Zhou, *Artificial Aesthetics and Ethical Ambiguity: Exploring Business Ethics in the Context of AI-driven Creativity*, *Journal of Business Ethics*, Springer, 2024.

1. Il gap nella culpability (imputabilità): quando un sistema di IA produce un danno, può essere impossibile identificare un individuo che abbia agito in modo moralmente biasimevole. Nessuno ha inteso il danno, nessuno lo ha previsto, eppure il danno c'è.
2. Il gap nella moral accountability (rendicontabilità morale pubblica): anche quando non c'è colpa, c'è un obbligo di spiegare e giustificare pubblicamente le proprie azioni. I sistemi opachi rendono questa rendicontabilità impossibile perché nessuno, nemmeno il produttore del sistema, è in grado di spiegare con precisione perché è stato prodotto un certo output.
3. Il gap nell'active responsibility (dovere di prevenzione): chi ha il dovere di anticipare i possibili danni e prevenirli? Quando la responsabilità è distribuita tra sviluppatori, deployer e utenti, questo dovere si disperde.
4. Il gap nella capability responsibility (competenza ad agire responsabilmente): anche chi vuole assumersi la responsabilità di un sistema di IA potrebbe non avere le competenze tecniche per farlo in modo efficace.

Queste quattro lacune si generano dalla combinazione di opacità, complessità e autonomia dei sistemi di IA. Il problema è aggravato dal fenomeno noto come many hands problem: nei sistemi sociotecnici complessi, la responsabilità si distribuisce tra così tanti attori (sviluppatori, addestratori, deployer, utenti) che nessuno sembra pienamente responsabile²²⁸.

4.2.6 Meaningful Human Control

²²⁸ Domus Web, *Intelligenza artificiale: libri fondamentali per capirla*, 2025.

La risposta più articolata al problema del responsibility gap è emersa attorno al concetto di meaningful human control (MHC), sviluppato originariamente nel contesto dei sistemi d'arma autonomi e successivamente esteso a diversi domini. L'MHC definisce le condizioni normative in presenza delle quali il controllo umano su un sistema di IA è sufficientemente robusto da permettere l'attribuzio-

ne di responsabilità. Non è sufficiente che un essere umano sia formalmente nel loop decisionale: occorre che abbia comprensione adeguata del sistema, capacità reale di intervenire e di contestarne le decisioni, e un contesto organizzativo che non svuoti il suo controllo di significato²²⁷.

Nel contesto del design, il meaningful human control implica che il designer non solo utilizzi gli strumenti di IA, ma comprenda sufficientemente i loro meccanismi per poter valutare criticamente gli output, identificare i casi in cui il sistema ha probabilmente sbagliato, e assumere piena responsabilità delle scelte progettuali finali.

Responsabilità culturale

“Coltivare, educare, umanizzare: tre verbi densi, che ci invitano a una responsabilità condivisa. Coltivare, inteso come prendersi cura del pensiero, dei linguaggi e dei legami. Educare, come gesto di apertura all'altro e di trasmissione consapevole. Umanizzare, come sfida per rendere ogni tecnologia una scelta di civiltà. La cultura italiana ha dimostrato, nei secoli, di saper coniugare eredità e invenzione. A noi, oggi, spetta raccogliere quel testimone e proiettarlo nel tempo dell'intelligenza artificiale.” - Gaetano di Tondo

L'Intelligenza Artificiale è spesso presentata quale forza autonoma, capace di apprendere, prevedere e sostituire. Tuttavia, il dibattito attuale spesso trascura la natura profondamente culturale del suo impatto. Ogni algoritmo, modello linguistico o rete neurale reca in sé, oltre all'impronta tecnica, una visione implicita del mondo. Adottare gli strumenti generativi senza interrogarli significa accettarne acriticamente i presupposti epistemologici e trasmetterli, attraverso i propri artefatti, al pubblico. In questo senso, l'IA non può essere governata solo con etichette di sicurezza o codici etici esterni: richiede una cultura umanistica viva, capace di generare uno sguardo critico sui sistemi automatizzati e di ricollocare al centro il soggetto umano nella sua complessità. Educare al design nell'era dell'IA significa innanzitutto educare allo sguardo critico, alla responsabilità e al dubbio: tre qualità che nessun sistema algoritmico può sostituire e che costituiscono, oggi più che mai, il nucleo irriducibile della competenza progettuale²²⁹.

²²⁹ Cervantes, Juan-Antonio et al., *Ascribing Responsibility for the Actions of Learning Automata, Philosophy & Technology*, Springer, 2021.

Vera Tripodi

Filosofa morale,
professoressa universitaria

La presenza di Vera Tripodi garantisce un necessario approfondimento sull'etica della tecnologia, distinguendo tra ciò che è tecnicamente possibile e ciò che è eticamente giustificabile. Dopo una formazione in filosofia e anni di ricerca accademica tra l'Università di Torino e collaborazioni internazionali, oggi insegna Etica delle Tecnologie Emergenti e Filosofia dell'Innovazione, concentrandosi proprio sull'impatto sociale dei sistemi automatizzati. Il suo contributo si focalizza sull'individuazione dei rischi legati ai bias e alle disuguaglianze sociali amplificate dall'IA.

Attraverso la sua attività di ricerca e insegnamento, approfondisce come muti la responsabilità professionale di progettisti e ingegneri nel momento in cui l'automazione incide profondamente sulle scelte di design, offrendo gli strumenti critici per orientare lo sviluppo tecnologico verso una direzione più equa e consapevole.



Nei suoi insegnamenti affronta il tema del rapporto tra ciò che è tecnicamente fattibile e ciò che è eticamente giustificabile: nel caso dell'intelligenza artificiale, come possiamo evitare che la possibilità tecnica guidi automaticamente le scelte progettuali?

Io credo che il punto centrale sia evitare una sorta di automatismo per cui, se qualcosa si può fare tecnicamente, allora si fa. Questo rischio esiste molto nel campo dell'intelligenza artificiale, dove l'innovazione corre più veloce della riflessione sulle sue conseguenze. Per evitare questo, penso sia fondamentale portare l'etica dentro il processo di progettazione fin dall'inizio, e non considerarla solo alla fine, come una verifica successiva. In altre parole, non si tratta di chiedersi "è giusto o sbagliato?" dopo aver costruito il sistema di IA, ma di far entrare questa domanda direttamente nel modo in cui il sistema viene pensato. Questo richiede anche un lavoro interdisciplinare, perché chi progetta tecnologia non può rispondere da solo a queste domande. E poi ci sono questioni molto concrete: quali valori stiamo incorporando in un sistema? Chi li decide? E in che modo un artefatto tecnologico finisce per riflettere quei valori? Alla fine, credo che la responsabilità di chi progetta sia centrale: ogni scelta tecnica, anche la più piccola, non è mai neutrale, ma porta con sé conseguenze sul piano sociale.

Considerando che l'AI può riprodurre o amplificare disuguaglianze e bias sociali già presenti, quali ritiene siano i principali rischi associati a questo fenomeno?

Il rischio principale, a mio avviso, è che l'intelligenza artificiale finisca per amplificare disuguaglianze che già esistono nella società, spesso in modo poco visibile.

Questo succede perché i sistemi impara-

no dai dati del passato, e quindi tendono a "ereditarne" anche i problemi, compresi i pregiudizi. Il punto è che questi bias non sono sempre intenzionali, ma possono comunque produrre effetti molto concreti. Un esempio è quello della sanità: se i dati non rappresentano bene tutta la popolazione, alcuni gruppi possono ricevere diagnosi meno accurate o trattamenti meno efficaci. Un altro ambito delicato è quello delle risorse umane, dove sistemi automatizzati possono finire per riprodurre stereotipi legati al genere o ad altre caratteristiche sociali. C'è poi un aspetto che considero particolarmente importante, cioè la mancanza di trasparenza: spesso non è facile capire come una decisione sia stata presa, e questo rende difficile anche correggere eventuali errori. Infine, c'è un rischio più sottile, ma secondo me molto rilevante: quello di considerare "neutrali" decisioni che in realtà sono il risultato di scelte umane precise, anche se nascoste dentro il sistema. che sono invece il risultato di scelte progettuali e di assunzioni valoriali implicite.

Un'altra tematica centrale del dibattito sull'AI è quella delle responsabilità professionali: come cambia oggi la responsabilità di progettisti e ingegneri quando lavorano con sistemi di intelligenza artificiale?

Secondo me, cambia in modo abbastanza profondo. Non basta più pensare alla responsabilità come "il sistema funziona correttamente oppure no".

Oggi bisogna considerare che un sistema di intelligenza artificiale è sempre parte di un contesto più ampio: dipende dai dati, da come viene usato, e dal contesto sociale in cui viene inserito. Quindi la responsabilità diventa inevitabilmente più ampia, più "socio-tecnica". Inoltre, non si tratta più di una responsabilità limitata alla fase di progettazione. I sistemi continuano a evolvere nel tempo,

quindi anche chi li sviluppa deve pensare a tutto il loro ciclo di vita: manutenzione, aggiornamenti, e in alcuni casi anche dismissione. E poi c'è un aspetto che considero decisivo: la consapevolezza che le scelte tecniche non sono mai solo tecniche. Per questo oggi servono competenze più ampie, non solo ingegneristiche ma anche etiche e sociali, per capire davvero l'impatto di ciò che si costruisce.

4.3. L'irriducibilità dell'umano

Nonostante le straordinarie capacità di generazione e analisi delle macchine, rimangono dimensioni della progettazione che sembrano resistere all'automazione. Esplorare ciò che rende l'intuito e il giudizio umano irriducibili permette di individuare il nucleo autentico della professione nel nuovo scenario tecnologico.

4.3.1 L'infrastruttura del potere algoritmico e la crisi della fiducia

Con fiducia epistemica si intende la fiducia nelle fonti di conoscenza e nelle procedure attraverso cui il sapere viene prodotto, validato e diffuso e costituisce uno dei beni sociali più preziosi e, al tempo stesso, più fragili della contemporaneità. L'avvento dell'intelligenza artificiale generativa la sottopone oggi a una pressione multidirezionale senza precedenti. Da un lato, tali sistemi sono in grado di produrre contenuti falsi ma altamente plausibili, rendendo sempre più difficile distinguere tra informazione autentica e manipolazione. Dall'altro, l'AI generativa contribuisce a una crescente concentrazione del controllo delle infrastrutture informative nelle mani di pochi attori privati globali, ridimensionando progressivamente il ruolo tradizionale degli esperti e delle istituzioni come mediatori del sapere.

i meccanismi tradizionali di fact-checking faticano strutturalmente a tenere il passo con la velocità e la massa dei contenuti generati. La sfida contemporanea non si limita più alla necessità di smentire singole falsità, bensì risiede nella gestione di un flusso costante e pervasivo di narrazioni sintetiche che imitano perfettamente le strutture del ragionamento umano. Nel momento in cui l'IA si emancipa dal ruolo di mero assistente e inizia a generare testo e informazione in modo del tutto autonomo, la distinzione tra una fonte autentica e una fabbricata comincia inevitabilmente a sfumare. Richiamando le dinamiche di potere analizzate da Michel Foucault, la ricerca avverte del rischio sistemico legato alla nascita di un nuovo monopolio epistemico: se le piattaforme tecnologiche private diventano gli unici arbitri di ciò che è considerato affidabile e veritiero, si instaura una governance della verità privatizzata. Al contrario, se la regolazione si rivela troppo ampia o inefficace, il panorama rischia di essere dominato da una disinformazione endemica, capace di minare la fiducia pubblica nella conoscenza stessa²³⁰.

4.3.2 La maschera dell'etica tecnologica: il fenomeno dell'ethics-washing

²³⁰ Ellen Hohma et al., *Investigating accountability for Artificial Intelligence through risk governance: A workshop-based exploratory study*, PMC / NCBI, 2023.

A fronte di questa crisi epistemica, l'ultimo decennio ha registrato una proliferazione massiccia di linee guida, dichiarazioni di intenti e carte etiche prodotte da organismi internazionali. Il problema che tali cornici teoriche spesso presentano è il fatto di rimanere confinate a un piano astratto, venendo strumentalizzate per finalità di ethics-washing, termine con il quale si indica *"l'uso strategico del linguaggio etico per guadagni reputazionali piuttosto che per una governance sostanziale"*. La produzione di principi privi di reali meccanismi di applicazione e sanzione finisce così per generare una retorica dietro la quale le pratiche dannose delle grandi piattaforme continuano indisturbate²³⁰.

Per approfondire questa distorsione, è stata analizzata la distribuzione non uniforme dei valori all'interno degli strumenti etici attualmente disponibili sul mercato. Se da un lato concetti come la trasparenza, la giustizia, l'equità e la responsabilità formale risultano onnipresenti, dall'altro valori più profondi e strutturali quali la libertà, la solidarietà o la sostenibilità sistemica vengono raramente operazionalizzati o tradotti in vincoli tecnici. Il risultato è una preoccupante riduzione dell'etica a una

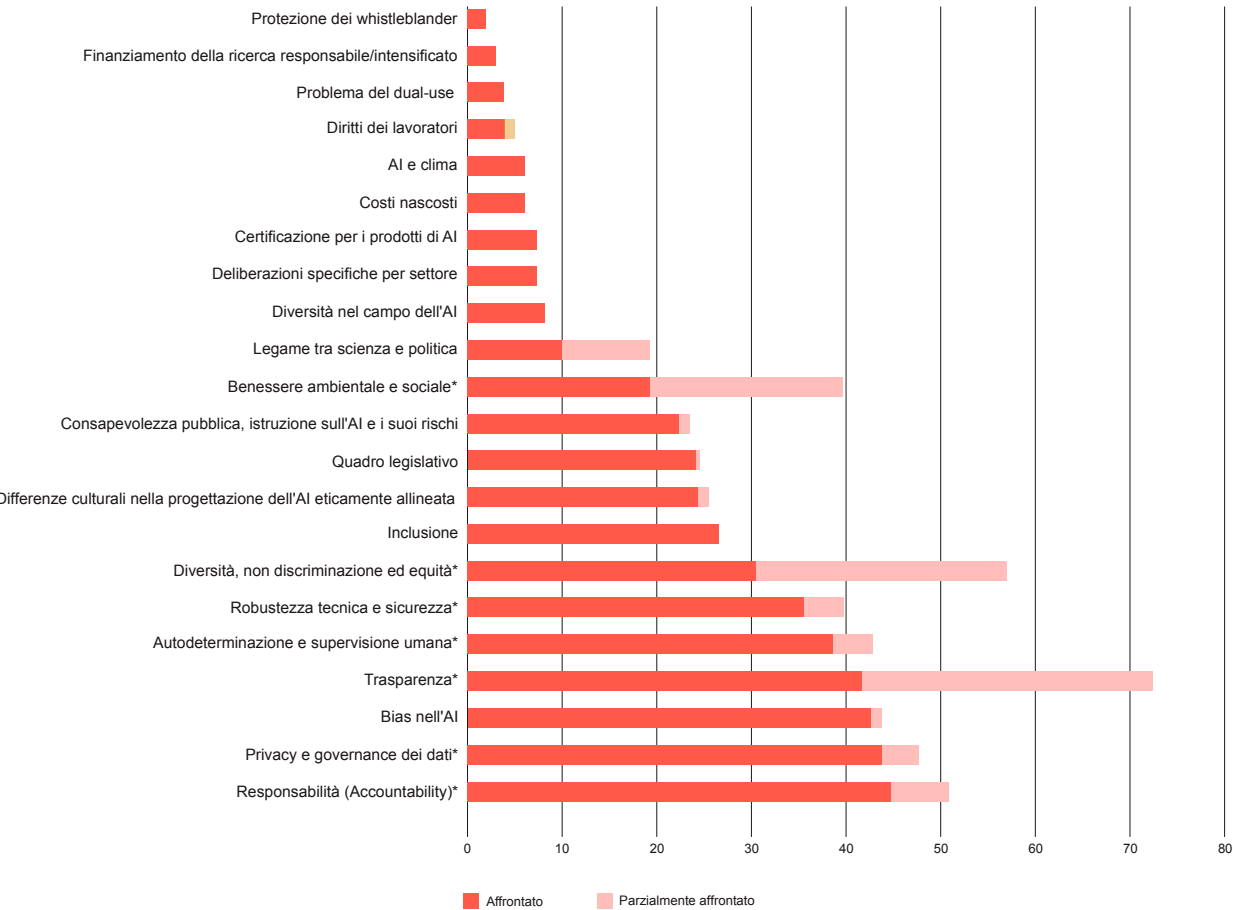
Figura 4.9: Frequenza dei principi etici integrati nei principali strumenti digitali di autovalutazione. Grafico rielaborato da “Decoding AI ethics: What do ethical tools tell us about ethics?”, Portegies T. et al.

²³¹ DiCulther, *Educare al futuro: l'incontro tra intelligenza artificiale e patrimonio culturale*, 2025.

mera pratica burocratica, una serie di caselle da spuntare (il cosiddetto box-ticking) che perde di vista la reale complessità del contesto sociale e politico.

Inoltre, gli autori evidenziano come la totalità degli strumenti esaminati si affidi esclusivamente a forme di autovalutazione per stimare la conformità ai principi specifici. Una simile evidenza suggerisce che l'etica, per come viene applicata in questi tool, sia auto-regolata dalle aziende stesse, sollevando seri dubbi sulla responsabilità oggettiva e sulla reale neutralità delle valutazioni prodotte²³¹.

Principi etici affrontati dagli strumenti etici:



Nota. La tabella include tutti i principi etici tratti dallo studio di Hagendorff (2020).

* I principi etici sono inclusi nelle Linee guida europee.

4.3.3 La funzione autore e la crisi dell'autorialità

Chiara De Vecchis e Paolo Traniello, nel loro saggio "La proprietà del pensiero. Il diritto d'autore dal Settecento a oggi" (2012), ricostruiscono storicamente come lo statuto dell'autore come proprietario legittimo dell'opera non sia un dato naturale, ma il frutto di una precisa evoluzione legislativa. Se in un primo tempo il riconoscimento giuridico era connesso principalmente all'investimento dell'imprenditore, è solo con la legge inglese Statute of Anne del 1710 che la figura dell'autore viene codificata nei termini moderni di proprietà intellettuale²³².

²³² Thierry Poibeau, *Disinformation, Misinformation and the Crisis of Trust in AI-Generated Content*, ResearchGate, 2025.

L'avvento dell'intelligenza artificiale generativa destabilizza radicalmente questa nozione rinascimentale di autore unico. Per comprendere la portata di tale crisi è utile recuperare la celebre concettualizzazione foucaultiana espressa nel saggio "Qu'est-ce qu'un auteur?" (1969). Foucault argomenta che il nome dell'autore non si limita a designare un individuo in carne e ossa, ma svolge una precisa funzione discorsiva: è uno strumento che permette di delimitare, classificare e differenziare i testi, stabilendo relazioni reciproche e controllando la circolazione stessa del significato all'interno della società. L'autore è, in sintesi, un principio di organizzazione del discorso e un meccanismo di controllo semantico.

L'applicazione di questa prospettiva ai modelli generativi produce un effetto profondamente destabilizzante. Se l'autore viene concepito non come origine sovrana del testo, ma come una funzione discorsiva che organizza, legittima e attribuisce valore ai contenuti, allora nulla impedisce, sul piano teorico, che tale funzione venga occupata anche da un sistema algoritmico. In questo modo, l'autorità autoriale tende a spostarsi dall'intenzionalità soggettiva alla capacità di produrre significato e valore economico.

La modernità occidentale ha tradizionalmente fondato l'idea di autorialità sul trittico intenzionalità, originalità e responsabilità. I modelli generativi contemporanei operano invece attraverso correlazioni matematiche e inferenze probabilistiche, prive di interiorità, coscienza o intenzione in senso umano. Di fronte a questa frattura, la giurisprudenza sul copyright tenta di preservare la centralità del soggetto umano trasferendo l'autorialità al controllore, al programmatore o al supervisore del processo algoritmico. Tale soluzione, tuttavia, non dissolve le tensioni ontologiche introdotte dall'AI generativa, ma ne ridefinisce semplicemente i confini, lasciando irrisol-



²³³ Tijs Portegies et al., *Decoding AI ethics: What do ethical tools tell us about ethics?*, Ethics and Information Technology, Springer, 2025.

Figura 4.10: A Recent Entrance to Paradise, l'opera generata autonomamente dal sistema algoritmico Creativity Machine. Immagine da "Un'opera creata dall'IA può essere protetta da diritto d'autore? La Corte di Washington verso la decisione", Lavagnini S.

ta la questione del rapporto tra produzione del testo, intenzionalità e responsabilità²³³.

Il dibattito giuridico internazionale ha già prodotto precedenti storici in materia. Il 18 marzo 2025, il tribunale distrettuale federale del District of Columbia ha confermato la decisione dell'Ufficio per il copyright degli Stati Uniti (USCO) nel celebre caso Thaler v. Perlmutter, stabilendo che un modello di IA non può godere della tutela del diritto d'autore per un'opera creata dal modello stesso. Stephen Thaler, sviluppatore del sistema di IA generativa denominato Creativity Machine, aveva tentato di registrare l'immagine sintetica intitolata A Recent Entrance to Paradise, indicando la macchina stessa come unico autore. L'USCO ha respinto la richiesta, e il tribunale ha successivamente confermato la decisione ribadendo un principio cardine del Copyright Act del 1976: un'opera, per essere oggetto di diritto e protezione, deve essere tassativamente il frutto della creatività e dell'ingegno umano²³⁴.

4.3.4 Estetica frammentata e l'illusione della creatività post-antropocentrica

²³⁴ Chiara De Vecchis, Paolo Traniello, *La proprietà del pensiero. Il diritto d'autore dal Settecento a oggi*, Carocci Editore, 2012.

Se il diritto tenta di tracciare un confine netto, la sociologia dell'arte propone definizioni di creatività che rompono da tempo con il mito romantico e individualista dell'artista isolato. La creatività viene infatti riletta come

un risultato socialmente costruito, derivante dall'interazione ambientale, dalla negoziazione di significati e dal giudizio estetico collettivo. Questi processi sono storicamente radicati in strutture di potere e conoscenza ben più ampie, dimostrando come la produzione creativa sia sempre stata un fenomeno collettivo e relazionale.

In questa cornice, alcuni studiosi propongono il concetto di creatività post-antropocentrica o post-creatività: una prospettiva teorica ampliata che trascende l'esclusività dell'azione umana per riconoscere sia i contributi collettivi della rete sociale (programmatori, annotatori di dati, utenti che istruiscono i modelli), sia il ruolo attivo di entità non umane, incluse la tecnologia coinvolta nel processo²³⁵.

²³⁵ Lala Hajibayova, *Authorship in the Age of Artificial Intelligence*, ResearchGate, 2025.

La prospettiva postumanista suggerisce così che sia la creazione artistica sia la correlata responsabilità etica debbano essere considerate come equamente ripartite e condivise tra agenti umani e non umani, richiedendo la formulazione di un discorso etico congiunto²²⁴.

“Se l'arte richiede espressione interiore e immaginazione, come nella visione espressionista, le macchine non possono davvero creare arte perché mancano di coscienza e di un sé autentico da esprimere; ma se è definita attraverso la mimesi o la rivelazione, possono essere considerate creatrici.” - Mark Coeckelbergh (2017)

Tuttavia, questa tendenza alla dissoluzione dell'umano trova un forte limite teorico nella distinzione fondamentale avanzata dal filosofo della tecnologia Mark Coeckelbergh (2017). Se l'arte viene definita attraverso la lente dell'espressionismo, intesa cioè come manifestazione di un'interiorità e di un vissuto, le macchine non possono in alcun modo creare arte, poiché strutturalmente prive di coscienza e di un sé autentico da esprimere. Se, al contrario, l'arte viene ridotta a mera mimesi, esecuzione formale o rivelazione combinatoria, allora i sistemi computazionali possono essere considerati creatori. Come argomentano diversi studiosi, i risultati prodotti dall'IA rimangono pur sempre riproduzioni tecniche basate su dataset preesistenti. L'antropomorfizzazione acritica di questi strumenti potrebbe sminuire gli artisti umani, svalutando il loro lavoro²³⁵.

La differenza tra macchina e mente umana non è accessoria: è ontologica. I sistemi di IA generativa producono risultati, non risposte. La questione se l'IA possa considerarsi capace di creare arte è emblematica. Se un arti-

sta utilizza l'IA per produrre suoni e li definisce come sua creazione, l'IA non è l'autore di quell'oggetto artistico, così come il fabbricante dei pennelli di Picasso o dell'orinatoio di Duchamp non era l'autore delle loro opere. Il fatto che l'IA possa generare un tessuto musicale non deve trarre in inganno: essa "sta forgiando senza coscienza del proprio agire, senza volontà, senza intenzione, senza aver maggior coscienza, volontà o intenzione abbia una pressa nel dar forma a una lamiera".

Come scrive il giurista Valter Fraccaro, *"La ricerca mentale, inclusa quella artistica, è volta a ottenere risposte: solo lo spirito critico della persona è in grado di stabilire se uno dei possibili risultati offerti dalla macchina sia una risposta accettabile"*²³⁵.

Margaret Boden, psicologa cognitiva britannica e una delle massime autorità sullo studio della creatività computazionale, sostiene che i modelli computazionali offrano preziose informazioni sui processi creativi umani e sulla natura strutturata del pensiero creativo. Tuttavia, avverte che l'IA "non può replicare la complessità e la profondità della mente umana". Sebbene l'IA possa generare idee che appaiono creative (ossia nuove, sorprendenti e preziose), la nostra incapacità di comprendere appieno il funzionamento della mente umana, un mistero che persiste "fin dagli antichi medici e filosofi greci", implica che nemmeno l'IA possa farlo²³⁶.

²³⁶ Wired Italia, *Intelligenza artificiale e diritto d'autore: i limiti*, 2025.



Figura 4.11: La scienziata cognitiva e filosofa Margaret Boden, figura fondamentale nella storia dell'intelligenza artificiale per i suoi studi pionieristici sulla creatività computazionale e sulla filosofia della mente. Immagine da "Margaret Boden (1936–2024)", Ferry G.

Questa tensione teorica si traduce in una destabilizzazione delle categorie dell'estetica tradizionale:

- Il controllo estetico: l'autorità del progettista sullo stile, la composizione e le scelte cromatiche vacilla nel momento in cui tali aspetti vengono co-determinati da un algoritmo che opera su pattern statistici estratti da milioni di opere altrui.
- Il controllo creativo: l'intrusione dell'IA lungo la filiera complica la tracciabilità della paternità, rendendo ambiguo il confine in cui termina il contributo dello strumento e inizia quello dell'autore.
- L'autenticità: storicamente legata all'espressionismo dell'artefice, entra in crisi di fronte a un output derivante da correlazioni probabilistiche di dati terzi.
- L'identità artistica: la fusione simbiotica rischia di diluire l'identità stilistica individuale, riducendo il designer a una mera funzione di supervisione del sistema.
- L'originalità: risulta compromessa, poiché ogni output generato rappresenta una ricombinazione statistica di materiale preesistente.

Nella ricerca qualitativa pubblicata nel 2025 nel *Journal of Business Ethics* condotta nell'ambito del design assistito dall'IA, gli artisti e i progettisti coinvolti hanno manifestato profonde difficoltà nel distinguere il proprio apporto creativo da quello fornito dall'IA, giungendo a mettere in discussione l'autenticità stessa del proprio lavoro. L'indagine dimostra chiaramente come queste criticità non possano essere risolte unicamente mediante risposte estetiche o tutele legali, rendendo la dimensione etica un presidio indispensabile per garantire pratiche eque e trasparenti nel design assistito dall'algoritmo²²⁶.

4.3.5 Il limite biologico e l'immaginazione morale

Storicamente, la tecnologia si è configurata come uno strumento al servizio dell'azione umana. Su questo specifico binomio tra l'essere umano e lo strumento si innescava la riflessione del filosofo Maurizio Ferraris, il quale ribalta le tesi tecnofobiche tradizionali per ridefinire l'uomo come un animale politecnico. Secondo Ferraris, la tecnica non è un elemento alieno che minaccia di deumanizzarci, bensì la base stessa del processo di uma-

²³⁷ VareseNews, *L'intelligenza artificiale secondo Maurizio Ferraris a Gallarate*, 2025.

²³⁸ Ferraris, Maurizio, *Non c'è niente di più umano dell'intelligenza artificiale*, Il Foglio, 2025.

²³⁹ Maria Danielsen, Sabine Roeser, *Why emotions and works of art are pertinent to the assessment of the ethical risks of AI*, Ethics and Information Technology, Springer, 2025.

nizzazione: l'animale umano nasce biologicamente debole, privo di un ambiente definito e strutturalmente disadattato, e ha trovato nella tecnica, dal primitivo bastone da scavo fino ai modelli di intelligenza artificiale, l'unica forma possibile di adattamento e di emancipazione da una originaria condizione di preda. In questa prospettiva, la delega tecnologica non è una perdita di sé, ma un potenziamento: l'uomo delega capacità alla tecnica e quest'ultima gliela restituisce amplificate, rendendolo, nel corso di una storia lunghissima, sempre più umano²³⁷.

Tuttavia, Ferraris introduce una linea di demarcazione netta che investe direttamente il concetto di intenzionalità: ai computer, per quanto evoluti, mancano radicalmente la vita, la volontà e la facoltà dei fini. I sistemi algoritmici non prendono alcuna iniziativa e sono privi di un mondo dello spirito o di un inconscio. Necessitano pertanto, in ogni istante, di un prompt o di un bisogno umano che li attivi. La vera discontinuità dell'intelligenza artificiale rispetto al bastone di legno primordiale risiede però nella natura della nostra interazione con essa. Ogni domanda che rivolgiamo alla macchina o ogni compito che le assegniamo non è un atto neutro, ma una produzione istantanea di valore che genera una duplice capitalizzazione: da un lato incrementa il capitale di dati complessivo, dall'altro fa fruttare l'immenso patrimonio del sapere accumulato dall'umanità che, senza la sollecitazione e l'interrogazione umana, rimarrebbe un archivio inerte e del tutto infruttuoso²³⁸.

Quando ci interroghiamo sulla possibile sostituzione dell'essere umano, è necessario non perdere di vista il fatto che i sistemi di IA non possono genuinamente ragionare secondo i valori fondamentali relazionali e morali della nostra specie. Le emozioni umane e l'empatia non sono interferenze cognitive, ma requisiti essenziali per comprendere l'impatto morale delle scelte progettuali.

Le decisioni del design richiedono quella che la ricerca chiama immaginazione morale: una capacità radicata nell'esperienza vissuta, nell'esposizione estetica, nella disponibilità a confrontarsi con la propria vulnerabilità. È precisamente questa capacità che nessun modello probabilistico potrà automatizzare, e che costituisce oggi il nucleo irriducibile della competenza progettuale²³⁹.

Maurizio Ferraris

Filosofo teoretico,
professore emerito

Filosofo di fama internazionale e Professore Emerito di Filosofia teoretica presso l'Università di Torino, Maurizio Ferraris è una delle figure più autorevoli del pensiero contemporaneo, noto per aver dato vita alla corrente del "Nuovo Realismo". La sua vasta produzione saggistica e la direzione del centro di ricerca LabOnt (Laboratorio di Ontologia) lo pongono come un teorico fondamentale per analizzare l'impatto della tecnica sull'essere umano e sulla società.

La sua riflessione si muove sul confine tra potenziamento e sostituzione, evidenziando come la volontà e l'iniziativa restino prerogative umane fondamentali. È stato scelto per la sua capacità di leggere criticamente i fenomeni tecnologici attuali e le loro conseguenze sul senso del lavoro e sulla dignità della persona.



La distinzione tra intelligenza naturale e artificiale, in base ad una sua affermazione, è principalmente una questione di "iniziativa" e "volontà", che alla seconda mancano. Tuttavia, un tema di forte rilevanza nel design è che per l'utente finale questa differenza rischia di diventare invisibile. Se ciò che distingue l'umano è la volontà e non l'output, come possiamo rendere rilevante questa dimensione affinché il lavoro umano non perda senso e dignità?

La differenza tra intelligenza naturale e artificiale non riguarda soltanto la capacità di produrre risultati, ma la struttura intenzionale che orienta l'azione. Una macchina può generare testi, immagini, simulazioni linguistiche o decisioni operative estremamente sofisticate, ma non desidera nulla, non teme nulla, non spera nulla. L'essere umano invece agisce dentro un orizzonte di fini, di aspettative, di conflitti e di interpretazioni. Anche quando svolge un compito apparentemente semplice, l'essere umano lo colloca dentro una storia personale e collettiva. Questa è la ragione per cui la volontà e l'iniziativa costituiscono il punto decisivo della distinzione.

Il rischio contemporaneo consiste nel fatto che, dal punto di vista dell'utente, la qualità dell'output tende a oscurare la differenza ontologica tra chi opera e ciò che opera. Se un sistema artificiale produce un risultato efficace, rapido e persino creativo, l'origine umana dell'iniziativa può diventare invisibile. Nel design questo problema è particolarmente evidente, perché l'utente entra in rapporto con interfacce sempre più fluide e persuasive.

Quando la relazione con la macchina assume una forma apparentemente dialogica, l'impressione è che ci sia una soggettività operante dietro il dispositivo. In realtà, ciò che appare come intenzionalità è il risultato di correlazioni sta-

tistiche e di processi computazionali.

Per evitare questa invisibilità non basta rivendicare astrattamente la centralità dell'umano. Occorre mostrare che il valore del lavoro umano non coincide con la sola esecuzione tecnica di un compito. La storia della tecnica dimostra che gli strumenti non sostituiscono semplicemente capacità precedenti, ma trasformano il significato delle attività. La scrittura non ha eliminato la memoria, la fotografia non ha eliminato la pittura, il computer non ha eliminato il pensiero. Ogni trasformazione tecnica ridistribuisce funzioni e attribuisce nuovo rilievo a dimensioni che prima restavano sullo sfondo.

Nel caso dell'intelligenza artificiale, ciò che rischia di passare in secondo piano è precisamente la dimensione della decisione e dell'assunzione di fini. Un algoritmo può ottimizzare un processo, ma non può stabilire perché un certo fine debba essere perseguito invece di un altro. Può suggerire la soluzione statisticamente più efficace, ma non può comprendere il valore umano, simbolico o storico di quella soluzione. Un esempio è offerto dalla medicina: un sistema artificiale può analizzare milioni di dati clinici e formulare diagnosi con altissima precisione, ma la decisione relativa a ciò che conta come cura, sofferenza o qualità della vita resta inscritta in un orizzonte umano. Allo stesso modo, nell'educazione un sistema artificiale può organizzare contenuti e valutare prestazioni, ma non può comprendere il significato esistenziale della formazione.

Il design allora assume una funzione decisiva perché non si limita a costruire strumenti funzionali, ma organizza l'esperienza degli utenti. Rendere visibile il contributo umano significa progettare sistemi che non occultino la genealogia del sapere che li rende possibili.

Ogni intelligenza artificiale è il risultato di enormi quantità di dati, di lavoro cognitivo accumulato, di scelte culturali e di infrastrutture sociali. Se tutto questo viene nascosto dietro l'illusione di una macchina autonoma, il lavoro umano perde riconoscibilità simbolica prima ancora che economica.

C'è inoltre un problema culturale più profondo. La modernità industriale aveva attribuito dignità al lavoro soprattutto attraverso la produttività. Oggi questo criterio rischia di diventare insufficiente, perché le macchine possono raggiungere livelli di efficienza superiori in numerosi ambiti. Occorre allora ridefinire il valore del lavoro umano non soltanto come prestazione tecnica, ma come capacità di iniziativa, interpretazione, cooperazione e costruzione di significati condivisi. In questo senso, il lavoro umano conserva una specificità perché implica sempre una relazione con il mondo comune e con altri esseri umani.

La dignità del lavoro non dipende dunque soltanto dalla produttività, ma dal fatto che nel lavoro si esprime una volontà capace di dare forma al mondo comune. Se dimentichiamo questo elemento, rischiamo di giudicare l'essere umano soltanto in base alla velocità e all'efficienza, cioè secondo criteri che favoriscono inevitabilmente le macchine. Se invece riconosciamo che la tecnica è uno strumento inscritto in un orizzonte umano più ampio, allora l'intelligenza artificiale può diventare un potenziamento delle possibilità umane e non la cancellazione del senso del lavoro.

Alle sei del mattino del 31 marzo Oracle ha licenziato tra i 12 e i 30 mila dipendenti con una mail, scelta probabilmente finalizzata a liberare miliardi di dollari da investire in infrastrutture per l'intelligenza artificiale. Se, come da lei di-

chiarato, la tecnica ha storicamente potenziato le capacità umane contribuendo a renderci sempre più performanti, come si può conciliare questo pensiero con il fenomeno sopracitato?

Il caso Oracle mostra un tratto tipico delle grandi trasformazioni tecnologiche: la tecnica aumenta enormemente la potenza operativa e, nello stesso tempo, produce ridistribuzioni spesso traumatiche del lavoro. Questo non è un fenomeno nuovo. Ogni rivoluzione tecnica ha eliminato professioni, ne ha trasformate altre e ne ha generate di nuove. La differenza attuale consiste nella velocità del processo, nella scala globale delle infrastrutture digitali e nella concentrazione del potere economico che accompagna lo sviluppo dell'intelligenza artificiale.

Quando affermo che la tecnica ha storicamente potenziato le capacità umane non intendo sostenere che ogni applicazione tecnica produca automaticamente benessere sociale. La tecnica aumenta possibilità operative, ma il modo in cui queste possibilità vengono organizzate dipende da scelte economiche, politiche e culturali. La macchina a vapore ha aumentato enormemente la produzione industriale, ma ha anche prodotto sfruttamento operaio; la rete digitale ha ampliato l'accesso all'informazione, ma ha anche favorito nuove concentrazioni monopolistiche; l'automazione industriale ha ridotto la fatica fisica, ma ha anche reso più fragile una parte del lavoro umano.

Il caso Oracle è emblematico perché mostra la logica dominante dell'attuale capitalismo tecnologico. Il licenziamento di migliaia di persone non viene presentato come una crisi dell'azienda, ma come una strategia di riallocazione delle risorse verso le infrastrutture dell'intelligenza artificiale. Questo significa che la priorità non è semplicemente ridurre costi, ma aumentare la capacità computazio-

nale e la posizione competitiva nel mercato globale dell'IA. Ci troviamo dunque davanti a un passaggio storico nel quale il capitale investe sempre più nelle macchine che trattano informazione e sempre meno nel lavoro umano tradizionale.

Questo produce un paradosso molto forte. Gli strumenti di intelligenza artificiale sono costruiti grazie a enormi patrimoni di conoscenza collettiva: dati prodotti dagli utenti, testi scritti da esseri umani, immagini generate da culture storiche, linguaggi sedimentati nel tempo. Tuttavia, il risultato economico di questo processo rischia di tradursi nell'espulsione dal lavoro proprio di coloro che hanno contribuito alla costruzione di quel patrimonio cognitivo. In altre parole, l'intelligenza artificiale nasce da una cooperazione sociale amplissima ma i benefici tendono a concentrarsi in poche piattaforme.

In questo senso, il problema non è la tecnica in quanto tale, ma l'uso economico e politico della tecnica. Pensare di arrestare lo sviluppo tecnologico sarebbe illusorio e probabilmente dannoso. La storia mostra che gli esseri umani hanno sempre sviluppato strumenti per aumentare le proprie capacità operative. Il punto decisivo consiste piuttosto nel comprendere come redistribuire i vantaggi prodotti da questi strumenti.

Se la produttività cresce enormemente grazie all'automazione, allora occorre interrogarsi su chi beneficia di quella crescita. Se i guadagni derivanti dall'intelligenza artificiale vengono concentrati esclusivamente nelle grandi corporation tecnologiche, la conseguenza sarà un aumento della disuguaglianza e della precarizzazione. Se invece quella produttività viene accompagnata da nuove forme di tutela sociale, di accesso alla formazione e di ridistribuzione economica, allora l'automazione può liberare tempo e risorse invece di produrre

esclusione.

C'è poi un altro elemento importante. Molte narrazioni contemporanee presentano il licenziamento tecnologico come un fenomeno inevitabile, quasi naturale. In realtà ogni trasformazione tecnica è mediata da decisioni istituzionali. La durata del lavoro, le forme contrattuali, l'accesso alla formazione, la fiscalità delle piattaforme e la regolazione dei dati non sono effetti automatici della tecnica, ma risultati di scelte collettive. Pensare il contrario significa attribuire alla tecnologia una forza metafisica che la tecnica non possiede.

Per questa ragione, la questione decisiva non consiste nel chiedersi se l'intelligenza artificiale sostituirà il lavoro umano, ma quale forma assumerà la cooperazione tra esseri umani e sistemi artificiali. La tecnica ha sempre trasformato il lavoro, ma il significato sociale di questa trasformazione dipende dalla capacità politica di orientarla. Più importante che arrestare la tecnica è costruire istituzioni capaci di trasformare l'aumento di potenza tecnica in aumento di possibilità sociali.

Lei ha menzionato che il "grande problema del rapporto distorto con la tecnologia è quello di togliere la speranza". Se, in ambito politico, come lei sostiene, questo può aprire la strada a derive autoritarie, quali conseguenze può avere nel contesto tecnologico e, in particolare, nello sviluppo dell'intelligenza artificiale?

Quando parlo della perdita della speranza come effetto di un rapporto distorto con la tecnologia, mi riferisco al fatto che molti individui percepiscono il progresso tecnico come un processo impersonale rispetto al quale non esiste possibilità di intervento. Se la tecnica appare come

una forza inevitabile che decide il destino del lavoro, delle relazioni e perfino della conoscenza, allora cresce la sensazione di impotenza. Gli individui incominciano a pensarsi non più come soggetti capaci di orientare il mondo tecnico, ma come elementi passivi che devono adattarsi a decisioni prese altrove.

In politica questa condizione favorisce derive autoritarie perché gli individui tendono ad affidarsi a figure che promettono protezione e semplificazione. Quando il futuro viene percepito come minaccia permanente, aumenta il desiderio di stabilità immediata anche a costo della libertà. La perdita della speranza produce allora una riduzione dell'orizzonte democratico, perché il cittadino rinuncia progressivamente alla convinzione di poter incidere sulle trasformazioni collettive.

Nel contesto dell'intelligenza artificiale il rischio assume una forma particolare. Se gli esseri umani incominciano a percepirsi come inferiori rispetto ai sistemi artificiali sul piano cognitivo e decisionale, può emergere una delega crescente delle capacità critiche. Non si tratta soltanto di affidare alle macchine compiti tecnici, ma di accettare che criteri di valutazione, selezione e interpretazione vengano definiti automaticamente. Un esempio è offerto dagli algoritmi che organizzano l'accesso alle informazioni: quando un sistema decide quali contenuti mostrare, quali relazioni rendere visibili e quali priorità imporre, la struttura dell'esperienza collettiva viene progressivamente modellata da procedure opache.

C'è poi un aspetto culturale ancora più profondo. L'intelligenza artificiale produce spesso un effetto di deresponsabilizzazione cognitiva. Se un sistema appare capace di rispondere a ogni domanda, di produrre immagini, diagnosi, sintesi o decisioni, cresce la tentazione di rinun-

ciare allo sforzo dell'interpretazione personale. Il rischio allora non consiste soltanto nell'automazione del lavoro, ma nell'automazione del giudizio. Una società che smette di esercitare capacità critiche diventa più vulnerabile alla manipolazione informativa e alla concentrazione del potere simbolico.

Questo fenomeno è particolarmente evidente nelle piattaforme digitali contemporanee. I sistemi algoritmici non si limitano a distribuire informazioni: modellano attenzione, desideri e comportamenti. Quando gli individui ricevono continuamente contenuti selezionati da

procedure automatiche, l'esperienza del mondo tende a diventare sempre più filtrata. In questo senso, l'intelligenza artificiale non modifica soltanto l'economia o il lavoro, ma interviene direttamente nella costruzione dell'immaginario collettivo.

La conseguenza più problematica non è dunque una ribellione delle macchine, scenario spesso cinematografico, ma la riduzione della capacità umana di iniziativa. Una società che considera inevitabile ogni decisione tecnica tende a rinunciare alla discussione pubblica sui fini. Se tutto viene interpretato come necessità tecnologica, allora non c'è più spazio per la deliberazione politica, per il conflitto democratico e per la scelta collettiva.

In questo senso, la speranza non è un sentimento ingenuo o puramente psicologico. La speranza è la convinzione che il mondo tecnico possa essere orientato e trasformato.

Conservare questa convinzione significa ricordare che l'intelligenza artificiale non è un destino metafisico: è un insieme di strumenti prodotti storicamente da esseri umani e dunque modificabili da esseri umani. La questione decisiva consi-

ste allora nel mantenere aperta la possibilità dell'intervento umano sulla tecnica.

Per questa ragione, il problema fondamentale non è se l'intelligenza artificiale diventerà più potente degli esseri umani, ma se gli esseri umani conserveranno la capacità di attribuire senso al proprio mondo tecnico. Una civiltà che delega integralmente alle macchine l'organizzazione della conoscenza e della decisione rischia di perdere progressivamente fiducia nella propria iniziativa storica. Una civiltà che invece riconosce la tecnica come prodotto umano può utilizzare l'intelligenza artificiale come strumento di ampliamento delle

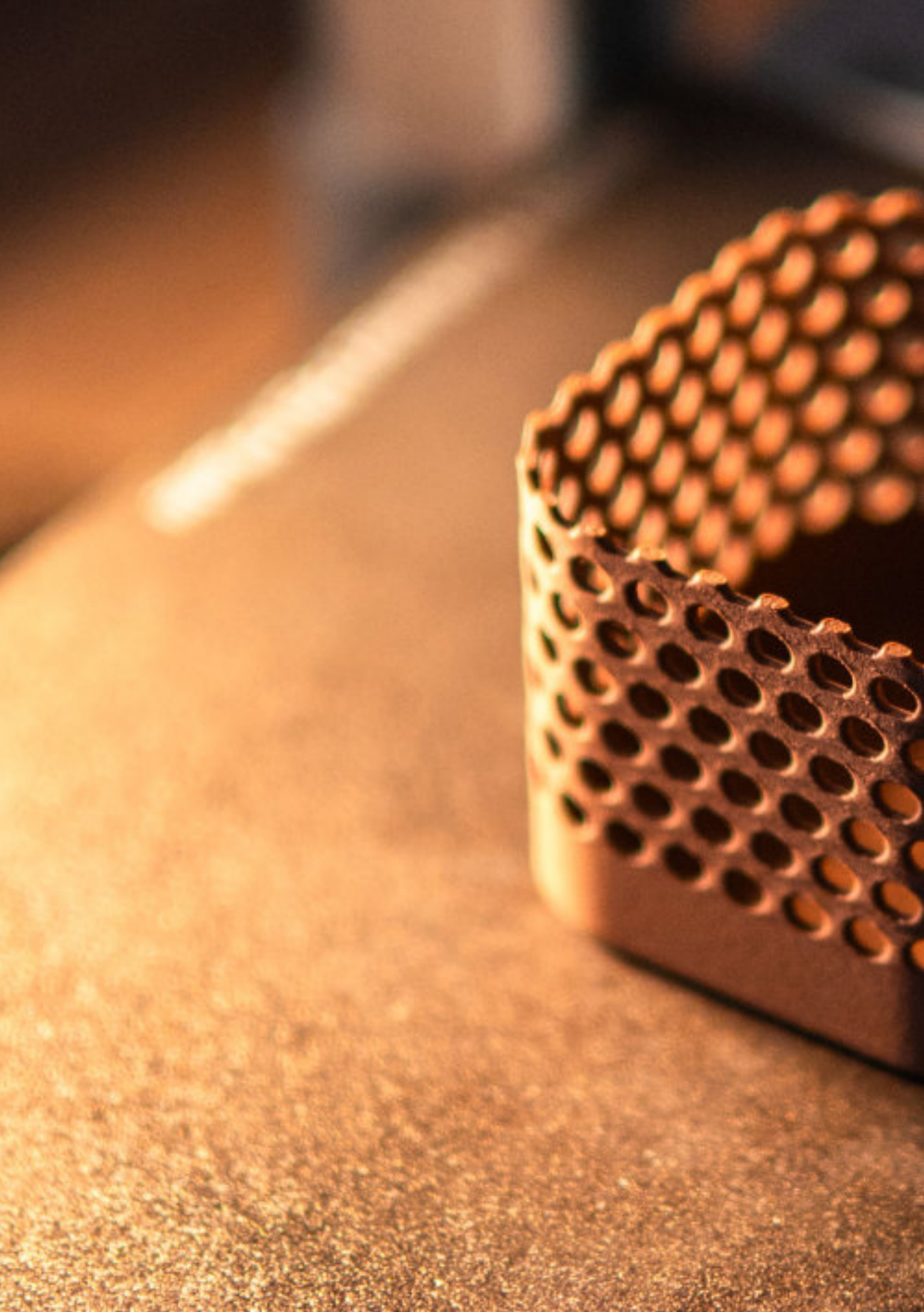
possibilità collettive invece che come principio di passivizzazione sociale.

4.4 Decidere è ancora un atto umano

Tre questioni trasversali emergono come particolarmente urgenti per il designer contemporaneo. Quando un'opera nasce dall'interazione tra intenzione umana e generazione algoritmica, chi ne risponde? Fin dove può spingersi la delega prima che il designer smetta di essere l'autore del proprio lavoro e diventi semplicemente l'occasione in cui il lavoro accade? Infine, in un paesaggio in cui pochi soggetti privati detengono il controllo delle infrastrutture creative, come si preserva l'autonomia professionale e culturale? Come si esercita quella responsabilità epistemica?

Abitare l'era post-AI significa allora, per il designer, assumere consapevolmente un ruolo critico di regia e mediazione. Il progettista non coincide più soltanto con l'esecutore materiale dell'opera, ma con il soggetto che orienta il processo, ne definisce i limiti e si assume la responsabilità delle visioni che esso incorpora. Ogni prompt, ogni scelta visiva e ogni decisione strutturale esprimono infatti una precisa idea del mondo e delle relazioni umane che si intendono rendere visibili e condivisibili.

Questa esigenza ha trovato di recente un'eco autorevole anche fuori dal perimetro disciplinare del design. Nella sua prima lettera enciclica, *Magnifica Humanitas* (15 maggio 2026), dedicata alla custodia della persona umana nel tempo dell'intelligenza artificiale, papa Leone XIV



Bambu Lab

progetto

5.1 Prima dell'algoritmo: cosa significava progettare

Per comprendere la portata del cambiamento attuale, è necessario volgere lo sguardo ai modelli metodologici che hanno definito il design nel secolo scorso. Analizzare i paradigmi tradizionali ci permette di cogliere le continuità e le rotture che l'intelligenza artificiale introduce nel flusso di lavoro creativo.

5.1.1 L'evoluzione degli strumenti di progettazione

Gli strumenti non sono mai neutrali. Ogni tecnologia di rappresentazione e modellazione che il designer ha adottato nel corso del tempo ha portato con sé non solo nuove possibilità operative, ma nuovi modi di pensare il progetto: ha modificato il ritmo del processo creativo, ridefinito le competenze richieste e spostato il confine tra ciò che era possibile immaginare e ciò che era possibile realizzare. Ricostruire questo paradigma è la premessa necessaria per valutare con precisione la portata del cambiamento che l'intelligenza artificiale sta introducendo nel design. Solo a partire da una comprensione solida del processo consolidato è possibile distinguere ciò che i modelli generativi stanno effettivamente trasformando da ciò che, invece, rimane, o dovrebbe rimanere, una prerogativa del giudizio umano.

Il disegno come pensiero

Disegnare non è mai stato soltanto rappresentare. Prima che esistesse qualsiasi strumento digitale, lo schizzo su carta era il luogo in cui il progetto prendeva forma e il processo stesso attraverso cui l'idea si chiariva. La matita costringeva a prendere decisioni di proporzione e di forma e ogni segno era insieme un'ipotesi e una verifica. Questa natura cognitiva del disegno spiega perché la padronanza tecnica della rappresentazione, come geometria descrittiva, resa materica, rendering manuale, costituissero il nucleo dell'identità professionale del designer industriale²⁴².

²⁴² Agnese Azzola et al., *Generative AI in the Design Process*, FrancoAngeli, 2025.

Il flusso di lavoro tradizionale era però lento per definizione. Gli schizzi su carta dovevano essere scannerizzati, rifiniti digitalmente e composti in layout di presentazione: tempi significativi dedicati non tanto alla progettazione quanto alla comunicazione delle idee. L'introduzione delle tavolette grafiche ha progressivamente affiancato il disegno analogico, offrendo maggiore flessibilità senza rinunciare alla gestualità del tratto²⁴³.

L'avvento del Computer-Aided Design

L'ingresso del computer nello studio di design non ha prodotto subito una rottura. Inizialmente era poco più di un tavolo da disegno digitale: stessi gesti, stesso risultato, schermo al posto della carta. Con il nuovo millennio, la crescente competenza informatica ha portato i designer a indagare i processi sottostanti al funzionamento degli strumenti digitali di uso quotidiano, e l'uso consapevole del computer ha stimolato un nuovo tipo di modellazione basato su algoritmi: procedure sistematiche in grado di esplorare varianti e relazioni che il disegno manuale non avrebbe potuto gestire²⁴².

Il risultato è un ecosistema software oggi molto articolato: strumenti vettoriali per la comunicazione visiva (Illustrator, Affinity Designer), modellatori parametrici per la geometria di prodotto (SolidWorks, Fusion 360, Rhinoceros, Grasshopper), modellatori poligonali per le forme organiche (Blender, Maya), software di rendering per la visualizzazione (KeyShot). A questi si aggiungono strumenti di editing fotografico, piattaforme collaborative e software per la creazione di presentazioni. Padroneggiare questo ecosistema è diventato parte integrante della competenza del designer, spostando il valore professionale dalla manualità del tratto alla capacità di navigare ambienti digitali diversificati²⁴³.

²⁴³ Esa Tammisto, *Usage of Artificial intelligence in industrial design processes*, Lappeenranta-Lahti University of Technology LUT, 2025

Una transizione stratificata

Sarebbe però un errore leggere questa evoluzione come una sostituzione lineare. Il disegno a mano non è scomparso dalla pratica del design industriale: resiste come strumento privilegiato nelle fasi esplorative, quando la velocità del gesto e la libertà dall'interfaccia consentono di generare e scartare idee con una rapidità che nessun software eguaglia. Allo stesso modo, il modello fisico rimane indispensabile per verificare dimensioni reali, ergonomia e funzionalità meccanica: aspetti che il digitale simula, ma che solo la materia rende tangibili e misurabili.

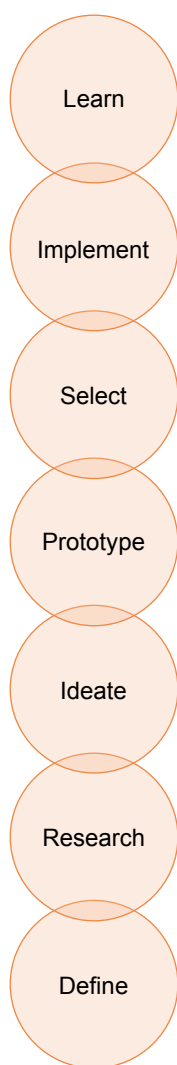
Il processo standard del design industriale è quindi ibrido per natura: gli schizzi migliori vengono trasferiti in modelli 3D, i modelli vengono renderizzati per la comunicazione con il cliente, i prototipi fisici vengono costruiti per testare ciò che lo schermo non può mostrare. Ogni fase richiede iterazioni, e il progetto avanza attraverso un continuo rimbalzo tra analogico e digitale, tra esplorazione e verifica²⁴³. È questo equilibrio consolidato, e non il solo strumento digitale, che l'intelligenza artificiale si trova oggi a ridisegnare.

5.1.2 L'uomo al centro: il paradigma dello Human-Centered Design

Il progresso degli strumenti non ha cambiato solo il modo in cui il designer lavora, ma ha contribuito a ridefinire anche la domanda fondamentale attorno a cui il progetto si organizza: non più soltanto come si fa, ma per chi si fa e perché. È in questo spostamento di prospettiva che affonda le radici il paradigma dello Human-Centered Design che mette l'essere umano, con i suoi bisogni, le sue fragilità e i suoi contesti, al centro del processo progettuale.

L'essenza del design

Il design si colloca all'intersezione tra arte, ingegneria e tecnologia, funzionando come spazio di traduzione tra linguaggi che altrimenti non si parlerebbero. Questa posizione ibrida definisce anche la sua complessità intrinseca: non è una disciplina che ricerca soluzioni univoche, ma che negozia tra variabili eterogenee e spesso parzialmente incompatibili. Il valore del progettista sta proprio nella capacità di tenere insieme questa tensione equilibrando i compromessi tra utenti, produttori e con-



²⁴⁴ Joana Cerejo, *Human-Centered AI Design Process – Part 1*, Medium, 2021.

Figura 5.1: Le sette fasi del processo di progettazione (Simon's 7 Stages of the design process): Define, Research, Ideate, Prototype, Select, Implement e Learn²⁴⁴. Grafico rielaborato.

testi d'uso senza creare nuove frizioni mentre ne risolve altre.

Al centro di questo processo sta l'empatia, non tanto come attitudine generica ma come strumento operativo: comprendere il nucleo di un problema significa saper osservare i punti di attrito che utenti e stakeholder sperimentano, spesso senza saperli articolare. Una soluzione progettuale di successo genera valore lungo l'intero ciclo di vita del prodotto, dalla produzione all'installazione, garantendo un elevato livello di usabilità e qualità dell'esperienza per tutti gli attori coinvolti nelle diverse fasi di interazione con il prodotto²⁴³.

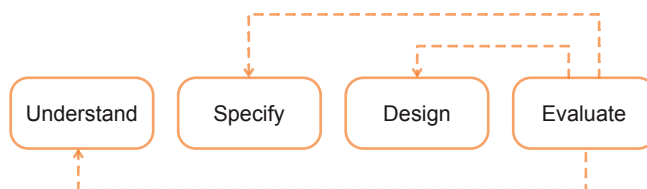
La nascita dello Human-Centered Design

Le radici dello Human-Centered Design affondano nella seconda metà del Novecento, quando i designer iniziarono a focalizzare il proprio sguardo dalla forma del prodotto ai bisogni di chi lo avrebbe usato. Un momento fondativo risale al 1969, quando il premio Nobel Herbert Simon pubblicò *The Sciences of the Artificial*, delineando per la prima volta un processo strutturato di Design Thinking in sette fasi che avrebbe influenzato i modelli progettuali per decenni²⁴⁴.



Figura 5.2: Lo scienziato politico ed economista Herbert Simon, premio Nobel e pioniere assoluto dell'intelligenza artificiale per i suoi studi sui processi decisionali e sulla simulazione del pensiero umano. Immagine da "Herbert Simon", Wikipedia.

Figura 5.3: Diagramma del processo di User-Centred Design, articolato nelle quattro fasi iterative²⁴³. Grafico rielaborato.



La prima fase consiste nell'analisi approfondita del contesto d'uso, con l'obiettivo di comprendere le caratteristiche degli utenti, le attività che svolgono, gli strumenti che utilizzano e l'ambiente in cui avviene l'interazione con il prodotto o il sistema. Questa attività permette di acquisire una conoscenza accurata delle esigenze, delle aspettative e delle criticità che caratterizzano l'esperienza d'uso. La seconda fase riguarda l'identificazione e la specificazione dei requisiti degli utenti. Qui vengono definiti i bisogni che la soluzione progettuale dovrà soddisfare, i problemi che dovrà risolvere e le modalità attraverso cui il prodotto potrà generare valore per gli utenti. I requisiti costituiscono quindi il riferimento fondamentale per orientare le successive decisioni progettuali. Nella terza fase vengono sviluppate una o più soluzioni progettuali, sotto forma di concept, prototipi o versioni preliminari del prodotto. Tali soluzioni devono rispondere in modo coerente ai bisogni e alle problematiche emerse durante le fasi precedenti, traducendo i requisiti individuati in proposte concrete. La quarta fase prevede la valutazione delle soluzioni progettuali rispetto ai criteri definiti attraverso l'analisi del contesto e dei requisiti utente. Attraverso attività di verifica e test con gli utenti, è possibile individuare eventuali criticità o aspetti non adeguatamente affrontati. Qualora emergano elementi da migliorare, il processo viene reiterato, consentendo di affinare progressivamente la soluzione progettuale.

Un elemento distintivo dell'approccio UCD è infatti il costante coinvolgimento degli utenti lungo l'intero processo di progettazione. I designer mantengono il focus sui

Figura 5.4: Diagramma di Venn sulle tre lenti dell'innovazione del Design Thinking: l'intersezione tra desiderabilità (Desirability), fattibilità (Feasibility) e sostenibilità economica (Viability))²⁴³. Grafico rielaborato.

bisogni, sugli obiettivi e sulle caratteristiche degli utilizzatori finali, integrando frequentemente attività di ricerca, osservazione e valutazione con gli utenti stessi. In questo modo, il prodotto finale risulta maggiormente allineato alle reali esigenze delle persone a cui è destinato²⁴³.

A partire dagli anni Novanta, User-Centered Design e Human-Centered Design vengono spesso usati in modo intercambiabile per indicare l'integrazione degli utenti finali nella progettazione. Progressivamente, tuttavia, il secondo inizia a distinguersi: mentre l'UCD mantiene un orientamento prevalentemente tecnico-funzionale, l'HCD sposta l'asse verso una dimensione più ampia, che include empatia, contesto sociale e coinvolgimento degli stakeholder oltre agli utenti diretti²⁴⁴.

5.1.3 Il Design Thinking come framework metodologico

È in questo clima che negli anni Novanta IDEO evolve l'approccio HCD nel Design Thinking, portandolo a una popolarità che negli anni successivi avrebbe attraversato il mondo del business, dell'educazione e della ricerca. Come ha sintetizzato Tim Brown, presidente esecutivo di IDEO, il Design Thinking è un approccio centrato sull'uomo all'innovazione che integra i bisogni delle persone, le possibilità della tecnologia e i requisiti per il successo aziendale. Questa definizione racchiude i tre pilastri fondanti del metodo il cui punto di intersezione definisce lo spazio progettuale ottimale: desiderabilità, fattibilità tecnologica e sostenibilità economica²⁴³.

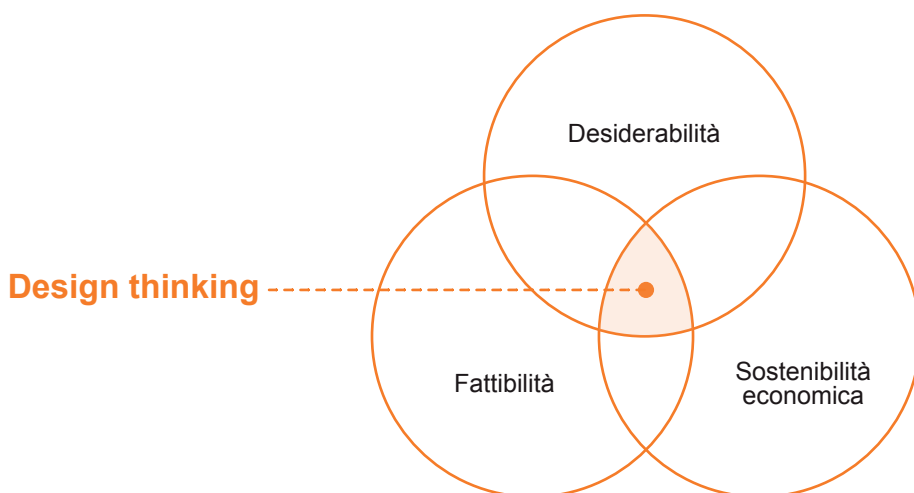




Figura 5.5: Il designer industriale Tim Brown, co-fondatore di IDEO e teorico del design thinking. Immagine da “Tim Brown”, Global Speakers Bureau.

L'integrazione di questi tre aspetti consente di sviluppare prodotti, servizi, processi e strategie capaci di rispondere efficacemente ai bisogni delle persone, mantenendo al contempo la propria realizzabilità e sostenibilità nel tempo. Per questa ragione, il Design Thinking è oggi considerato uno strumento trasversale applicabile a contesti differenti e non esclusivamente riservato alla professione del designer.

“Design thinking is a human-centered approach to innovation that draws from the designer’s toolkit to integrate the needs of people, the possibilities of technology, and the requirements for business success.” — Tim Brown, Executive Chair di IDEO

È importante precisare che, nonostante la sua ampia diffusione in ambito progettuale e manageriale, il Design Thinking non dispone di una definizione univoca e universalmente condivisa. Può essere interpretato come un

approccio, una metodologia, una strategia di innovazione o, più in generale, come una modalità di pensiero orientata alla risoluzione dei problemi attraverso processi creativi e collaborativi.

Un principio centrale dell'approccio consiste nell'adozione di quella che viene spesso definita una "mentalità da principiante", ovvero la capacità di mantenere curiosità, apertura e disponibilità all'apprendimento continuo. Ciò implica evitare conclusioni premature, considerare l'ambiguità come una risorsa progettuale e accogliere prospettive differenti durante il processo di sviluppo delle soluzioni. Il Design Thinking si configura inoltre come un processo intrinsecamente iterativo, basato sulla sperimentazione continua e sul progressivo affinamento delle idee. In tale prospettiva, il fallimento non viene interpretato come un esito negativo, bensì come una preziosa opportunità di apprendimento. Testare rapidamente ipotesi e prototipi nelle fasi iniziali del progetto consente infatti di individuare criticità e correggere tempestivamente eventuali errori, riducendo il rischio di investire risorse significative in soluzioni inefficaci²⁴³.

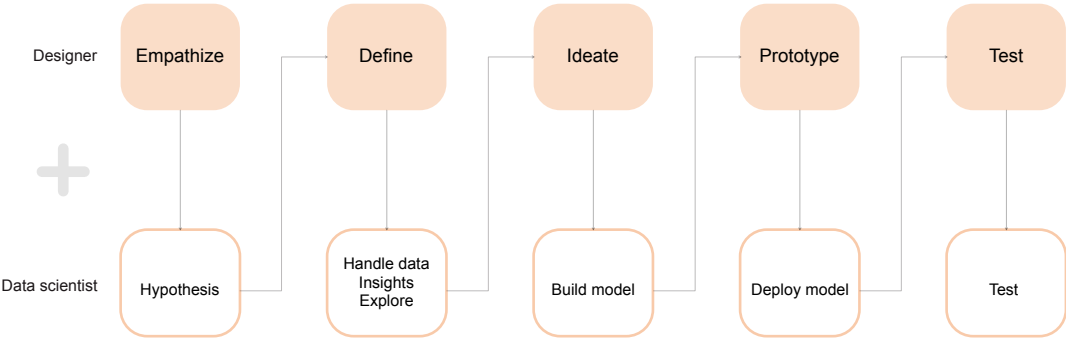


Figura 5.6: Mappatura del processo di Human-centered AI Design: parallelismo e interazione tra le fasi metodologiche dei Designer (linea superiore) e le corrispondenti attività tecniche dei Data Scientists (linea inferiore)²⁴⁵. Grafico rielaborato.

Il modello che ne ha determinato la diffusione globale è quello sviluppato dall'Hasso-Plattner Institute of Design di Stanford (d.school), articolato in cinque fasi iterative.

La prima fase, Empathize, ha l'obiettivo di sviluppare una comprensione approfondita degli utenti e del contesto in cui essi operano. Attraverso attività di osservazione, interviste, shadowing e altre tecniche di ricerca qualitativa, i progettisti raccolgono informazioni sui comportamenti, sulle motivazioni, sui bisogni e sulle difficoltà delle persone. L'empatia rappresenta il fondamento dell'intero approccio, poiché consente di superare supposizioni e preconcetti per comprendere il problema dal punto di vista degli utenti.

La seconda fase, Define, consiste nella sintesi e nell'interpretazione dei dati raccolti durante la ricerca. Le informazioni vengono analizzate e organizzate per individuare pattern ricorrenti, bisogni significativi e opportunità di intervento. Il risultato di questa fase è una definizione chiara e focalizzata del problema progettuale, spesso formulata come una problem statement o una domanda progettuale che guiderà le attività successive.

La terza fase, Ideate, è dedicata alla generazione di idee e possibili soluzioni. In questo momento del processo viene privilegiato il pensiero divergente, incoraggiando la produzione di un ampio numero di proposte senza giudizi o limitazioni premature. Tecniche come il brainstorming permettono di esplorare prospettive differenti e di individuare approcci innovativi alla risoluzione del problema definito.

Nella quarta fase, Prototype, le idee considerate più promettenti vengono trasformate in rappresentazioni concrete e sperimentabili. I prototipi possono essere semplici schizzi cartacei, modelli fisici, mock-up o prototipi digitali interattivi. L'obiettivo non è realizzare una soluzione definitiva, ma rendere tangibili le idee per verificarne rapidamente il potenziale e raccogliere feedback significativi.

La quinta fase, Test, prevede la valutazione dei prototipi attraverso il coinvolgimento diretto degli utenti finali. L'osservazione delle loro interazioni e la raccolta delle loro opinioni consentono di identificare punti di forza, criticità e opportunità di miglioramento. I risultati ottenuti possono confermare le ipotesi iniziali oppure evidenziare la necessità di rivedere il problema stesso, generando nuove conoscenze che alimentano ulteriori iterazioni del processo.

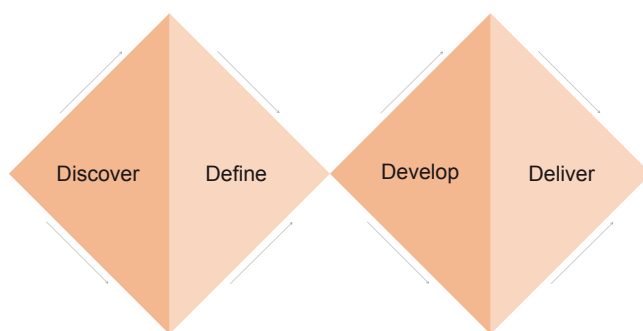
Sebbene vengano spesso presentate in modo sequenziale, tali fasi non costituiscono un percorso lineare ma, anzi, presentano continui ritorni e revisioni che consentono di affinare progressivamente la comprensione del problema e la qualità delle soluzioni proposte²⁴⁵.

5.1.4 Il Double Diamond: origine e architettura del modello

²⁴⁵ Joana Cerejo, *Human-Centered AI Design Process – Part 2: Empathize & Hypothesis*, Medium, 2021.

Il Double Diamond è probabilmente il modello di processo progettuale più conosciuto e utilizzato nello sviluppo di prodotto. La sua rappresentazione grafica è semplice,

Figura 5.7: Visualizzazione del framework del Double Diamond sviluppato dal Design Council, strutturato nelle quattro fasi sequenziali di Discover, Define, Develop e Deliver²⁴³. Grafico rielaborato.



concisa e facilmente comprensibile. Il metodo è stato sviluppato dal Design Council britannico nel 2003, a seguito di un'attività di ricerca e analisi dei diversi processi di design impiegati da aziende e organizzazioni. I titoli principali del grafico sono ampi e indicano le macro-fasi del progetto, ciascuna delle quali comprende diverse attività dipendenti dalla tipologia del progetto e del lavoro svolto²⁴³.

Il modello si articola in quattro fasi distribuite su due diamanti, ciascuno dei quali rappresenta un ciclo di pensiero divergente seguito da un pensiero convergente. La prima fase, Discover, prevede la messa in discussione del problema iniziale e conduce alla ricerca dei bisogni degli utenti. La seconda fase, Define, consiste nell'analizzare i risultati della scoperta per produrre un design brief accurato che definisca chiaramente le sfide da affrontare. La terza fase, Develop, include lo sviluppo, il test e il perfezionamento di molteplici soluzioni in parallelo, per garantire un terreno fertile all'innovazione. La quarta fase, Deliver, prevede il restringimento delle idee in un'unica soluzione, raffinata e preparata per il lancio²⁴³.

5.1.5 Il passaggio dallo Human-Centered Design allo Human-Centered Artificial Intelligence

L'avvento dell'intelligenza artificiale non si innesta su un paradigma progettuale immobile. Lo Human-Centered Design aveva già spostato il baricentro dalla forma al bisogno, dal prodotto alla persona. Ma l'AI introduce una pressione ulteriore, interrogando il ruolo stesso della tecnologia nel processo e imponendo di decidere se l'intelligenza artificiale debba essere uno strumento nelle mani del designer o un sistema da progettare, a sua volta, attorno all'essere umano.

I limiti del paradigma pre-algoritmico

Il processo progettuale centrato sull'essere umano, nella sua configurazione tradizionale, presenta limiti intrinseci che la crescente complessità dei sistemi progettati ha reso sempre più evidenti. Le tecniche tradizionali di ricerca utente, pur fondamentali, richiedono settimane o mesi per produrre insight significativi. L'esplorazione dello spazio delle soluzioni è limitata dalla capacità cognitiva del singolo progettista o del team. I cicli di iterazione comportano tempi e costi che la pressione competitiva rende sempre più difficili da sostenere. La ricerca ha dimostrato che l'adozione dell'intelligenza artificiale nel design riduce i tempi di ciclo fino al 25%, il numero di iterazioni di prototipo del 37,5% e aumenta la soddisfazione degli utenti del 20%²⁴⁶.

²⁴⁶ Yingchao Feng, *Artificial Intelligence-Assisted Product Design: From Concept to Prototype*, Pioneer Academic Publishing Limited, 2024.

Il modello classico del Double Diamond, con la sua struttura a due cicli di divergenza e convergenza, appare progressivamente rigido di fronte alle nuove dinamiche progettuali. I risultati empirici supportano la tesi centrale secondo cui i designer accolgono positivamente i benefici dell'AI generativa nelle fasi iniziali con nuove prospettive, ed efficienza accelerata, ma, una volta prodotta un'abbondanza di idee, incontrano difficoltà nel gestirne la quantità e nel prendere decisioni. Aggiungere semplicemente strumenti AI al classico processo Double Diamond non è sufficiente per risolvere questa criticità: i designer desiderano chiaramente una maggiore strutturazione per la fase di convergenza²⁴⁷.

²⁴⁷ Xilin Tang et al., *Human-AI Collaboration Within Industrial Design*, AHFE, 2025.

Il passaggio allo Human-Centered Artificial Intelligence

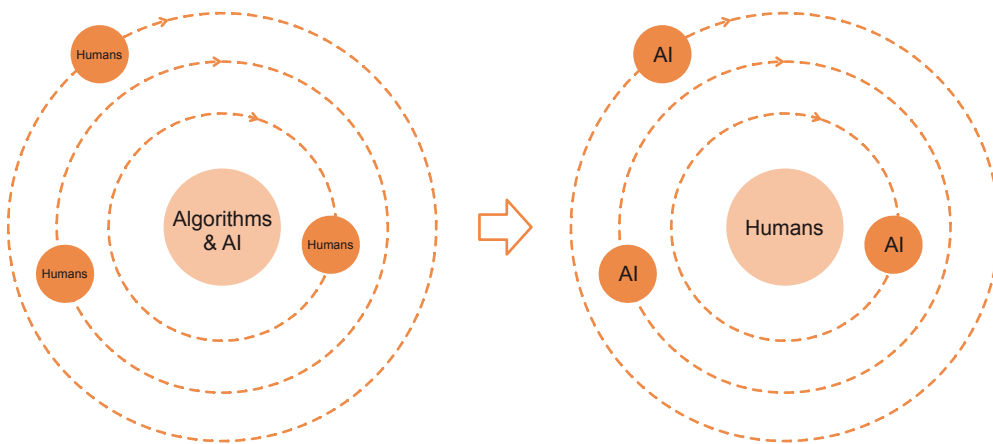
La transizione dallo Human-Centered Design allo Human-Centered Artificial Intelligence (HCAI) rappresenta un passaggio concettuale di portata fondamentale: non si tratta più solo di progettare pensando alle persone, ma di sviluppare sistemi di intelligenza artificiale che mettano l'essere umano al centro, affrontando esplicitamente questioni di equità, responsabilità, interpretabilità e trasparenza.

Fino a tempi recenti, ricercatori e sviluppatori si sono concentrati sulla costruzione di algoritmi e sistemi AI, misurando le prestazioni delle macchine e celebrando ciò che l'intelligenza artificiale era in grado di fare autonomamente. Lo Human-Centered Artificial Intelligence (HCAI) inverte questa prospettiva: sposta il centro di attenzione dagli algoritmi agli esseri umani, mettendo l'esperienza utente al primo posto, misurando le presta-

zioni delle persone e valorizzando le nuove capacità che la tecnologia mette a loro disposizione. Non si tratta di una revisione tecnica, ma di un cambio di paradigma: i sistemi AI del futuro dovranno promuovere dignità umana, uguaglianza e sicurezza, e questo non sarà possibile finché il focus resterà sulle metriche algoritmiche piuttosto che sui bisogni delle persone.

Figura 5.8: Il cambio di paradigma introdotto dall'HCAI dal modello "tecnocentrico" (a sinistra), in cui l'essere umano orbita attorno agli algoritmi, al modello "umanocentrico" (a destra), in cui l'Intelligenza Artificiale è posta al servizio delle persone²⁴⁴. Grafico rielaborato.

Questo passaggio non invalida i principi del Design Thinking, che restano centrati sulle persone, ma ne estende la portata. L'intelligenza artificiale consente di superare i limiti storici dei processi progettuali ad alta intensità umana: limiti di scala, di velocità, di capacità di elaborare grandi quantità di dati. La sinergia tra design e AI non è quindi una minaccia alla disciplina, ma una condizione che richiede nuove pratiche multidisciplinari, costruite insieme a data scientist e ingegneri²⁴⁴.



5.1.6 La complessità come condizione permanente del progetto

Cosa succede quando si cerca di automatizzare un processo che, per definizione, non ha una soluzione giusta? Il design ha sempre operato in uno spazio di ambiguità strutturale: problemi mal definiti, vincoli contraddittori, obiettivi che cambiano mentre si lavora per raggiungerli. Prima di capire cosa l'intelligenza artificiale trasforma nella pratica progettuale, vale la pena chiedersi cosa stava cercando di fare il designer e perché farlo non è mai stato, né sarà, una questione di algoritmi.

La ricerca progettuale e la risoluzione dei problemi

La ricerca sui metodi progettuali ha attraversato tre fasi storiche distinte, ciascuna delle quali ha ampliato la comprensione del processo di design. In una prima fase, la risoluzione dei problemi era fondata sul ragionamento logico e sull'uso di algoritmi computazionali per simulare il pensiero umano: un approccio che Simon stesso ha contribuito a teorizzare, esplorando i modelli razionali e intuitivi alla base del problem-solving complesso. A partire dagli anni Ottanta, l'attenzione si è spostata sui meccanismi cognitivi del designer: come struttura i problemi, come genera soluzioni, come gestisce l'incertezza. Infine, il crescente utilizzo dei computer ha rinnovato questi modelli, arricchendo la comprensione del pensiero progettuale e aprendo nuove prospettive per la formazione e la pratica²⁴⁸.

²⁴⁸ Yangluxi Li et al., *A Review of Artificial Intelligence in Enhancing Architectural Design Efficiency*, Applied Sciences MDPI, 2025.

La ricerca di filosofia e ontologia del design ha chiarito che il nucleo dell'attività creativa è la strutturazione del problema, non la sua soluzione. Simon la descrive come un processo continuo: la definizione degli obiettivi e dei sotto-obiettivi non precede il progetto, ma lo accompagna dall'inizio alla fine. Bryan Lawson, architetto e ricercatore britannico, distingue tra vincoli esterni, determinati dalle condizioni oggettive, e vincoli interni, derivati dall'esperienza del progettista. Il teorico del design Nigel Cross sottolinea la natura irriducibilmente ibrida del design, che richiede di integrare scienza e arte, rigore analitico e sensibilità creativa²⁴⁸. È questa complessità strutturale, ovvero il fatto che i problemi progettuali siano per definizione "mal definiti", interconnessi e privi di soluzioni univoche, a rendere il design un terreno particolarmente impegnativo per qualsiasi tentativo di automazione.

AI maturity model

Comprendere a quale livello di maturità si trovi l'intelligenza artificiale è fondamentale per valutare con realismo le trasformazioni in atto nel design. L'AI Maturity Model di Gartner articola questa traiettoria in cinque livelli progressivi. Il primo, Awareness, corrisponde a una fase di interesse esplorativo, ancora caratterizzata da aspettative eccessive e scarsa consapevolezza dei limiti reali della tecnologia. Il secondo, Active, è quello della sperimentazione: l'AI viene testata principalmente in contesti di data science, senza ancora integrarsi nei flussi di lavoro produttivi. Il terzo livello, Operational, segna il passaggio dalla sperimentazione alla produzione: l'AI inizia a ge-

nerare valore concreto attraverso l'ottimizzazione dei processi e l'innovazione di prodotti e servizi. Il quarto, Systemic, rappresenta un salto qualitativo: l'AI viene utilizzata in modo pervasivo per la trasformazione digitale dei processi e delle catene del valore, abilitando nuovi modelli di business. Il quinto e ultimo livello, Transformational, descrive una condizione in cui l'AI non è più uno strumento esterno ma parte integrante del DNA organizzativo, capace di ridefinire strutturalmente il modo in cui un'organizzazione opera e crea valore.

La maggior parte degli studi e delle organizzazioni di design si colloca oggi tra il secondo e il terzo livello: in una fase di sperimentazione attiva, con i primi strumenti generativi in produzione, ma ancora lontana da un'integrazione sistemica. Questo posizionamento ha implicazioni concrete per la professione: le decisioni prese ora su quali strumenti adottare, come formare i designer, quali competenze valorizzare e quali processi automatizzare sono orientamenti strategici che determineranno la traiettoria della disciplina nei prossimi anni. Chi si trova oggi al livello due e non costruisce consapevolmente un percorso verso il livello tre rischia di subire la trasformazione invece di guidarla²⁴⁹.

²⁴⁹ Liu, Y. et al., *How Can Artificial Intelligence Transform the Future Design Paradigm and Its Innovative Competency Requisition: Opportunities and Challenges*,

Il paradigma pre-algoritmico come punto di partenza
La ricostruzione del paradigma progettuale pre-algori-



per governare criticamente l'introduzione dell'AI nella pratica progettuale e per rispondere alla domanda che il settore non potrà evitare di porsi: quale valore specifico apporta il designer in un mondo in cui la macchina genera, e l'umano decide.

Figura 5.9: Il modello di maturità dell'Intelligenza Artificiale articolato in cinque livelli evolutivi, dall'interesse iniziale alla completa integrazione nel DNA aziendale²⁴⁹. Grafico rielaborato.



5.2 Dopo il modello generativo: cosa significa decidere

L'ingresso dei modelli generativi ha trasformato il processo progettuale da una sequenza di azioni lineari a un sistema di interazioni circolari. Progettare oggi significa saper navigare l'incertezza e governare l'abbondanza di opzioni, spostando il focus dalla creazione della forma alla regia del senso.

5.2.1 Nuovi strumenti di intelligenza artificiale nel processo progettuale

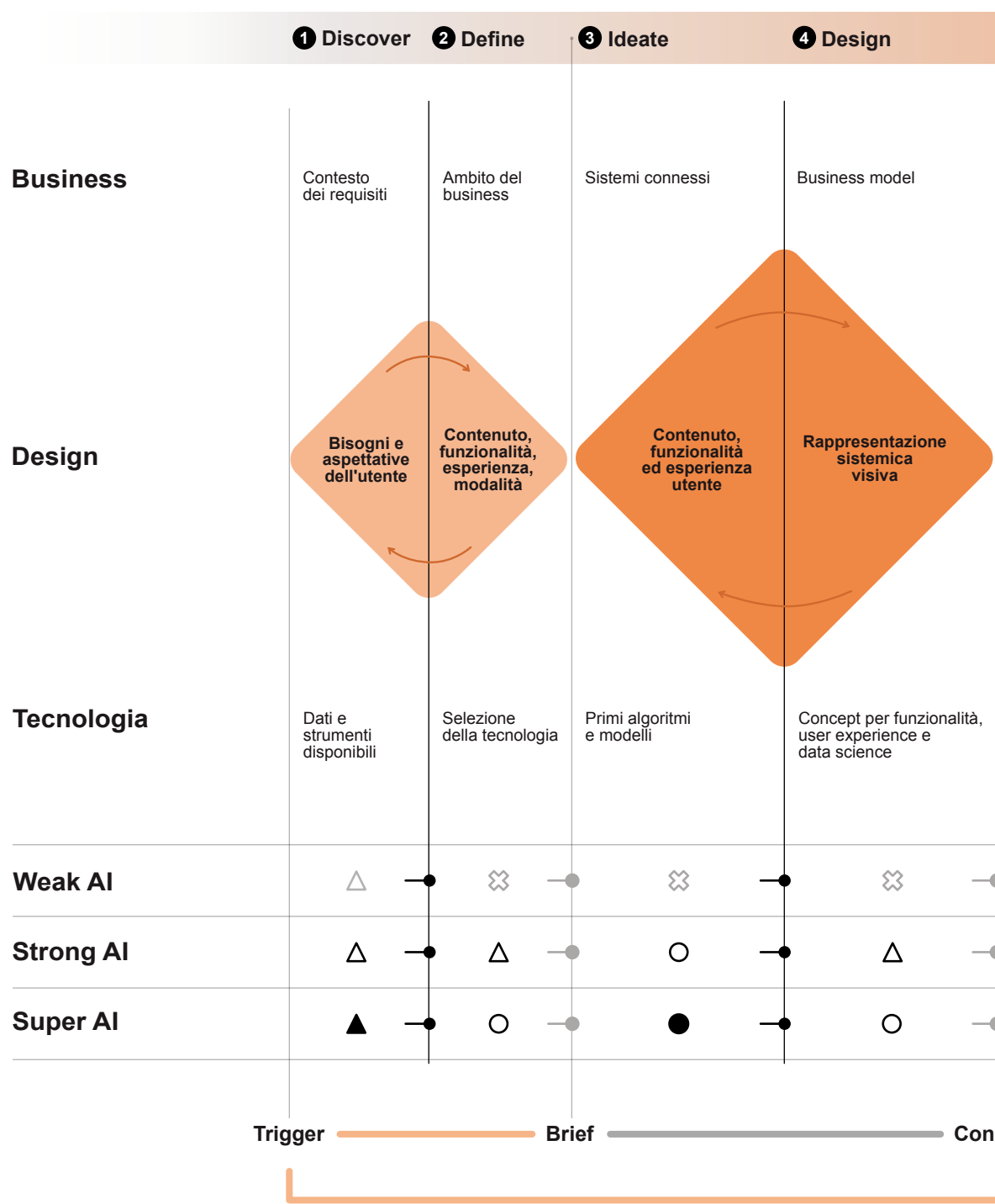
Comprendere come l'intelligenza artificiale entra nel processo progettuale richiede di scendere dal livello delle dichiarazioni generali a quello delle fasi concrete: non chiedersi se l'AI cambia il design, ma dove lo cambia, in che misura e con quali implicazioni per il ruolo del designer. La mappatura che segue non è un esercizio teorico, ma uno strumento per leggere con precisione una trasformazione già in corso.

L'AI nelle fasi del processo: una mappatura operativa

Il modello AI-Embedded Design Thinking Process espande la struttura del Double Diamond in dieci fasi sequenziali, mappando per ciascuna il contributo congiunto di tre dimensioni (business, design e tecnologia) e distinguendo tre livelli crescenti di intervento dell'IA (Weak,

Figura 5.10: Matrice del processo di AI-Embedded Design Thinking e relativo ruolo tecnologico. Lo schema mappa 10 fasi sequenziali incrociandole con le prospettive di Business, Design e Technology, e valutando l'impatto algoritmico in base al livello di complessità dell'IA (Weak, Strong, Super AI)²⁴⁹. Grafico rielaborato.

Il processo di Design Thinking con AI integrata



Significato dei simboli:

- △

effetto strumentale debole
- ▲

effetto strumentale forte
- effetto di co-collaborazione
- punto di selezione e decisione
- △

effetto strumentale medio
- ⊗

nessun effetto
- effetto di co-creazione
- punto di selezione e decisione



Strong, Super). Il livello debole indica un effetto strumentale di supporto, il livello forte una collaborazione attiva, il livello super una co-creazione in cui la macchina partecipa alla generazione stessa delle soluzioni. Questa gradazione non è uniforme lungo il processo: alcune fasi risultano ancora prevalentemente umane, altre sono già oggi terreno di intervento profondo dell'IA, e altre ancora rappresentano scenari prospettici verso cui il settore si sta muovendo. I punti di selezione e decisione disseminati lungo il processo ricordano la natura iterativa del progetto, segnalando i momenti in cui il team valuta se procedere, correggere o tornare indietro²⁴⁹.

La prima fase, Discover, ha come oggetto progettuale la comprensione dei bisogni e delle aspettative degli utenti, mentre sul piano tecnologico si lavora con i dati e gli strumenti già disponibili. In questa fase la Weak AI esercita un effetto strumentale basso, la Strong AI uno medio e la Super AI alto: tutte contribuiscono all'elaborazione di grandi volumi di dati provenienti da piattaforme digitali e social media. Strumenti di Natural Language Processing consentono di analizzare simultaneamente migliaia di recensioni, post e risposte a sondaggi, identificando pattern emotivi e bisogni latenti che le interviste tradizionali da sole difficilmente intercetterebbero.

La seconda fase, Define, traduce i risultati della scoperta in una direzione progettuale condivisa, lavorando su contenuto, funzionalità, esperienza e modalità di interazione, mentre sul piano tecnologico avviene la selezione delle tecnologie da impiegare. Qui la Weak AI non interviene, mentre la Strong AI esercita un effetto strumentale medio e la Super AI introduce per la prima volta un effetto di collaborazione: è il momento in cui sistemi avanzati possono contribuire alla sintesi degli insight raccolti e alla formulazione del brief, supportando il team nella definizione delle priorità progettuali. Questa fase si chiude con la produzione del Brief, primo grande punto di convergenza del processo.

Le fasi tre e quattro, Ideate e Design, corrispondono al secondo diamante del modello tradizionale e rappresentano il cuore generativo del processo. Nella fase di Ideate, l'oggetto progettuale riguarda contenuto, funzionalità ed esperienza utente, mentre tecnologicamente si lavora con i primi algoritmi e modelli. La Weak AI è assente, la Strong AI ha un effetto di collaborazione, la Super AI un effetto di co-creazione piena: è qui che modelli generativi come Midjourney e DALL-E vengono impiegati per produrre rapidamente immagini, mockup e concept visivi a

partire da input in linguaggio naturale, comprimendo i tempi di esplorazione e ampliando lo spazio delle possibilità. Nella fase di Design, l'attenzione si sposta sulla rappresentazione visiva e sistemica delle soluzioni, con la tecnologia orientata alla modellazione per funzionalità, esperienza utente e data science. La Weak AI resta assente, la Strong AI esercita un effetto strumentale medio, la Super AI uno di collaborazione: il processo si fa più selettivo, e le idee generate nella fase precedente vengono strutturate e valutate. Questa coppia di fasi si chiude con il Concept, secondo punto di convergenza del modello.

La quinta fase, Prototype, vede il concept trasformarsi in artefatti tangibili per la verifica dell'esperienza e del contenuto, impiegando prototipi realizzati con tecnologia reale. Weak AI e Strong AI esercitano rispettivamente un effetto strumentale basso e medio, mentre la Super AI contribuisce con un effetto di collaborazione. In questa fase algoritmi di ottimizzazione topologica supportano la riduzione dello spreco di materiale riprogettando la struttura interna dei componenti: secondo i dati del McKinsey Global Institute, l'applicazione dell'IA nella progettazione e ingegneria di prodotto può ridurre il numero di iterazioni di prototipo fino al 30%, un dato che non rappresenta un semplice incremento di efficienza, bensì un ripensamento strutturale del flusso di lavoro.

La sesta fase, Test, sottopone i prototipi a verifica attraverso test utente, interviste e workshop, con tecnologie di A/B testing. La Weak AI non interviene, mentre Strong AI e Super AI esercitano rispettivamente un effetto strumentale medio e di collaborazione. L'IA automatizza la raccolta dati, cattura in tempo reale interazioni dettagliate attraverso telecamere e sensori, e processa grandi volumi di feedback con algoritmi di NLP, identificando problemi di esperienza utente non precedentemente scoperti. I risultati attestano un aumento dell'efficienza della raccolta dati del 50% e un miglioramento della profondità dell'analisi del sentiment. Questa fase si chiude con i Development Requirements, terzo punto di convergenza.

Le fasi sette e otto, Develop e Implement, portano il progetto verso la realizzazione definitiva. In Develop, il focus è sullo sviluppo finale di contenuto ed esperienza, supportato da algoritmi e data feed: Weak AI e Strong AI hanno intervengono con un effetto strumentale basso e medio, mentre la Super AI di collaborazione. In Implement, l'oggetto progettuale riguarda i test finali e la Su-

per AI ha effetto forte: è la fase in cui i sistemi si connettono e il prodotto acquista la sua forma operativa definitiva. Il processo raggiunge lo step Go Live, quarto punto di convergenza del modello.

Le ultime due fasi, Operate e Scale, segnalano una discontinuità concettuale rispetto al processo tradizionale, che tipicamente si concludeva con il lancio del prodotto. In Operate, il sistema viene alimentato da feedback continui per migliorare contenuto ed esperienza utente, con il contributo soprattutto della Super AI. In Scale, il progetto si espande verso nuovi mercati e aree di business attraverso nuovi modelli e la moltiplicazione di quelli esistenti: qui la Weak AI è assente, la Strong AI mantiene un effetto strumentale medio, e la Super AI raggiunge il suo massimo intervento con un effetto forte. Il processo si "chiude" con gli Updates, che riconducono il flusso verso l'inizio in un ciclo continuo di aggiornamento²⁴⁹.

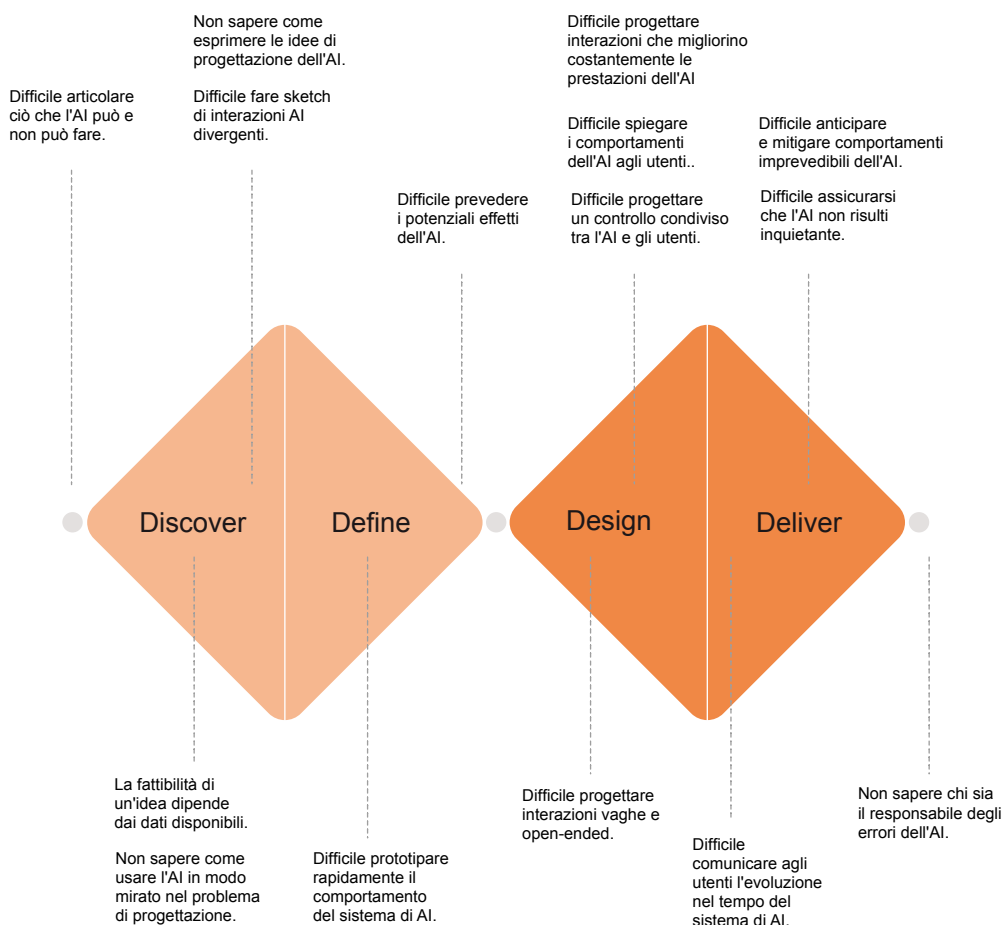
5.2.2 Problemi emergenti dall'utilizzo dell'AI nelle fasi del processo

L'adozione dell'intelligenza artificiale nei processi di design non si esaurisce nella promessa di nuove possibilità: porta con sé, in egual misura, una serie di tensioni che mettono alla prova tanto gli strumenti quanto chi li utilizza. Prima ancora di padroneggiare le potenzialità dei modelli generativi, il designer si trova a confrontarsi con limiti inattesi, resistenze operative e sfide che non trovano ancora risposta nei metodi consolidati della disciplina. I problemi emergenti non sono marginali né transitori: si collocano al cuore del processo progettuale e ne interrogano le fondamenta.

Difficoltà tecniche e progettuali

L'utilizzo dell'intelligenza artificiale nella progettazione non produce soltanto vantaggi: genera al contempo una serie di attriti che si manifestano a livelli diversi e si distribuiscono lungo l'intero flusso di lavoro. Tornando a prendere come riferimento il Double Diamond, si osserva che nella fase di Discover il designer fatica a capire come impiegare l'IA in modo intenzionale, e la fattibilità delle idee risulta vincolata alla disponibilità dei dati: una limitazione che si manifesta già nella fase generativa, quando la natura generalista dei modelli, privi di comprensione della terminologia e del contesto specifico del design, genera inefficienza e rielaborazione continua²⁵⁰.

²⁵⁰ O'Regan, S., *The Design Difficulties of Human-AI Interaction*, 2024.



²⁵¹ Sedef Süner-Pla-Cerdà et al., *Designer experiences and perspectives on the role of generative AI in industrial design*, Springer Nature, 2025.

Figura 5.11: Mappatura delle complessità e delle sfide di progettazione nell'interazione uomo-IA (Human-AI Interaction) collocate lungo le quattro fasi del framework Double Diamond (Discover, Define, Design, Deliver)²⁵⁰. Grafico rielaborato.

In Define diventa difficile tradurre le intenzioni progettuali in forme comprensibili ai modelli e prevedere il comportamento del sistema: l'interazione basata su prompt verbali risulta spesso inadeguata per professionisti abituati a pensare e comunicare in modo visivo, e i modelli faticano a mantenere coerenza visiva tra più generazioni, rendendo ardua la costruzione di narrazioni progettuali con elementi formali uniformi²⁵¹.

Con il passaggio al secondo diamante i problemi continuano ad aumentare. Una sequenza di tentativi di sviluppo formale di un cruscotto automobilistico attraverso ControlNet Lora e DreamStudio illustra con evidenza il limite della consistenza: le varianti generate dall'IA propongono angolazioni, altezze e gerarchie di layout differenti per ogni generazione, senza che il designer possa controllare con precisione quali elementi conservare e quali modificare. Il risultato è un output utilizzabile come riferimento formale per movimenti di superficie, ma non come base dimensionalmente affidabile per ergonomia,

Dashboard design [depth LORA]



bus interior cockpit design, futuristic, simplicity, layered design language, closer to dashboard, birdview



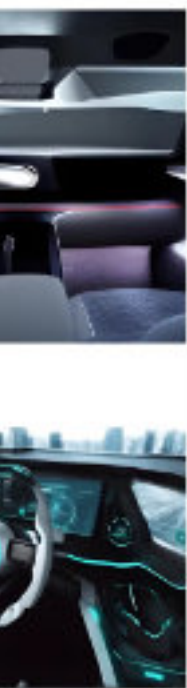
²⁵² Yu-Min Fang, *The Role of Generative AI in Industrial Design: Enhancing the Design Process and Education*, International Conference on Innovation, Communication and Engineering, 2023.

Figura 5.12: Sperimentazione e prove formali per il design di plance e cruscotti d'autoveicoli (dashboard form trials) attraverso l'uso di Intelligenza Artificiale Generativa: confronto tra i risultati ottenuti con ControlNet LoRA (in alto) e DreamStudio (in basso)²⁵¹.

linee di giunzione e dettagli produttivi²⁵⁰.

A questo si aggiunge una tendenza strutturale dei modelli a privilegiare l'innovazione estetica rispetto a quella funzionale: i generatori producono output visivamente accattivanti ma spesso privi di considerazioni ergonomiche, costruttive o di producibilità. Questa asimmetria è ben documentata: nella fase divergente l'IA risulta efficace per la generazione rapida di idee e il supporto a stili espressivi, mentre nella fase convergente è meglio affidata agli strumenti esistenti e all'expertise del designer. Le sfide principali includono la difficoltà nel distinguere tendenze di design sottili, l'inefficienza per ottimizzazioni e il rischio di prodotti strutturalmente e stilisticamente troppo maturi ma funzionalmente vuoti, in cui la forma ha superato la sostanza progettuale²⁵². Progettare interazioni aperte e non deterministiche, immaginare sistemi che migliorino continuamente le proprie prestazioni, spiegare il comportamento dell'IA agli utenti e definire forme di controllo condiviso tra macchina e persona rimangono sfide ancora in larga parte irrisolte.

In Deliver emergono infine le questioni più spinose: l'imprevedibilità dei comportamenti, il rischio che le interazioni risultino percepite come invasive, la difficoltà di comunicare agli utenti l'evoluzione del sistema nel tempo



e l'assenza di una chiara attribuzione di responsabilità in caso di errore²⁵⁰.

Il sovraccarico decisionale

Se il problema tecnico riguarda la qualità degli output, quello cognitivo ne riguarda la quantità. L'IA è straordinariamente efficace nell'espandere le possibilità, generando un numero di alternative che eccede di gran lunga la capacità umana di valutazione. In uno studio condotto con sette studenti di design industriale, tre su sette hanno riportato di sentirsi "sopraffatti" o "incerti su come filtrare" le opzioni generate dall'IA. Un partecipante ha osservato che "è incredibile avere così tante opzioni disponibili immediatamente, ma ho bisogno di un modo per selezionare e ottimizzare le migliori senza perderne nessuna", mentre un altro ha notato che "il tempo che dedichiamo a organizzare gli output dell'IA potrebbe essere superiore al tempo che impiegheremmo a elaborare meno idee da soli". Il dato è significativo: cinque studenti su sette hanno stimato un risparmio di tempo di almeno il 25%, ma il beneficio viene parzialmente neutralizzato dalla complessità della fase di convergenza²⁴⁷.

Il progettista si trova così di fronte a un paradosso: dispone di nuovi strumenti più potenti, ma non di nuove competenze per governarne l'output²⁵³.

Problemi comunicativi

Uno specifico ordine di problemi riguarda il rapporto tra il pensiero visivo del designer e il linguaggio testuale richiesto dai modelli. Scomporre intenzioni visive complesse come composizione, materiali e dettagli formali in descrizioni testuali coerenti impone un'operazione cognitiva controintuitiva, che genera una distanza sistematica tra ciò che viene immaginato e ciò che viene prodotto²⁴².

Il problema ha già stimolato i primi tentativi di soluzione: modelli di predizione delle parole categorizzano automaticamente ciascun termine del prompt in classi semantiche (medium, composizione, soggetto) rendendo visibili le categorie mancanti e suggerendo alternative. Rimane però aperta una questione più profonda: l'IA non si comporta come un esecutore fedele, ma come un interlocutore con un proprio lessico, radicato nei pattern dei dati di addestramento, con cui il designer deve imparare a negoziare²⁵⁴.

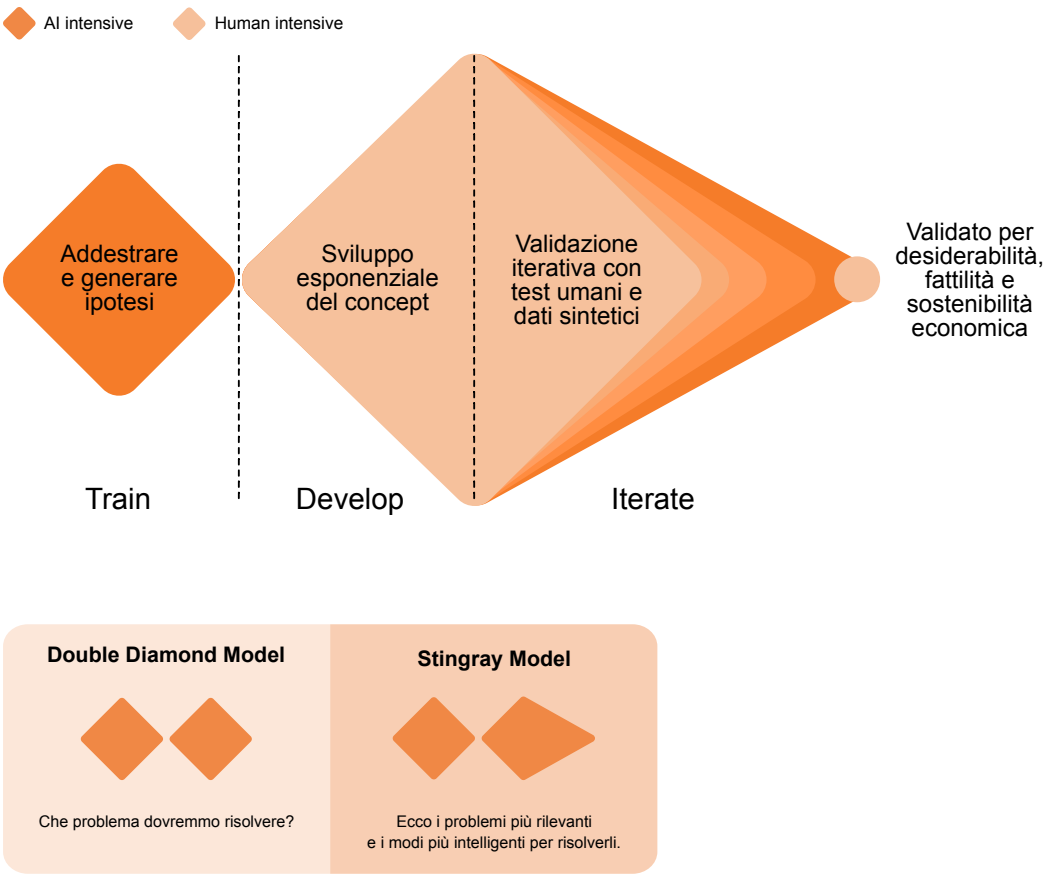
²⁵³ Jeremiah Folorunso et al., *Product Design: The Evolving Role of Generative AI in Creative Workflows*, International Journal of Scientific and Management Research, 2025.

²⁵⁴ Henriikka Vartiainen et al., *Emerging human-technology relationships in a co-design process with generative AI*, Elsevier, 2024.

5.2.3 Lo Stingray Model

Figura 5.13: Stingray Model sviluppato da Board of Innovation per l'integrazione dell'IA nel design. Lo schema illustra il passaggio dal classico Double Diamond a un flusso asimmetrico suddiviso nelle macro-fasi di Train, Develop e Iterate, evidenziando l'equilibrio tra attività ad alta intensità umana e algoritmica²⁵⁵. Grafico rielaborato.

Il Double Diamond ha il merito della chiarezza e della semplicità, e rimane un riferimento fondamentale nella didattica e nella pratica professionale. Tuttavia, l'introduzione dell'IA ne mette in evidenza un limite strutturale: la sequenza lineare di soli due cicli divergenza-convergenza risulta progressivamente rigida in un'epoca che richiede maggiore agilità e iterazione in tempo reale. L'IA eccelle nelle fasi divergenti, generando un volume e una diversità di idee senza precedenti, ma la sfida si concentra nella convergenza: decidere quali direzioni perseguire tra le centinaia di alternative prodotte richiede strutture di filtraggio e sintesi che il Double Diamond tradizionale non prevede esplicitamente. La ricerca conferma che semplicemente aggiungere strumenti di IA al processo classico non è sufficiente a risolvere questo squilibrio: i designer richiedono con chiarezza una maggiore strutturazione per le fasi convergenti²⁴⁷.



²⁵⁵ Netgroup, *The New Playbook for Innovation: Stingray Model Design*, 2024.

Lo Stingray Model si propone come alternativa al Double Diamond, costruita specificamente per l'era dell'IA generativa. Il modello è adattivo, rapido e articolato in tre fasi: Train, in cui il designer definisce obiettivi deliberati e addestra l'IA ad accogliere i bisogni del consumatore, consentendo ai team di concentrarsi sullo sviluppo di concept desiderabili, fattibili e sostenibili; Develop, in cui problemi e soluzioni vengono esplorati simultaneamente grazie all'ideazione assistita dall'IA, permettendo un'analisi esponenziale ed esaustiva dello spazio del problema; Iterate, in cui le soluzioni tangibili vengono validate attraverso processi iterativi che incorporano test sintetici e strumenti abilitati dall'IA per valutare i concept rispetto a fattori diversi, tra cui sostenibilità e tendenze di mercato. La differenza fondamentale rispetto al Double Diamond risiede nella concezione del rapporto con i dati: nel modello tradizionale la ricerca empatica è centrale e richiede mesi di indagine qualitativa, mentre nello Stingray Model l'IA sintetizza vaste quantità di dati sui consumatori in poche ore, evidenziando pattern e opportunità che gli esseri umani potrebbero non cogliere²⁵⁵.

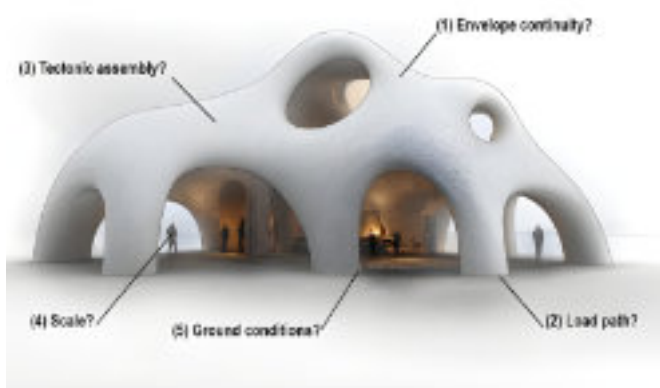
5.2.4 L'IA non semplifica il progetto, ma lo modifica

L'idea che l'intelligenza artificiale semplifichi il processo progettuale è una semplificazione essa stessa. In realtà, l'IA ristruttura le condizioni entro cui il progettista opera, spostando il focus dalla definizione diretta della soluzione alla costruzione dei sistemi e dei criteri attraverso cui le soluzioni emergono. Il progettista non genera più una singola soluzione, ma definisce un ambiente controllato di possibilità, all'interno del quale la logica generativa, anziché la forma finale, diventa l'oggetto della progettazione²⁵³.

Il plausibility gap: quando l'immagine precede la sostanza

Una conseguenza particolarmente rilevante della ristrutturazione del processo riguarda il plausibility gap: il divario tra la plausibilità visiva di un output generato dall'IA e la sua effettiva traducibilità in un artefatto costruibile. Per rendere questo problema misurabile, è utile ricorrere a una classificazione in quattro regimi della progettazione digitale, basandosi su un confronto di tipo ontologico: come viene prodotta la prima evidenza del progetto, come questa acquista legittimità e se rimane verificabile lungo una catena di prove. Il Regime 0 corrisponde

all'oggetto fisico, il Regime 1 al modello digitale, il Regime 2 al codice parametrico, il Regime 3 allo spazio latente ovvero il regime proprio dell'IA generativa.



Nei Regimi 1 e 2 esiste una continuità interna garantita da un generatore leggibile, ovvero il modello 3D o il codice parametrico, che media tra forma e fabbricazione attraverso una sequenza verificabile: dalla morfogenesi alla generazione di varianti, dalla selezione alla raziona-

Categoria di Attrito	Cosa rimane indeterminato nell'artefatto basato prima sull'immagine	Cosa deve essere deciso/ricostruito (Esempi di impegni richiesti)
(1) Ambiguità geometrica	L'immagine suggerisce un guscio continuo, ma la geometria esatta (continuità della curvatura, strategia dello spessore, condizioni dei bordi in corrispondenza di oculi e archi) rimane indecidibile solo dall'aspetto visivo.	Definire superfici/solidi espliciti (es. mesh NURBS), classe di continuità, logica dello spessore del guscio, profili delle aperture e tolleranze geometriche per la fabbricazione.
(2) Ambiguità strutturale	Il padiglione "si legge" come stabile, eppure il percorso del carico (dove poggia, come le forze si trasferiscono attraverso il guscio, cosa funge da arco/costolone/diaframma) non è specificato.	Stabilire il sistema strutturale (funzionamento a guscio vs. telaio a costoloni vs. ibrido), supporti/fondazioni, strategia di rinforzo e un diagramma verificabile di trasferimento del carico.
(3) Ambiguità materiale/tettonica	Gli indizi superficiali implicano un assemblaggio simile a muratura o piastrelle, ma la materialità è semanticamente plausibile piuttosto che tettonicamente definita (unità, giunti, fissaggi, impermeabilizzazione, logica termica).	Specificare il sistema dei materiali (unità/pannelli, giunti, sottostruttura), sequenza di assemblaggio, dettaglio delle giunzioni, impermeabilizzazione/drenaggio e requisiti prestazionali.
(4) Ambiguità metrica/di scala	Le figure umane forniscono una scala approssimativa, ma moduli, campate, raggi e tolleranze non sono ancorati metricamente; l'immagine non definisce un sistema dimensionale affidabile.	Fissare una griglia/modulo dimensionale, campate e spazi liberi, rapporti spessore-luce, tolleranze e vincoli costruttivi che stabilizzino il progetto al di là degli indizi di scala pittorici.
(5) Ambiguità di dettaglio	Le giunzioni critiche sono visivamente lisce, ma i dettagli chiave non sono risolti (contatto con il suolo, bordi degli oculi, transizioni tra rivestimento interno e pelle esterna, percorsi di drenaggio).	Risolvere i dettagli dei bordi, le condizioni di base, la strategia per l'umidità/drenaggio, la continuità delle finiture, le interfacce tra i sistemi e i requisiti di manutenibilità/sicurezza.

Figura 5.14: Processo di validazione tettonica nell'architettura generativa. A sinistra: l'artefatto originale selezionato dallo spazio latente. A destra: la relativa Friction map che evidenzia i cinque punti di ambiguità costruttiva necessari per la ricostruzione geometrica e la validazione reale²⁵⁶.

lizzazione tettonica. Nel Regime 3 questa continuità si interrompe: il punto di partenza è un'immagine ad alta plausibilità visiva prodotta senza geometria sottostante, e il lavoro progettuale si inverte. Ciò che nei regimi precedenti era il punto di arrivo (un artefatto geometricamente definito, strutturalmente coerente, metricamente leggibile) diventa qui un problema da risolvere a posteriori. Un'analisi degli output generati, annotata secondo cinque categorie di ambiguità (geometrica, strutturale, materica, scalare e di dettaglio) trasforma il gap da affermazione generica in procedura ispezionabile: l'immagine resta visivamente convincente mentre l'oggetto rimane incompleto come catena di evidenze²⁵⁶.

5.2.5 Il cambio di paradigma nell'interazione con le macchine: da command-based a intent-based

²⁵⁶ José Carlos López Cervantes & Cintya Eva Sánchez Morales, *Design in the Age of Predictive Architecture: From Digital Models to Parametric Code to Latent Space*, MDPI, 2026.

Tabella 5.1: Checklist della mappa d'attrito per la diagnosi del divario tra plausibilità visiva e realizzabilità tettonica. La tabella definisce gli elementi indeterminati negli artefatti image-first generati dall'IA e i relativi requisiti decisionali richiesti per la validazione costruttiva²⁵⁶.

Ogni tecnologia ridefinisce il modo in cui l'essere umano si relaziona con essa. L'intelligenza artificiale non fa eccezione: la sua introduzione nei processi progettuali non rappresenta semplicemente un aggiornamento degli strumenti disponibili, ma segna una discontinuità nel modo stesso in cui il designer interagisce con la macchina. Comprendere la natura di questo cambiamento, e le implicazioni che porta con sé, è condizione necessaria per orientarsi consapevolmente in un panorama professionale che si sta riconfigurando con rapidità.

La nuova relazione con il sistema

Per decenni l'interazione con i computer ha seguito un paradigma command-based: l'utente fornisce un input diretto per ottenere un risultato diretto, proseguendo in un ciclo di feedback bidirezionale fino al raggiungimento dell'esito desiderato. L'introduzione dell'intelligenza artificiale ha innescato un passaggio a un paradigma intent-based: anziché impartire una sequenza di comandi, il progettista esprime un intento al sistema e collabora con esso, raffinando gli input e guidando l'IA mentre questa interpreta, adatta e risponde in modo dinamico. Non si tratta più di interazioni puntuali, ma di relazioni continue in cui il sistema mantiene una memoria contestuale e un modello dell'utente, trasformando l'interfac-

cia da strumento reattivo a interlocutore proattivo²⁵⁷.

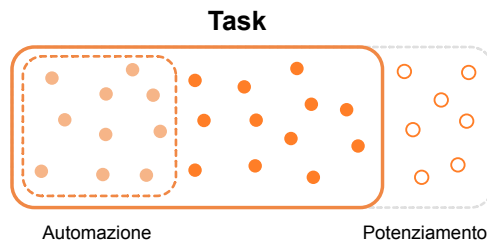
Questa trasformazione è visibile nella traiettoria dei principali software di progettazione. Il Neural CAD inaugurato da Autodesk con Fusion 2026 ne è l'esempio più esplicito: attraverso il linguaggio naturale è possibile tradurre intenti progettuali in geometrie solide editabili, mentre funzionalità come AutoConstrain assegnano automaticamente i vincoli meccanici. Figma consente di descrivere un'interfaccia a parole e ottenere un prototipo funzionale con codice generato. Copiloti come l'Ansys Engineering Copilot permettono di navigare processi di simulazione complessi attraverso una conversazione. In tutti questi casi, la competenza richiesta non è più "so usare questo strumento" ma "so dire a questo strumento cosa voglio ottenere": una distinzione apparentemente sottile, ma profondamente diversa sul piano cognitivo²⁵³.

Augment versus automate

Con il passaggio al paradigma intent-based emerge una tensione fondamentale che attraversa l'intero campo del design assistito dall'IA: la scelta tra potenziamento e sostituzione. La distinzione non è puramente semantica, ma implica visioni radicalmente diverse del ruolo della macchina nel processo creativo. Il modello elaborato da Brynjolfsson rappresenta questa tensione in modo efficace: i task si distribuiscono lungo un continuum che

²⁵⁷ Denny Royal, *Designing for an AI World*, Tietoevry, 2025.

Figura 5.15: Modello evolutivo dei task in seguito all'introduzione dell'Intelligenza Artificiale, che mostra attività destinate all'automazione completa e attività soggette ad aumento o potenziamento delle capacità umane²⁵⁸. Grafico rielaborato.



va dall'automazione completa in cui l'IA esegue senza supervisione umana fino al potenziamento, in cui la macchina amplifica le capacità del designer mantenendo il controllo umano al centro, fino a un'area ancora fuori dalla portata dell'IA, in cui il giudizio umano rimane insostituibile²⁵⁸.

Per comprendere con maggiore precisione dove si colloca il controllo umano in ciascuna interazione con la macchina, è utile fare riferimento al framework dei dieci livelli di automazione proposto da Sheridan e Verplank. La scala procede dal livello 1, in cui il sistema non offre alcu-

²⁵⁸ Pablo Fernández Vallejo, *Redefine Your Design Skills to Prepare for AI*, Nielsen Norman Group, 2024.

Tabella 5.2: I dieci livelli di automazione secondo il modello di Sheridan e Verplank, disposti in ordine crescente dal controllo umano totale (Livello 1) all'autonomia decisionale e operativa completa della macchina (Livello 10)²⁵⁹.

Livello	Descrizione dell'automazione
1	Il sistema non offre alcuna assistenza: gli esseri umani devono prendere tutte le decisioni e compiere tutte le azioni.
2	Il sistema offre un insieme completo di possibili decisioni/azioni che l'essere umano deve approvare.
3	Il sistema offre un insieme ristretto di possibili decisioni/azioni che l'essere umano deve approvare.
4	Il sistema offre una decisione/azione e un'alternativa che l'essere umano deve approvare.
5	Il sistema offre una singola decisione/azione che l'essere umano deve approvare.
6	Il sistema concede all'essere umano un tempo limitato per porre il veto prima dell'esecuzione automatica.
7	Il sistema esegue l'azione automaticamente, poi informa necessariamente l'essere umano.
8	Il sistema informa l'essere umano solo se espressamente richiesto.
9	Il sistema informa l'essere umano solo se decide di farlo.
10	Il sistema decide tutto, agisce in modo autonomo, ignorando l'essere umano.

²⁵⁹ Cerejo, J., *Designing for Automation vs Augmentation*, Medium, 2021.

na assistenza e tutte le decisioni spettano all'essere umano, fino al livello 10, in cui il sistema decide e agisce autonomamente ignorando l'utente. I livelli intermedi descrivono gradazioni significative: il sistema può offrire un insieme completo di opzioni tra cui scegliere, restringerle a una selezione, proporre un'unica alternativa, eseguire automaticamente lasciando un margine di veto, o informare l'utente solo a posteriori. Nella pratica del design attuale, la maggior parte degli strumenti IA si colloca tra i livelli 2 e 5, mantenendo il designer in una posizione di approvazione e selezione. La direzione di sviluppo, tuttavia, spinge progressivamente verso livelli più alti, rendendo urgente la definizione consapevole di dove si vuole che il controllo umano risieda²⁵⁹.

5.2.6 Il cambio di ruolo del designer: da creator a director

Se la generazione delle alternative non è più un'attività esclusivamente umana, quale diventa il contributo specifico del designer? La risposta non risiede nella sostituzione della figura professionale, bensì nella trasformazione delle sue responsabilità.

Dalla produzione alla curatela

Se tradizionalmente il processo progettuale richiedeva un coinvolgimento diretto nella produzione e nello sviluppo delle soluzioni, oggi una parte crescente delle attività generative può essere supportata o automatizzata dagli algoritmi. Strumenti capaci di produrre rapidamente immagini, varianti formali e proposte concettuali consentono infatti di esplorare un numero molto elevato di alternative in tempi ridotti, modificando il rapporto tra progettista e produzione dell'output. Il valore del designer non risiede più esclusivamente nella capacità di generare artefatti, ma sempre più nella capacità di orientare, selezionare e interpretare i risultati prodotti dall'IA. Il professionista assume quindi una funzione di curatela, intesa come attività di scelta, organizzazione e attribuzione di significato tra molteplici possibilità generate automaticamente. La progettazione si configura sempre meno come un processo lineare di produzione e sempre più come un'attività di direzione e governo della complessità, nella quale il designer definisce gli obiettivi, stabilisce i criteri di valutazione e costruisce coerenza tra le diverse proposte emerse.

Questa evoluzione si articola in cinque funzioni specifiche che il designer è chiamato ad assumere nella collaborazione con l'IA. Il progettista assume innanzitutto un ruolo analitico, fondamentale per identificare il problema progettuale, definire gli obiettivi e interpretare le esigenze degli utenti e del mercato. La qualità degli output generati dall'IA dipende infatti dalla capacità del designer di costruire un quadro progettuale chiaro e coerente, all'interno del quale gli strumenti generativi possano operare in modo efficace.

A questa dimensione si affianca il ruolo istruttivo, che riguarda la progettazione dell'interazione con l'Intelligenza Artificiale. Il designer è chiamato a tradurre obiettivi e requisiti in prompt efficaci, definendo vincoli, riferimenti stilistici e criteri di generazione. La competenza nel

prompt engineering diventa quindi una nuova forma di progettazione, attraverso la quale il professionista orienta il comportamento del sistema e ne sfrutta il potenziale creativo.

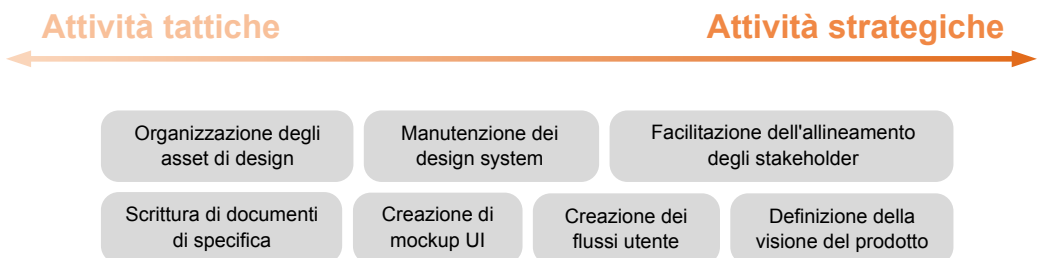
Una volta generati i risultati, emerge il ruolo valutativo, che consiste nell'analizzare criticamente le proposte prodotte dall'IA. Il designer deve selezionare le soluzioni più pertinenti, verificandone la qualità estetica, la coerenza con il brief, la fattibilità tecnica e l'allineamento con gli obiettivi strategici del progetto. L'Intelligenza Artificiale è infatti in grado di produrre numerose alternative, ma non possiede la capacità di attribuire valore progettuale alle diverse opzioni generate.

Accanto alla valutazione si sviluppa il ruolo sintetico, attraverso il quale il designer integra, modifica e combina i risultati ottenuti per trasformarli in proposte progettuali coerenti e significative. Più che accettare passivamente gli output prodotti dalla macchina, il progettista opera una rielaborazione critica che consente di collegare intuizioni differenti, risolvere contraddizioni e costruire una visione progettuale unitaria.

Infine, il designer assume un ruolo curatoriale, volto a garantire continuità, coerenza e identità all'intero processo creativo. Tale funzione non riguarda soltanto la selezione delle singole proposte, ma anche la capacità di costruire una narrazione progettuale coerente nel tempo, mantenendo allineati linguaggio formale, valori del brand e obiettivi strategici. In questa prospettiva il designer agisce come direttore creativo del sistema uomo-macchina, orchestrando i contributi dell'IA e assicurando che gli output generati convergano verso una visione progettuale condivisa²⁴².

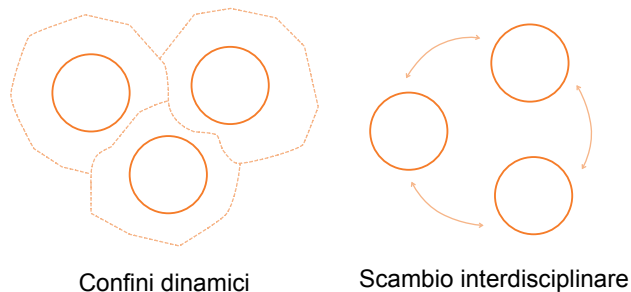
Task affidati all'AI e decisioni strategiche riservate al designer

Figura 5.16: Mappatura dei compiti del design in base alla complessità strategica. Lo schema evidenzia come le attività sul polo sinistro (tattico-esecutive) siano maggiormente esposte all'automazione algoritmica rispetto a quelle sul polo destro (strategico-relazionali)²⁵⁸. Grafico rielaborato.



La redistribuzione dei compiti tra IA e designer segue una logica che lo spettro delle attività progettuali rende visibile con chiarezza: da un lato le attività tattiche, come organizzare asset, scrivere specifiche, produrre mockup, tendono progressivamente verso l'automazione. Dall'altro le attività strategiche, tra cui facilitare l'allineamento tra stakeholder, definire la visione di prodotto, costruire la direzione creativa, rimangono saldamente in capo al designer. Come già detto, questa migrazione implica che il valore professionale del designer si sposti dalla capacità esecutiva alla capacità di giudizio, di sintesi e di visione. Non è l'intera professione ad essere messa in discussione dall'IA, ma soltanto i suoi aspetti esecutivi, lasciando intatto il nucleo creativo e metodologico. Per fare un paragone, si può dire che è come un interior designer che progetta lo spazio ma delega ad altri la tinteggiatura delle pareti²⁵¹.

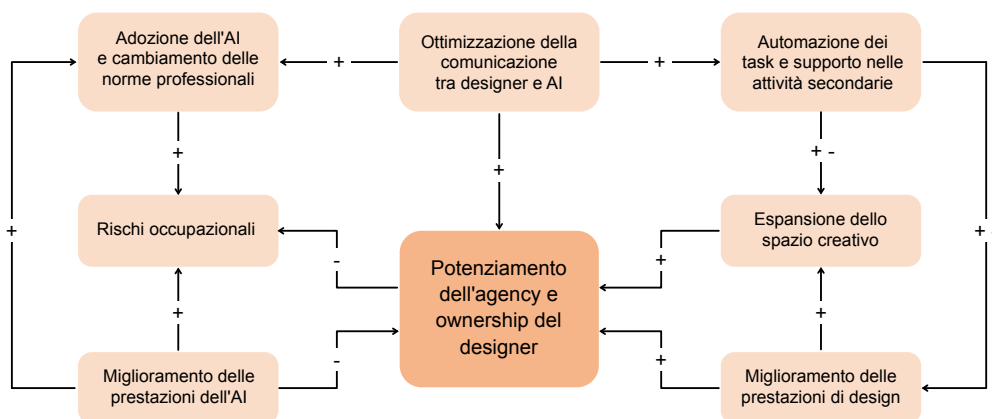
Figura 5.17: A sinistra: l'espansione e la sfocatura dei confini delle competenze individuali. A destra: l'accelerazione dei flussi di scambio e comunicazione tra domini conoscitivi differenti²⁵⁸. Grafico rielaborato.



L'IA trasforma anche la dimensione collettiva del lavoro. I confini tra discipline, tradizionalmente definiti e stabili, diventano fluidi e sovrapposti: ingegneri, designer, sviluppatori e ricercatori si trovano a condividere spazi di lavoro e competenze in modo crescente, con scambi circolari tra domini che prima comunicavano in modo sequenziale. Questa porosità non è priva di tensioni: ridefinisce responsabilità, ruoli e gerarchie, ma apre anche la possibilità di processi più integrati e reattivi²⁵⁸.

Come evidenziato nello schema analizzato, il fulcro di questa trasformazione è rappresentato dall'enhanced designer agency and ownership, ovvero dall'aumento dell'autonomia decisionale e del controllo del progettista sul processo creativo. Grazie al miglioramento delle prestazioni dei sistemi di IA e all'ottimizzazione della comunicazione uomo-macchina, il designer non è più chiamato a realizzare direttamente ogni singolo output, ma a definire obiettivi, strategie, vincoli e criteri qualitativi che guidano il lavoro dell'algoritmo. L'automazione delle attività secondarie e ripetitive consente infatti di liberare tempo e risorse cognitive, favorendo l'esplorazione

Figura 5.18: Mappatura tematica delle prospettive e delle esperienze dei designer industriali nell'adozione della GenAI²⁵¹. Grafico rielaborato.



5.2.7 Tensioni irrisolte

Se da un lato l'Intelligenza Artificiale amplia le possibilità operative del designer, dall'altro introduce una serie di questioni che rimangono tuttora aperte. Le promesse di velocità, efficienza e supporto creativo si accompagnano infatti a interrogativi più profondi che non rappresentino semplici limiti temporanei destinati a scomparire con l'evoluzione degli strumenti, ma evidenziano alcune tensioni intrinseche al rapporto tra progettazione e automazione.

Il deskilling

Come spiegato in precedenza, storicamente, l'alfabetizzazione al progetto si è strutturata attraverso processi manuali e lenti: lo schizzo sul taccuino, la manipolazione volumetrica dei modelli fisici, lo scontro con i vincoli geometrici del CAD. Nel momento in cui un designer in formazione adotta l'IA come interfaccia primaria per la

generazione di concept o per la sintesi visiva, l'atto dello "scontro con il limite" viene eluso. L'algoritmo offre soluzioni formali immediate e già risolte dal punto di vista estetico. Il rischio reale è l'instaurarsi di un apprendimento superficiale, in cui il professionista junior si trasforma in un selettore di opzioni preconfezionate, privo della competenza profonda necessaria per decostruire l'output ricevuto, individuarne le debolezze strutturali o strutturare varianti autonome che non siano semplici derivazioni del prompt iniziale.

I bias non intenzionali

I modelli di intelligenza artificiale operano su base statistica e tendono perciò a proporre configurazioni oggettuali che rispecchiano le tassonomie dominanti dei paesi in cui tali tecnologie vengono sviluppate e addestrate. Nel product design, ciò si traduce in un'evidente resistenza del software a elaborare soluzioni che escano ad esempio dai canoni occidentali. Quando si richiede la generazione di scenari d'uso, di interni o di interfacce utente, le macchine tendono a escludere le subculture visive, le tradizioni artigianali locali e le segmentazioni demografiche non standardizzate. Il designer che non esercita un'azione di disturbo o di ri-orientamento del software finisce per agire come veicolo di diffusione di un immaginario visivo monoculturale, normalizzando specifiche configurazioni estetiche ed escludendo la complessità del reale dall'orizzonte del progetto.

L'omologazione estetica

La disponibilità diffusa delle medesime architetture neurali rischia di generare una saturazione del mercato basata sulla ridondanza formale. Poiché i modelli generativi estraggono regolarità matematiche da ciò che è già stato ampiamente catalogato e immesso sul web, la loro spinta propulsiva è intrinsecamente conservativa. L'apparente ricchezza visiva degli output spesso iperrealistici maschera in realtà il riciclo di archetipi esistenti. Questo fenomeno tocca da vicino la progettazione strategica: se l'ispirazione e la prima visualizzazione di un semilavorato dipendono da piattaforme condivise da migliaia di studi concorrenti, la capacità di un marchio di distinguersi si riduce drasticamente. Si assiste a una convergenza verso soluzioni di compromesso statistico. La sfida per la disciplina non risiede più nel saper produrre l'immagine "bella" o d'impatto, ma nel possedere gli strumenti storici e critici per scardinare la piacevolezza superficiale dell'output algoritmico, ricercando attivamente

te quella rottura linguistica che storicamente definisce l'innovazione disruptive.

L'atrofizzazione cognitiva

Il pericolo più insidioso per l'autonomia intellettuale del progettista risiede nella progressiva esternalizzazione dell'atto decisionale. Un processo di co-design uomo-macchina non coordinato può provocare uno scivolamento verso l'indebolimento del pensiero abduttivo, ovvero la capacità umana di formulare ipotesi inedite partendo da indizi frammentari. Il rischio di non-innovazione è una conseguenza diretta: la creatività richiede la rottura con il preesistente, la capacità di formulare problemi che nessun dataset ha mai contemplato²⁵². La ricerca evidenzia che l'IA è effettivamente efficace nel generare variazioni e nelle fasi divergenti, ma la fase di convergenza, quella in cui il progettista seleziona, giudica e decide, deve restare il dominio privilegiato dell'intelligenza umana²⁴¹.

La disuguaglianza nell'accesso

Un'ultima tensione irrisolta riguarda l'accesso agli strumenti. Le infrastrutture di calcolo ad alte prestazioni e i software dotati di filtri avanzati per il controllo formale sono soggetti a barriere economiche d'ingresso elevate. Ciò rischia di generare una frattura competitiva: da un lato, i grandi studi in grado di permettersi strumenti proprietari e ad alta fedeltà; dall'altro, professionisti indipendenti o realtà emergenti confinate a interfacce commerciali standardizzate, caratterizzate da una maggiore rigidità stilistica e da una persistenza marcata dei bias descritti in precedenza.

Le tensioni irrisolte che accompagnano questa transizione (il deskilling, i bias algoritmici, l'omologazione estetica, l'atrofizzazione cognitiva, la disuguaglianza nell'accesso) non sono effetti collaterali risolvibili con soluzioni tecniche, ma sono dimensioni strutturali della nuova condizione progettuale che richiedono risposte culturali, formative e istituzionali. La risposta a queste tensioni risiede nella costruzione di un nuovo sistema di competenze che integri la fluency tecnologica con il vantaggio propriamente umano: la capacità di immaginare, giudicare e assumersi la responsabilità del senso.

Massimo Banzi

Interaction designer,
imprenditore

Massimo Banzi è il co-fondatore di Arduino, la piattaforma hardware e software open-source che ha rivoluzionato il mondo del design e della prototipazione elettronica. Considerato una figura chiave del movimento Maker a livello globale, il suo percorso unisce la competenza tecnica alla visione educativa: è stato infatti professore associato presso l'Interaction Design Institute Ivrea (IDII) e continua a insegnare presso il Copenhagen Institute of Interaction Design (CIID).

La sua carriera è guidata dalla missione di democratizzare la tecnologia, rendendo complessi sistemi computazionali strumenti accessibili e creativi per designer, artisti e non esperti. La sua esperienza può fornire uno sguardo critico sui sistemi di IA spesso centralizzati e opachi, portando una voce autorevole sulla necessità di un design che mantenga l'utente al centro del processo e sulla difesa della creatività umana di fronte all'automazione.



L'esperienza di Arduino dimostra l'importanza di progettare a partire dagli utenti. Molti sistemi di AI, al contrario, sono oggi potenti ma opachi e spesso progettati senza un reale coinvolgimento degli utenti: quanto il problema è tecnologico e quanto invece è legato al modo in cui li progettiamo e li raccontiamo?

L'esperienza di Arduino si basa proprio sull'idea di fare un'innovazione che è principalmente esperienza d'uso: dove la tecnologia magari non è la più avanzata o la più sofisticata, ma l'esperienza d'uso è il valore principale.

Spesso i sistemi AI contemporanei sono prodotti che vengono da ricercatori e ingegneri, dove costruire la piattaforma più sofisticata è la cosa più importante; chi si occupa di centrare l'innovazione sulle persone arriva più avanti nel processo di realizzazione del prodotto.

È chiaro che viviamo in un momento in cui intorno alle AI si è creata una corsa all'oro che ricorda tutte le bolle speculative del passato. Quando c'è la promessa di guadagni enormi, si comincia a fare un sacco di scorciatoie: oltre agli utenti, gli impatti che queste tecnologie possono avere sulla vita delle persone non sono sempre considerati. I legislatori europei provano a starci dietro, ma come diceva giustamente un professore dell'MIT qualche giorno fa, qualunque grande azienda di AI americana è meno regolamentata di un qualsiasi negozio di panini: chi vuole aprire un negozio di panini deve seguire molte più regole e stare molto più attento alle norme di chiunque voglia costruire un modello di AI avanzato.

È ovvio che quando sei in una bolla, e ci sono tanti soldi in gioco, l'impatto sulla vita delle persone diventa davvero una nota a piè di pagina. Sicuramente è un momento difficile per chi tiene alla privacy delle persone.

L'approccio di Arduino ha reso la tecnologia accessibile e aperta a designer, maker e non esperti. Oggi, nel campo dell'AI, prevale invece un modello centralizzato e proprietario: quali ostacoli impediscono una democratizzazione simile? Vede la possibilità di un'evoluzione verso modelli più aperti e accessibili, o lo sviluppo dell'AI resterà legato a grandi piattaforme private?

L'AI che vediamo tutti è principalmente fatta di grossi modelli, costruiti da grandissime aziende, che girano in cloud e per cui dobbiamo pagare per utilizzarli. In realtà esistono molti modelli di AI anche molto sofisticati, persino delle stesse aziende che fanno questi prodotti grossi, rilasciati come open source, proprio come Arduino.

Se mi doto di un computer con una scheda madre abbastanza potente — ma basta davvero un Mac delle ultime generazioni, che è pieno di acceleratori AI — posso scaricare gratuitamente dei modelli e farli girare sul mio computer, senza bisogno del cloud. Certo, un modello di frontiera magari ha 750 miliardi di parametri, mentre quello che scarico sul mio computer ne ha 14: non è il massimo, però per tantissime applicazioni può essere molto utile.

Un esempio: ho usato uno di quei modelli grossi in cloud, gli ho mostrato una serie di contenuti che mi interessavano e altri che non mi interessavano, e gli ho chiesto di sviluppare un prompt che fungesse da filtro. Poi ho dato questo filtro a un modello Mistral — un'azienda francese, open source, 14 miliardi di parametri — e Mistral, con il prompt generato dal modello più grande, è in grado di guardare degli articoli online e dirmi con un certo livello di precisione se mi interessano.

Ci sono quindi molte cose che ci portano a immaginare uno spazio, oggi, per i modelli e gli strumenti open source. Esiste

tanto software open source anche dal lato utente: per esempio Open WebUI, e molti altri che riproducono l'interfaccia dei grandi modelli, ma dietro hanno un modello open source. La possibilità c'è, si può fare; è chiaro però che queste cose non vengono molto viste da chi non è un esperto. Forse il lavoro che potrebbe fare qualcuno che lanciasse un "Arduino" di questo campo è costruire uno strumento che metta insieme tutto e lo renda facilissimo, utilizzabile anche dall'utente finale che non è esperto di AI.

Da professore, ha sempre incoraggiato i giovani designer a mantenere l'ottimismo. Come si mantiene questo spirito propositivo di fronte alla prospettiva che l'IA possa automatizzare parte del lavoro progettuale? E cosa direbbe a chi teme una perdita di spazio per la creatività umana?

Devo dire che già adesso si sta creando un movimento interessante: la gente guarda le cose fatte con le AI, comincia a riconoscerle e comincia a trovarle sgradevoli, il che è molto sano. Quindi un minimo di speranza c'è. È chiaro che per molte applicazioni dove c'era lavoro molto manuale, prima o poi l'AI probabilmente porterà via una serie di lavori, quelli di più bassa manovalanza, per dirla in maniera un po' grezza. Però rimane sempre uno spazio. Lo diceva perfino Ada Lovelace quando, nell'Ottocento, guardando i progetti della macchina analitica di Babbage si rese conto che un sistema basato sulla conoscenza di un essere umano può potenzialmente fare solo ciò che era già possibile fare a un essere umano.

Se noi designer accettiamo l'idea che l'AI diventa uno strumento, la nostra individualità — da dove veniamo, la nostra cultura, il mix di culture ed esperienze che abbiamo fatto nella vita — ci dà la possibilità di vedere le cose da un angolo di-

verso. Utilizzando l'AI possiamo quindi produrre lavori molto interessanti che l'AI da sola non riuscirebbe a creare: lì lo spazio dell'essere umano c'è ancora.

Alla fine muore la parte più manuale del design, mentre la parte creativa resta. Se avete visto, ci sono degli spot girati di recente — uno della Apple — dove hanno pubblicato online il dietro le quinte per far vedere che era tutto fatto a mano. Le cose fatte a mano diventeranno l'equivalente dell'alta moda: ci saranno le schiuffezze fatte con l'AI che piaceranno a certi clienti, e poi altri clienti disposti a volere un prodotto garantito, fatto a mano.

Purtroppo ogni generazione che introduce automazioni nella tecnologia fa perdere dei ruoli, però in qualche modo se ne creano altri. Conosco giovani designer che hanno abbracciato l'AI in maniera molto sofisticata e che, usando tool che combinano diversi modelli tra loro, fanno cose molto avanzate. Il tema è diventare quelle persone che riescono a mettere insieme tutti questi modelli e farne un cocktail che produce qualcosa di innovativo, dove l'elemento umano è quello che comprende l'elemento umano.

Faccio una risposta corollario. Io, per esempio, uso l'AI per programmare; ma programmo da una vita, e ogni tanto l'AI mi dice delle castronerie pazzesche. Per poter sfruttare al massimo questa tecnologia devo sapere qual è il risultato che deve arrivare alla fine ed essere in grado di capire se quel risultato è buono o fa schifo. Secondo me succederà lo stesso con il designer: tutti potranno chiedere a un modello "fammi l'invito di compleanno della zia Pina", ma solo un bravo designer saprà guardare l'output e dire "ok, questa è una schifezza, facciamo una cosa diversa". L'elemento umano sarà sempre, finché dura, quel pezzo in più che distingue una cosa veramente creativa e innovativa da uno stampino.

5.3 Nuovo sistema di competenze

La trasformazione dei processi richiede un aggiornamento profondo del set di abilità del designer. Accanto alla padronanza tecnica degli strumenti, emergono nuove competenze trasversali legate alla gestione dell'informazione, alla critica degli output e alla capacità di tradurre intenti complessi in istruzioni algoritmiche.

5.3.1 La formula del designer AI-native

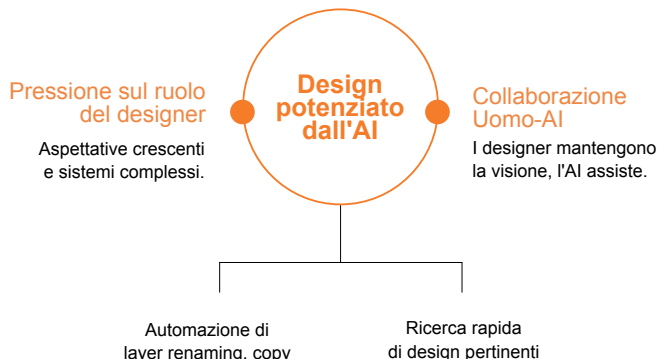
Se le attività operative tendono progressivamente a essere automatizzate e il valore del progettista si sposta verso funzioni di indirizzo, valutazione e coordinamento, diventa necessario interrogarsi su quali capacità consentano di operare efficacemente in questo nuovo scenario. Il concetto di designer AI-native nasce proprio da questa esigenza: descrivere un profilo professionale capace di integrare competenze tecnologiche, comprensione critica dei sistemi intelligenti e qualità distintivamente umane all'interno di un'unica identità progettuale.

Dai nuovi strumenti alla nuova identità professionale

Il passaggio dall'era pre-algoritmica a quella post-generativa non si esaurisce nell'adozione di nuovi software. Come si è visto nella sezione precedente, l'intelligenza

²⁶⁰ Dhanwani, R., *Designing the Future with AI: Trends & Insights for 2026*, Parallel, 2026.

Figura 5.19: Dinamiche di transizione verso l'AI-Powered Design: bilanciamento tra l'aumento della complessità sistemica e i benefici della collaborazione integrata per l'automazione dei compiti di routine e la ricerca rapida di soluzioni²¹¹. Grafico rielaborato.

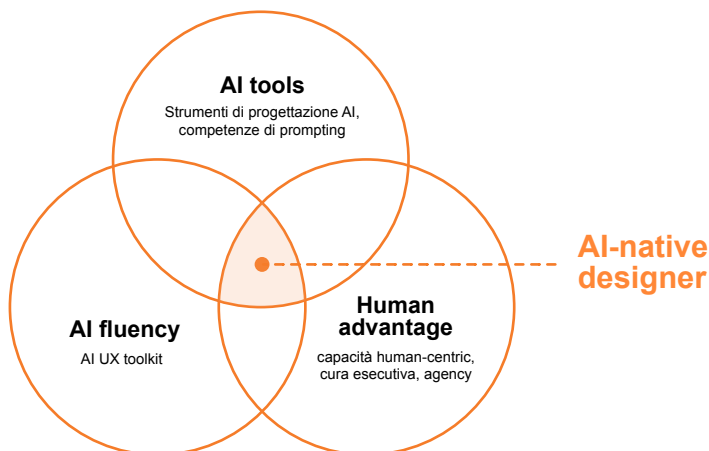


Questa transizione implica che la competenza del designer non possa più essere misurata soltanto in termini di abilità operative, ma debba includere la capacità di articolare visioni progettuali in forme comprensibili dall'intelligenza artificiale e, al tempo stesso, di valutare criticamente le risposte che il sistema fornisce. Il denominatore comune che emerge dall'analisi delle strategie aziendali, dai percorsi formativi internazionali e dalla letteratura di ricerca è uno spostamento del valore professionale dalla competenza tecnica operativa alla capacità di istruire, orientare e governare sistemi intelligenti²⁵⁹.

La formula AI tools + AI fluency + human advantage = AI-native designer, elaborata nel dibattito contemporaneo sul design digitale, offre una cornice interpretativa efficace per comprendere questa evoluzione²⁶¹.

²⁶¹ Sharma, S., *AI tools + AI fluency + human advantage = AI-native designer*, UX Collective, 2025.

Figura 5.20: Le macro-competenze per definire il profilo dell'AI-native designer²⁶¹. Grafico rielaborato.



Il primo termine, *AI tools*, designa la padronanza pratica degli strumenti: conoscere le piattaforme di generazione di immagini, i software di modellazione assistita, i copiloti ingegneristici, le interfacce conversazionali per la prototipazione. Il secondo termine, *AI fluency*, va oltre l'uso operativo e indica la comprensione dei principi di funzionamento dei modelli generativi, dei loro limiti strutturali, dei bias incorporati nei dati di addestramento e delle implicazioni etiche dell'output algoritmico. Il terzo termine, *human advantage*, identifica quelle qualità cognitive, relazionali ed etiche che restano irriducibilmente umane e che, proprio nell'era dell'automazione cognitiva, acquisiscono un valore strategico senza precedenti²⁶⁰. La somma di queste tre dimensioni non produce un professionista semplicemente "aggiornato", ma una figura qualitativamente nuova: il designer *AI-native*, capace di operare all'interno di ecosistemi ibridi in cui intelligenza umana e intelligenza artificiale si integrano in modo strutturale²⁴⁹.

5.3.2 Le competenze umane che faranno la differenza

Quando la produzione formale ed esecutiva viene parzialmente delegata ai sistemi automatizzati, l'identità del progettista smette di coincidere con la pura maestria tecnica e si rifonda su facoltà puramente umane. In questo nuovo scenario professionale, le abilità che fanno la differenza non si misurano più sulla padronanza del singolo software, ma diventano meta-competenze trasversali capaci di tradurre l'efficienza degli algoritmi in senso culturale e valore sociale.

Le previsioni sulle competenze fondamentali nel 2030

L'orizzonte professionale del 2030 delinea una radicale inversione di tendenza nei requisiti di accesso al mercato del lavoro, spostando il baricentro dalle abilità puramente esecutive verso facoltà di matrice strategica, relazionale e cognitiva.

Il World Economic Forum evidenzia infatti che le competenze più richieste nel 2030 non saranno più di natura tecnica. Al contrario, con l'avanzare dell'automazione nel mondo del lavoro, i datori di lavoro daranno la priorità a capacità strategiche e incentrate sulle persone, quali il pensiero analitico, il pensiero creativo, l'alfabetizzazione tecnologica e la resilienza²⁶².

Il report McKinsey presenta grafici che documentano l'introduzione dell'intelligenza artificiale nelle professioni

²⁶² World Economic Forum, *The Future of Jobs Report 2025*, Ginevra, 2025.

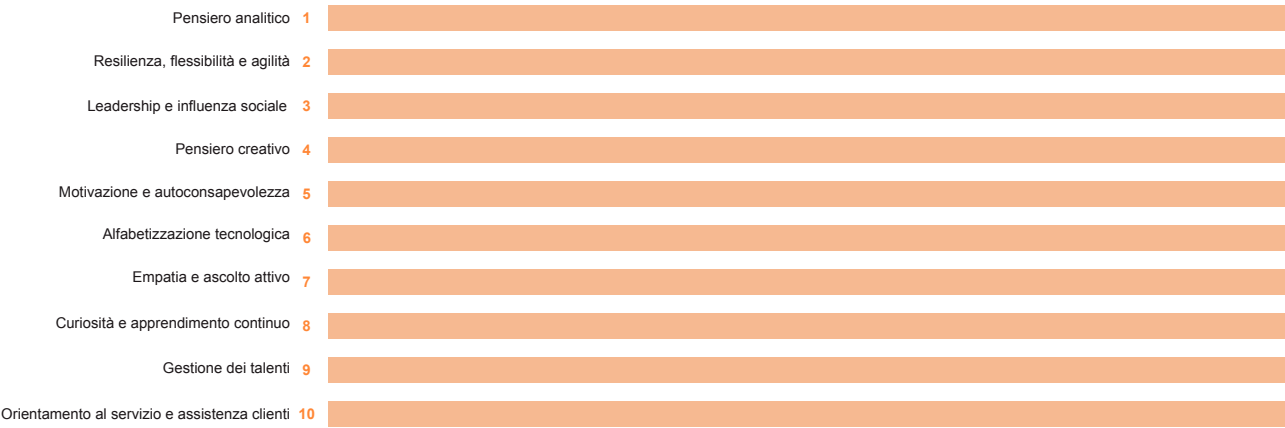


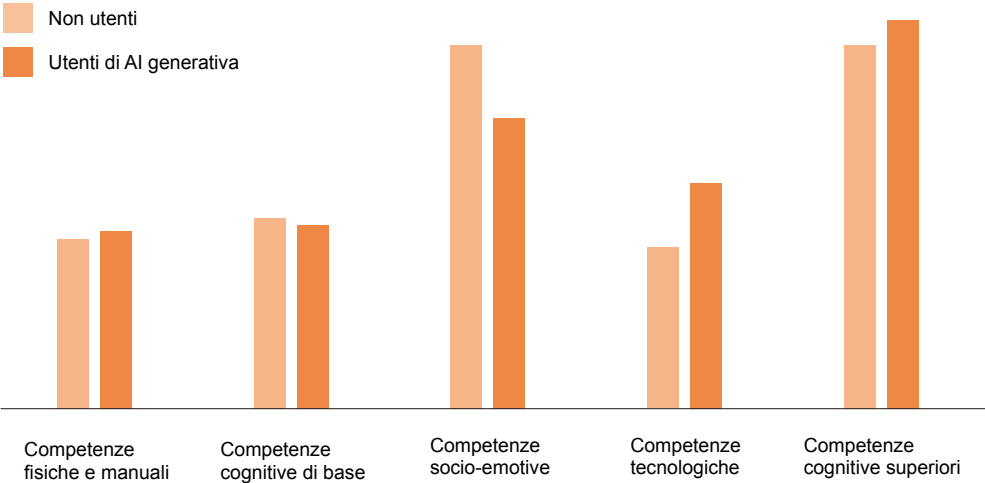
Figura 5.21: Graduatoria delle prime 10 competenze fondamentali richieste dal mercato, secondo il World Economic Forum²⁶¹. Grafico rielaborato.

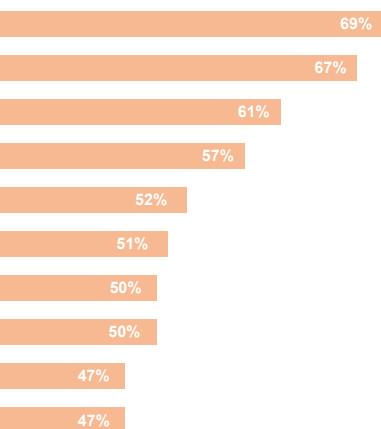
Figura 5.22: Percezione dell'importanza delle competenze lavorative a confronto tra utenti e non utenti di IA generativa. Rielaborazione da dati McKinsey & Company²⁶². Grafico rielaborato.

in generale e, soprattutto, l'importanza crescente delle competenze umane nel panorama lavorativo proiettato verso il 2030. Le visualizzazioni mostrano come le competenze tecnologiche di base, pur necessarie, saranno progressivamente commoditizzate dall'automazione, mentre le competenze sociali, emotive e cognitive superiori (ragionamento analitico, leadership, influenza sociale, pensiero critico e creatività) acquisiranno un peso percentuale crescente nella domanda di mercato. In particolare, i grafici McKinsey articolano le competenze in cinque macro-categorie: competenze fisiche e manuali, competenze cognitive di base, competenze cognitive superiori, competenze sociali ed emotive e competenze

Le competenze cognitive e socio-emotive sono più importanti di quelle tecnologiche in diverse categorie, sia per gli utenti che per i non utenti di AI generativa.

Dipendenti che hanno indicato ciascuna competenza come 1 delle loro 2 principali su 5 presentate, in termini di importanza per svolgere efficacemente il proprio lavoro attuale, %:





²⁶³ *The Human Side of Generative AI: Creating a Path to Productivity*, McKinsey & Company, 2024.

²⁶⁴ *The Intersection of Design Thinking and AI: Enhancing Innovation*, IDEO U, 2025.

Figura 5.23: Integrazione delle capacità umane e artificiali nei processi creativi. L'area di intersezione centrale rappresenta la finestra di opportunità in cui la combinazione dei due distinti set di valori amplifica il potenziale innovativo complessivo²⁶⁴. Grafico rielaborato.

tecnologiche. La proiezione al 2030 indica un decremento significativo della domanda nelle prime due categorie e un incremento marcato nelle ultime tre, con le competenze sociali ed emotive che registrano il tasso di crescita più elevato²⁶³.

Le domande che solo un umano può farsi

Se l'intelligenza artificiale eccelle nel fornire risposte rapide a problemi ben definiti, la specificità dell'intelligenza umana risiede nella capacità di formulare le domande giuste prima ancora di cercare risposte²⁴⁹. Questa capacità interrogativa costituisce il nucleo irriducibile della competenza progettuale nell'era algoritmica.

Il framework elaborato da IDEO rappresenta visivamente questa complementarità: da un lato i valori umani (empatia, creatività, intuizione, autoconsapevolezza) dall'altro i valori dell'intelligenza artificiale (velocità di elaborazione, consistenza, ampiezza di prospettive, multitasking). I due insiemi non si sovrappongono per caso: è precisamente nella loro intersezione che si genera l'innovazione amplificata dall'AI ed è il designer il soggetto che sa quando e come combinare le due parti²⁶⁴.

Innovazione potenziata dall'AI



Per il nuovo designer, la capacità di interrogarsi eticamente sugli output generati assume una rilevanza operativa concreta. Le nuove responsabilità che emergono dalla progettazione con l'intelligenza artificiale comprendono non soltanto la comprensione degli utenti, ma anche la riflessione approfondita sulla futura collaborazione AI-umano che si desidera creare e sul futuro che si vuole vedere realizzato. Quando si definiscono i requisiti

del progetto, occorre contemporaneamente definire le capacità dell'intelligenza artificiale che si intende sfruttare e comprendere se siano sufficientemente mature per essere impiegate. I sistemi di intelligenza artificiale ereditano i bias dai dati di addestramento, e tali bias possono emergere nel prodotto in modi che danneggiano specifici utenti. Il designer non può sempre controllare i dati di addestramento, ma può condurre ricerca inclusiva, portare voci sottorappresentate nel processo di design, comunicare trasparentemente i potenziali bias e promuovere linee guida etiche all'interno della propria organizzazione²⁶⁰.

Il pensiero sistemico

Tra le competenze umane destinate ad acquisire centralità crescente, il pensiero sistemico occupa una posizione privilegiata. Nell'era dell'intelligenza artificiale, il designer non opera più in isolamento su un singolo artefatto, ma all'interno di ecosistemi digitali che connettono concept, simulazione, produzione e distribuzione. I digital twin non sono più strumenti di validazione a posteriori, ma ambienti in cui il progetto nasce, evolve e viene continuamente ottimizzato, mentre gli agenti AI orchestrano flussi di lavoro complessi, validano automaticamente la conformità normativa e anticipano problemi di fabbricabilità. In questo scenario, il designer assume il ruolo di system integrator: colui che definisce gli obiettivi strategici, stabilisce i parametri di valore, supervisiona il lavoro degli agenti e garantisce la coerenza complessiva del sistema. Come suggerisce il report Capgemini TechnoVision 2026, si passa dal "creare strumenti" all'"esprimere intenti"²⁴⁹.

5.3.3 Nuove competenze da acquisire

Le competenze che questo scenario richiede si articolano su tre livelli distinti. Il primo è operativo e risponde alla domanda più immediata: come si usa l'intelligenza artificiale? Il secondo è relazionale e critico, e riguarda una questione più sottile: come ci si sta dentro al rapporto con il sistema, con quale postura, con quale grado di fiducia e di distanza critica? Il terzo è cognitivo e creativo, e tocca la dimensione più profonda: chi è il designer che governa tutto questo, quali capacità lo rendono irriducibile all'automazione?

Competenze operative

AI literacy e comprensione dei modelli

L'alfabetizzazione all'IA non è una competenza accessoria, ma un requisito trasversale che deve permeare l'intero percorso formativo del designer²⁴⁹. Non si tratta di conoscere le interfacce o le funzionalità degli strumenti, ma di comprendere i principi di funzionamento dei modelli generativi, i loro limiti intrinseci e le implicazioni dei dati di addestramento sugli output.

La differenza sostanziale rispetto agli altri strumenti digitali, come il CAD ad esempio, risiede nella natura probabilistica: le piattaforme AI-driven introducono un elemento di serendipità che trasforma il processo di design in un dialogo dinamico, più simile al co-design con un collaboratore che all'uso di uno strumento²⁴². Comprendere questa differenza è il punto di partenza di ogni competenza operativa successiva.

Il prompt engineering

La pratica del prompting richiede una comprensione profonda delle capacità del modello e delle strategie per orientarne l'output verso risultati coerenti con l'intento progettuale. È una pratica formale perché si concentra sulla composizione della richiesta più efficiente, e iterativa perché richiede di attraversare gli stessi prompt più volte fino al raggiungimento del risultato desiderato. La selezione delle parole, la struttura sintattica e il contesto della richiesta devono allinearsi con precisione alla visione progettuale²⁴².

Il prompt engineering introduce inoltre una dimensione cognitiva inedita per il designer: la necessità di navigare simultaneamente tra il dominio visuo-spaziale e quello linguistico. I creativi visivi, abituati all'immediatezza dell'espressione visiva, devono tradurre idee, composizioni ed emozioni in descrizioni testuali coerenti. Un processo di doppia cognizione che rappresenta una sfida peculiare per chi ha sempre operato prevalentemente per immagini²⁴².

L'esperienza di ricerca condotta presso il Politecnico di Milano con strumenti come Midjourney ha mostrato come questa competenza sia tutt'altro che generica: la strutturazione dei prompt per il design industriale richiede di impiegare il brief progettuale come prompt preliminare, di incorporare i requisiti come attributi specifici e di selezionare la terminologia attraverso benchmarking e analisi delle tendenze. [Generative AI in Product Design] Il legame tra conoscenza disciplinare e qualità del prompting è indissolubile: senza basi solide in storia del

design, terminologia tecnica e linguaggio progettuale, la capacità di interagire efficacemente con i sistemi generativi rimane superficiale²⁵².

Istruire lo strumento

Per gran parte della storia del design digitale, la competenza professionale si è misurata sulla padronanza tecnica degli strumenti: sapere dove si trovava ogni comando, conoscere le scorciatoie da tastiera, dominare i flussi di lavoro specifici di ogni software. L'intelligenza artificiale rompe questa dinamica e tale trasformazione è visibile nei nuovi paradigmi degli strumenti professionali.

Istruire significa innanzitutto saper costruire contesto. A differenza di un software tradizionale, che opera sempre nello stesso modo indipendentemente da chi lo usa, un modello generativo produce risultati radicalmente diversi in funzione del contesto che riceve. La ricerca empirica con designer industriali documenta come l'interazione risulti faticosa proprio quando manca questa consapevolezza: il sistema ha bisogno di sapere chi è il suo interlocutore, quale output si vuole produrre e per quale progetto specifico²⁵¹.

Istruire significa inoltre saper definire i vincoli prima ancora dei risultati. Il valore del designer non sta nel generare la soluzione, ma nel delimitare lo spazio delle soluzioni accettabili: materiali, proporzioni, registro estetico, vincoli normativi, coerenza con il sistema di prodotto esistente. Tradurre tutto questo in istruzioni comprensibili per un sistema che non ha accesso alla storia del progetto né alle preferenze tacite del committente è una competenza che richiede tanto rigore analitico quanto capacità di astrazione²⁴⁹.

Istruire significa infine saper valutare. Un output generativo non è mai né giusto né sbagliato in assoluto: è più o meno coerente con l'intento che lo ha prodotto. Il designer legge criticamente ciò che il sistema restituisce, identifica dove l'interpretazione si è discostata dall'intenzione e riformula le istruzioni di conseguenza. La ricerca di Autodesk documenta che questo approccio può ridurre i tempi di iterazione fino al 40%, ma il guadagno si realizza solo quando il designer è in grado di guidare il processo con sufficiente chiarezza²⁶⁰.

Competenze relazionali e critiche

Collaborazione critica e consapevole

Sviluppare una collaborazione critica significa qualcosa di più sottile che limitarsi a verificare gli output: significa costruire consapevolmente le condizioni entro cui la collaborazione può essere produttiva. La ricerca mostra che le coppie umano-AI non migliorano naturalmente attraverso la ripetizione, ma solo quando viene introdotto un intervento deliberato che orienta i partecipanti a sviluppare le idee congiuntamente, concentrandosi sul perfezionamento delle idee esistenti piuttosto che sulla generazione continua di idee nuove²⁴⁹. Il futuro del design non risiede nel prompt migliore, ma nella qualità della relazione generativa tra progettista e sistema.

Calibrare la fiducia

Un'altra competenza essenziale consiste nel saper calibrare il grado di fiducia accordato al sistema in funzione del contesto. Si tratta di sviluppare un giudizio situato e sapere quando un output può essere accettato con relativa sicurezza e quando richiede verifica o riformulazione della richiesta. La fiducia mal calibrata ha due forme di errore simmetriche. Accordarne troppa espone all'*automation bias*: accettare acriticamente output plausibili ma errati, delegare decisioni che richiedono giudizio umano. Accordarne troppo poca produce un blocco operativo: il designer che non riesce a fidarsi di nessun output non riesce nemmeno a sfruttare le reali capacità del sistema.

Competenze cognitive e creative

A queste competenze operative e relazionali, si sommano le capacità cognitive e creative che definiscono il profilo del designer nel lungo periodo.

La Competency Wheel propone che i designer debbano sviluppare le seguenti sette capacità: percezione sottile, immaginazione infinita, visione acuta, decision-making lungimirante, capacità di pianificazione sistemica, competenza computazionale ibrida e creatività a catena completa²⁴⁹.

Percezione sottile (cogliere i bisogni degli utenti)

La capacità di cogliere i bisogni degli utenti è l'espressione della sensibilità e dell'empatia del designer. Si articola su tre livelli progressivi: dalla comprensione dei bisogni espliciti, all'esplorazione di quelli impliciti, fino alla creazione di nuovi bisogni capaci di orientare gli stili di vita. È la dimensione dell'empatia profonda, che l'IA può al più

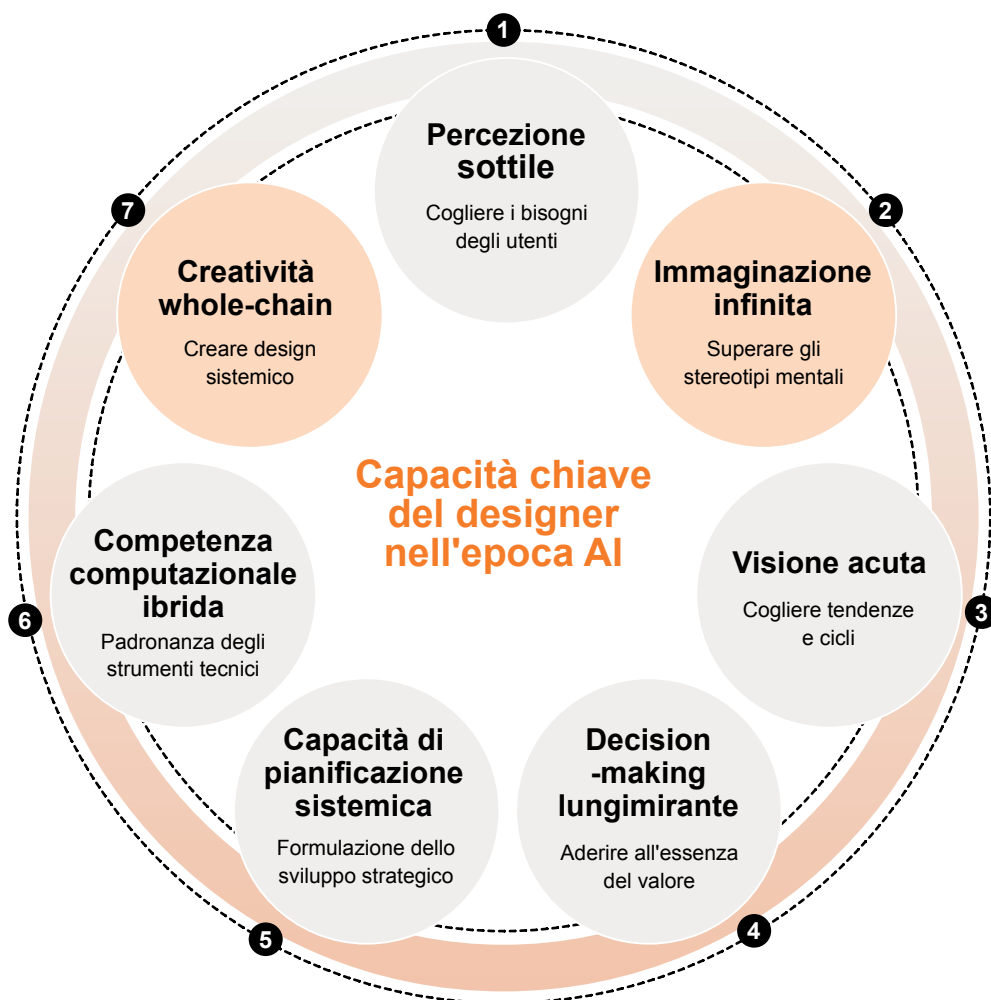


Figura 5.24: La ruota delle competenze per l'alfabetizzazione del futuro del design. Il framework illustra le sette capacità chiave necessarie per navigare la transizione professionale.²⁴⁹ Grafico rielaborato.

supportare ma non sostituire²⁴⁹.

Immaginazione infinita (superare gli stereotipi mentali)

L'immaginazione è la capacità di concepire oggetti e situazioni che non esistono, e ogni designer dovrebbe coltivare diversi tipi di immaginazione per potenziare le proprie competenze e approfondire il design. L'immaginazione infinita richiede che i designer sappiano utilizzare il pensiero controfattuale per superare gli stereotipi cognitivi. Il pensiero controfattuale è definito come la capacità di immaginare che l'esito delle cose possa essere diverso dalla situazione reale, immaginando una realtà alternativa. Poiché le persone non possono dare per scontato che il futuro sia uguale al presente, molte innovazioni nascono dalla rottura delle "vecchie cose" e dalla "costruzione" di "cose nuove"²⁴⁹.

Visione acuta (cogliere tendenze e cicli)

La visione acuta è la capacità di distinguere tra le infinite possibilità immaginabili quelle coerenti con le tendenze e i cicli tecnologici, ecologici, commerciali, normativi, per creare soluzioni progettuali con significato positivo e in linea con le esigenze dello sviluppo innovativo. In questo, l'IA può fornire ai designer risorse di dati e basi analitiche, ma come combinare efficacemente i due elementi con i bisogni reali degli utenti richiede il giudizio autonomo del designer²⁴⁹.

Decision-making lungimirante (aderire all'essenza del valore)

Il potere decisionale è la capacità che i designer nell'era dell'IA stanno già ampliando notevolmente. Per quanto avanzata possa diventare l'IA, il suo design non potrà mai sostituire il processo decisionale umano: il potere decisionale è un diritto che gli esseri umani devono preservare. Da questa prospettiva, i designer del futuro assomiglieranno maggiormente a dei design manager, con diversi strumenti che svolgono il ruolo di designer per ciascuna posizione specifica. Lo standard per il processo decisionale richiede ai designer di "aderire all'essenza del valore" che deve essere costantemente considerata nei prodotti e nei servizi²⁴⁹.

Capacità di pianificazione sistemica (formulazione dello sviluppo strategico)

I designer dovranno concentrarsi sulla strategia e sulla pianificazione, acquisire una comprensione approfondita del posizionamento e degli obiettivi del prodotto e del brand, formulare strategie e piani di design più efficaci e creare maggior valore per le imprese. L'AIGC è solo uno strumento di design e non può sostituire la formulazione e l'esecuzione della strategia e della pianificazione progettuale da parte del designer²⁴⁹.

Competenza computazionale ibrida (padronanza degli strumenti tecnici)

La competenza computazionale ibrida è la capacità di padroneggiare continuamente i nuovi strumenti digitali e le tecnologie emergenti, aggiornandosi senza soluzione di continuità. Non contraddice quanto detto finora: la familiarità tecnica rimane necessaria, ma come condizione abilitante, non come fine²⁴⁹.

Creatività a catena completa (creare sistemi di design)

L'essenza del design è l'innovazione, e la capacità ultima richiesta a un designer è la creatività: dalla scoperta del problema alla creazione di soluzioni con valore significativo e impatto reale. Si tratta di un processo sistematico che va dalla strategia all'implementazione. Se l'IA può sostituire i designer nello svolgimento di attività creative di un singolo nodo, o anche di una certa fase del processo creativo, risulta però inadeguata per il lavoro a catena completa che nei contesti aziendali costruisce prodotti, servizi ed esperienze da zero²⁴⁹.

5.3.4 La biforcazione delle competenze: due profili emergenti

L'analisi condotta nei paragrafi precedenti non converge verso un unico futuro profilo professionale, ma suggerisce due traiettorie che rispondono a esigenze diverse del mercato e a vocazioni diverse del progettista. Questi due profili non sono mutuamente esclusivi, né definitivamente separati. Definiscono piuttosto due polarità entro cui ciascun designer troverà la propria posizione, in funzione del contesto disciplinare, del settore professionale e della propria inclinazione. Ciò che li accomuna è la consapevolezza che il valore del progettista non risiede più nell'esecuzione, ma nella direzione.

Il designer come direttore creativo algoritmico

Il primo profilo è quello del designer che sfrutta le capacità generative dell'intelligenza artificiale per esplorare spazi risolutivi vastissimi e produrre varianti a velocità impensabili nel paradigma pre-algoritmico. In questo ruolo, il valore professionale risiede nella capacità di visione, nell'intuizione creativa e nella direzione strategica: il designer non esegue, ma orienta il sistema verso esiti significativi, seleziona tra le possibilità generate e imprime coerenza culturale ed estetica al processo. [Internazionale dell'Intelligenza Artificiale nel Design] La creatività del designer non si misura più sulla singola soluzione, ma sulla qualità del sistema generativo che ha saputo costruire²⁶⁰.

È in questo contesto che acquista pieno senso la formula coniata da Claire Silver, artista e figura di riferimento nel dibattito sulla creatività AI-collaborativa: "taste is the new skill". Quando gli strumenti abbassano la soglia tecnica fino a renderla accessibile su larga scala, ciò che di-

²⁶⁵ Rodriguez, H., *Think about what's happening here. Hundreds, if not thousands of micro-decisions*, AlxCreative, 2024.

stingue il direttore creativo algoritmico non è la capacità di eseguire ma quella di scegliere, rendendo il gusto il nucleo della competitività²⁶⁵.

Il designer come garante della qualità e dell'etica

Il secondo profilo è quello del designer che assume il ruolo di validatore critico degli output algoritmici: colui che filtra, verifica e orienta i risultati dell'intelligenza artificiale secondo criteri di fattibilità tecnica, sostenibilità ambientale, coerenza normativa e responsabilità etica. Questo ruolo richiede rigore analitico, conoscenza approfondita dei materiali e dei processi, e una sensibilità acuta verso le implicazioni sistemiche delle scelte progettuali²⁴⁹.

È una figura che non compete con l'algoritmo sul piano dell'esecuzione, ma lo governa sul piano dei valori e che trova nella propria formazione umanistica e progettuale la fonte di un'autorità che nessun sistema automatizzato può replicare.

5.3.5 Da dove inizia il cambiamento?

Il designer che emerge da questa analisi è una figura qualitativamente nuova. Non più soltanto un esecutore della forma, ma un architetto di relazioni tra intelligenza umana e artificiale: un professionista che formula le domande prima ancora di cercare risposte, che pensa per sistemi anziché per artefatti isolati, che assume la responsabilità non solo di ciò che produce ma delle visioni del mondo che i suoi progetti incorporano e trasmettono²⁴⁹.

Questa trasformazione pone una domanda inevitabile: come si forma questo designer? Se le competenze richieste sono di natura diversa rispetto al passato allora anche i luoghi, i metodi e le logiche della formazione devono cambiare. La sfida che attende la disciplina non è quella di sopravvivere alla trasformazione algoritmica, ma di guidarla e guidarla richiede che i sistemi educativi sappiano preparare progettisti consapevoli²⁶⁰.

Don Norman

Ingegnere informatico,
dottorato in psicologia

Considerato il padre della user experience (UX), Don Norman rappresenta il ponte ideale tra la psicologia cognitiva e l'ingegneria del design. Proprio grazie alla sua esperienza come Vice President of Advanced Technology presso Apple negli anni '90 e come professore emerito di scienze cognitive presso la University of California (UCSD) e successivamente alla Northwestern University, ha potuto definire i pilastri dell'interazione uomo-macchina che utilizziamo ancora oggi.

Partendo da questo solido background tecnico e psicologico, la sua visione si è evoluta verso un approccio che supera l'oggetto fisico per abbracciare i sistemi complessi e la loro sostenibilità sociale. La sua presenza è dunque fondamentale per discutere come la nozione di UX debba trasformarsi in un'epoca dominata da tecnologie adattive, predittive e spesso poco trasparenti.



Lei ha contribuito a dare forma al design per come lo conosciamo, ma oggi la pratica del designer sta cambiando rapidamente a causa dell'intelligenza artificiale. Cosa dovrebbero sapere i designer, adesso, su questa trasformazione, e che impatto avrà sul lavoro e sulle competenze costruite in anni di pratica?

La cosa più importante da capire sul lavoro che si fa oggi è che siamo di fronte a qualcosa di nuovo. Se guardiamo alla storia, ogni nuova tecnologia impiega cinque o dieci anni dalla sua introduzione prima di essere davvero compresa e prima che le persone capiscano come sfruttarla nel modo più efficace. Ed è dirompente, assolutamente. Questa, a mio avviso, è la trasformazione tecnologica più importante da chissà quando. A volte dico: dall'invenzione della scrittura. È potente. Le tre grandi rivoluzioni che hanno cambiato profondamente il mondo sono la scrittura, la rivoluzione industriale e questa. E non si tratta di uno strumento: è qualcosa con cui si può collaborare. Ed è proprio questo che molti di noi stanno cercando di garantire — che rimanga uno strumento di collaborazione.

La preoccupazione che i designer esprimono più spesso è: perderò il lavoro? In realtà, se ne dovrebbero preoccupare maggiormente i programmatori. Ma la risposta onesta è: se continuate a lavorare esattamente come fate oggi, sì, rischiate di perderlo. O, per chi è ancora studente, di non trovarlo. Questo significa che bisogna già adesso iniziare a imparare, usare e sperimentare con l'IA. E oggi la componente più potente non sono tanto i modelli linguistici in sé, quanto gli agenti. Gli agenti sono ancora più recenti, e sono loro che portano a termine il lavoro concreto. Cambieranno tutto.

Il mio partner Jakob Nielsen ora scrive

newsletter interamente dedicate all'IA, e ve le consiglio: ha idee molto interessanti su come i designer dovrebbero approcciarsi a questi strumenti. L'argomento centrale è questo. Nella programmazione, le competenze tecniche vengono rapidamente rimpiazzate dai sistemi di IA. Alcune delle principali aziende del settore dichiarano apertamente che i loro dipendenti non scrivono più codice: usano Claude — l'IA di Anthropic — per farlo. Ma hanno ancora bisogno dei programmatori per gestire e supervisionare Claude. E ora la stessa cosa sta accadendo nel design: le competenze che avete impiegato anni ad acquisire — imparare a disegnare, conoscere i materiali, padroneggiare i processi — non sono più strettamente necessarie. Ciò che diventa necessario è il giudizio.

Perché quello che sta avvenendo è una sostituzione delle competenze esecutive con il giudizio. E il giudizio è più importante e più difficile, perché pone domande come: stiamo facendo la cosa giusta? Questo richiede un'interazione continua durante il processo: intervenire, fermare, dire "no, rifallo in questo modo." E per farlo, bisogna comunque essere esperti di design — sapere cosa è possibile, cosa funziona, cosa no. Solo che non è più necessario eseguirlo da soli. Si ha a disposizione un assistente straordinariamente capace ma estremamente ingenuo, con cui bisogna imparare a lavorare. Un assistente che a volte inventa cose, quindi non ci si può fidare ciecamente di ciò che produce. Ma anche questo rientra nel giudizio: capire di cosa fidarsi, cosa verificare, quanto tempo vale la pena investire nel controllare che qualcosa sia davvero corretto o funzionante.

Non fissatevi troppo sui problemi attuali: molti di quelli di cui parliamo oggi non esisteranno più tra qualche mese, perché sono problemi già noti e già in corso di soluzione. Le risposte dovranno venire

da voi — non da me. Da chi oggi sta lavorando su questi temi.

Ecco il mio consiglio concreto per rispondere alla vostra domanda: dovete usare l'IA ogni giorno. Sperimentare. Imparare a costruire con gli agenti. E ripensare completamente il modo in cui lavorate.

A proposito di persone che vengono licenziate perché sostituite dall'intelligenza artificiale, vorremmo approfondire questo argomento e parlare del concetto di disumanizzazione. In un'intervista con Kevin McCullagh lei suggerisce che l'AI potrebbe rendere il lavoro più umano e meno ripetitivo. In quali condizioni progettuali questo scenario si realizza realmente, e in quali casi invece l'AI finisce per aumentare la disumanizzazione del lavoro?

La parte peggiore della rivoluzione industriale è stata la corsa all'efficienza. Abbiamo preso lavoratori specializzati — persone che progettavano abiti, cucinavano, costruivano oggetti con grande competenza e conoscenza, e che potevano essere orgogliose del risultato finale — e abbiamo detto: se dividiamo il lavoro in piccoli componenti separati, siamo molto più efficienti. L'esempio classico è la produzione di spilli: una persona sola ne produceva quaranta o cinquanta al giorno, ma dividendo il processo — chi tira il filo, chi lo taglia, chi lo affila, e così via — se ne potevano produrre migliaia. Il prezzo era che ogni persona faceva una sola cosa, tutto il giorno, all'infinito. Ed è lì che abbiamo smesso di trattare gli esseri umani come tali — e le persone lo odiavano.

Nei prossimi anni, saranno i robot a fare quello. Le nuove tecnologie di produzione e i robot sempre più intelligenti stanno già occupando gran parte di quelle

fabbriche. Quindi rispondo alla vostra domanda: credo che il giudizio umano diventerà la competenza centrale. Le macchine faranno le operazioni — montare, assemblare, eseguire. Ma saremo noi a dover capire cosa devono fare le macchine, e perché abbiamo progettato qualcosa in un certo modo. Forse esiste un metodo migliore. I veicoli elettrici, ad esempio, hanno meno componenti di un'automobile tradizionale. Capire come costruirli e valutarli — quello è il giudizio. Ed è per questo che voglio che diventiate bravi nell'usare l'IA: se la padronegiate davvero, sarete più preziosi di prima.

Dobbiamo abbattere i confini. Non parlo più solo di design, ma di design più azione — perché se non si realizza, non conta nulla. E questo sta iniziando ad accadere. Come sottolinea Jakob: oggi, quando finite il vostro design, potete chiedere all'IA di scrivere il programma di base — non sarà necessariamente quello definitivo, ma funzionerà esattamente come avete immaginato. I programmatori capiranno meglio cosa state cercando di fare, perché hanno davanti qualcosa che funziona. Potrà essere meno efficiente, ma sistamarlo è compito loro. Il punto è che i lavori si sovrappongono. Lo stesso vale per la produzione: posso fornire informazioni dettagliate sui componenti, i costi, i materiali, perché il sistema lo fa. E chi si occupa di produzione capirà meglio cosa state cercando di realizzare.

La separazione netta tra i diversi ruoli sta cambiando. Ma io sostenevo questa idea già prima dell'IA: pensavo che dovessimo avere team multidisciplinari in cui non lavorano solo designer diversi tra loro, ma anche marketing, produzione, vendite — perché ognuno ha requisiti diversi, e se non li conoscete, non progetterete bene.

Questo cambierà. Avremo un design più olistico e più collaborativo. E in parte sarà più facile, perché molti dei dettagli

che assorbivano tempo enorme li gestirà la macchina.

Quanto conta ancora, secondo lei, il lavoro individuale del designer in un contesto in cui i progetti sono sempre più frutto di team multidisciplinari?

In realtà, non ha più senso parlare del singolo designer. Ormai si parla di team.

Purtroppo, però, fin dal primo anno di scuola, si impara a fare le cose da soli e si viene valutati su quanto bene si riescano a fare da soli. All'università, chiedere aiuto non è permesso — non si può copiare, non ci si può far aiutare dagli altri. Anni fa ho scritto un articolo in cui dichiaravo di essere favorevole al "copiare". Perché l'unico posto in cui bisogna fare tutto da soli è la scuola. Mentre nel mondo reale, quello che conta è il risultato finale — non importa di chi sia l'idea.

Penso che dobbiamo formare le persone in modo diverso. Anche nell'educazione al design: una volta che si sceglie se diventare graphic designer, industrial designer, interaction designer, si parla solo con persone della propria specializzazione. Non si impara a lavorare insieme. Anzi, spesso si compete. No: bisogna imparare a crescere insieme. E non basta imparare a lavorare con altri designer — bisogna imparare a lavorare con chi si occupa di marketing, vendite, produzione. È questa l'educazione che dobbiamo costruire.

E l'IA aiuterà. Nel mondo esistono moltissime specializzazioni, e il problema è che se sei un ottimo graphic designer, probabilmente non conosci bene l'interaction design o l'industrial design — perché ogni specializzazione richiede una conoscenza così approfondita che nessuna singola persona può padroneggiarle tutte. Non è una colpa — è semplicemente impossibile. Ebbene, ora abbiamo

un'intelligenza artificiale che conosce praticamente tutto. È straordinariamente capace, ma è molto ingenua. Non sa come fare uso di ciò che sa, non ha giudizio. Quindi dobbiamo sfruttare la sua conoscenza. Quando devo approfondire un nuovo argomento — cosa che mi capita spesso, perché amo imparare cose nuove — lo spiego a Claude in dettaglio, descrivendo cosa sto cercando. È come avere un assistente brillantissimo. Ma ripeto: è un assistente molto ingenuo, che non capisce davvero le persone. Conosce la letteratura sul comportamento umano, ma non le persone specifiche per cui state progettando. Ed è lì che il giudizio umano sarà sempre insostituibile.





formazione

6.1 Un sistema in ritardo: la formazione che non tiene il passo

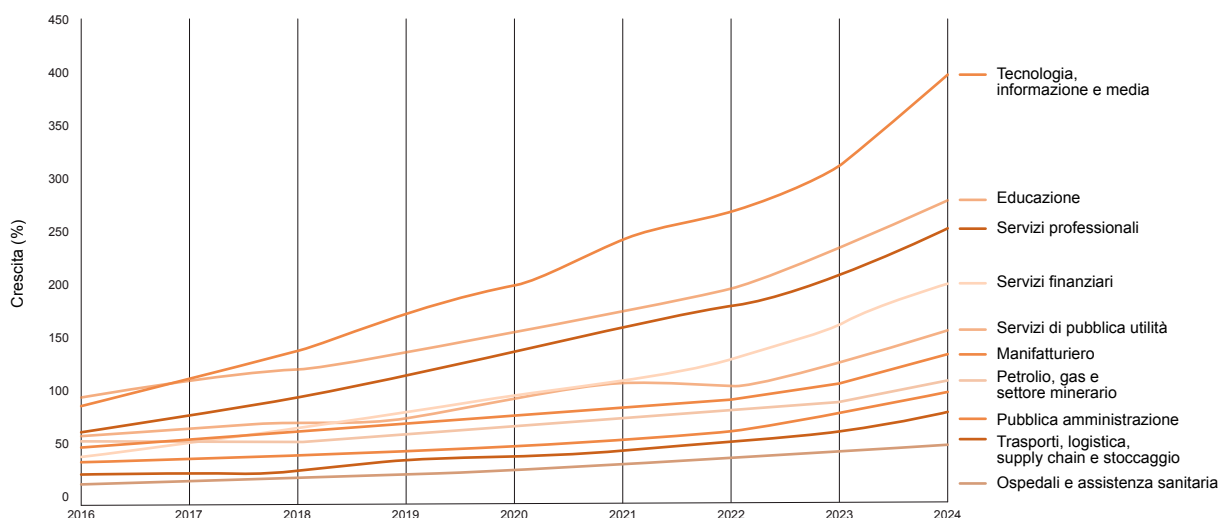
Mentre la tecnologia corre a velocità esponenziale, le istituzioni educative faticano a riformare i propri curricula con la stessa rapidità. Questo divario temporale crea una zona d'ombra in cui gli studenti si trovano spesso a utilizzare strumenti avanzati senza una guida metodologica e critica adeguata.

6.1.1 La differenza di velocità

Il divario strutturale tra la velocità dell'evoluzione degli strumenti di intelligenza artificiale e l'aggiornamento dei curricula formativi rappresenta la sfida più urgente per la formazione del designer contemporaneo. Il settore dell'educazione è oggi uno degli ambiti maggiormente interessati dalla diffusione delle tecnologie di IA. Come evidenziato dal Future of Jobs Report del World Economic Forum, tra il 2016 e il 2024 l'istruzione si colloca al secondo posto per crescita relativa della concentrazione di tecnologie IA, preceduta soltanto dal settore Technology, Information and Media. Questo dato conferma come l'intelligenza artificiale stia trasformando rapidamente anche i contesti formativi, pur in presenza di sistemi educativi che spesso faticano ad adeguarsi alla velocità del cambiamento tecnologico²⁶⁶.

²⁶⁶ World Economic Forum, *Why AI literacy is now a core competency in education*, 2025.

Il World Economic Forum documenta inoltre che il 40% delle competenze richieste dalla forza lavoro globale



²⁶⁷ Digital Education Council, *AI adoption is nearly universal among students, but confidence is not*, DEC Insights, 2025.

Figura 6.1: crescita della concentrazione delle tecnologie di intelligenza artificiale nei diversi settori economici (2016–2024). Il settore dell'educazione emerge come il secondo ambito per crescita dell'adozione e della presenza di competenze legate all'IA²⁶⁶. Grafico rielaborato.

cambierà entro i prossimi cinque anni, rendendo necessaria una risposta sistemica che va ben oltre l'aggiunta di singoli moduli o corsi sull'intelligenza artificiale. Il dato è particolarmente allarmante se si considera che il 92% degli studenti universitari già utilizza strumenti di IA generativa, spesso senza alcuna guida istituzionale su come farlo in modo critico e consapevole²⁶⁶.

Infatti, il 65% degli studenti teme che l'intelligenza artificiale possa rendere l'apprendimento troppo superficiale e scoraggiare il pensiero critico e la creatività²⁶⁷.

Preoccupazione degli studenti riguardo alla profondità e alla qualità dell'apprendimento con l'AI.

Affermazione:

Mi preoccupa che l'uso dell'AI possa rendere l'apprendimento troppo superficiale e scoraggiare il pensiero critico e la creatività.

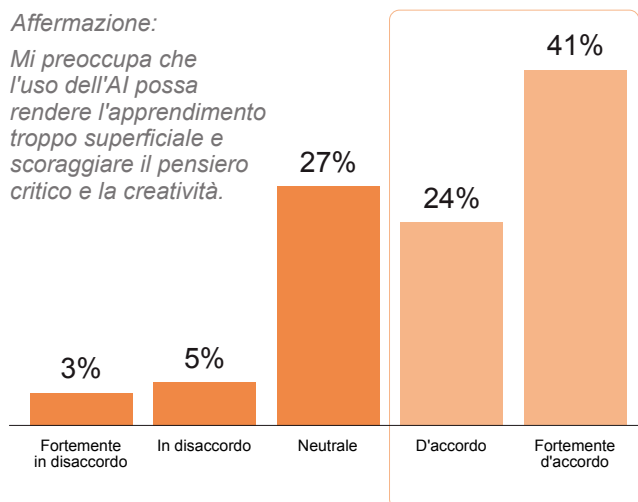


Figura 6.2: Percentuale di accordo degli studenti riguardo all'impatto negativo dell'intelligenza artificiale sulla profondità e sulla qualità dell'apprendimento²⁶⁷. Grafico rielaborato.

6.1.2 L'impreparazione istituzionale e il modello formativo obsoleto

Nel contesto italiano, la proposta avanzata da Confindustria nel maggio 2026 per un grande piano di formazione sull'intelligenza artificiale già dalla scuola superiore parte dall'industria e non dall'università, segnalando una disconnessione tra il sistema produttivo e quello educativo. Il presidente Orsini ha indicato la necessità di un progetto di formazione da iniziare nel ciclo delle superiori di secondo grado per tutti i giovani, con il sistema industriale pronto a mettersi in gioco in prima persona. La risposta positiva del governo conferma la consapevolezza del problema, ma l'assenza di una strategia organica per il design rivela come questo settore specifico rimanga ai margini della pianificazione formativa nazionale²⁶⁸.

²⁶⁸ L'Espresso, *Ministero dell'Istruzione, intelligenza artificiale, storia e geografia nei licei*, 2026

La denuncia che emerge dal confronto con il mondo accademico è netta: non esiste una strategia comune per le università italiane e più in generale per il sistema dell'istruzione del Paese. Alcuni atenei istituzionalizzano l'uso di ChatGPT senza rendersi conto del travaso di conoscenza verso piattaforme commerciali, mentre altri restano prudenti nella totale assenza di linee guida significative. Questa frammentazione non è soltanto un problema di coordinamento tra istituzioni, ma il sintomo di un'inadeguatezza più profonda, che riguarda il modello pedagogico stesso. Il modello formativo, progettato come un percorso lineare e incrementale di conoscenza, è sempre più inadeguato. Nell'era dell'intelligenza artificiale, alla specializzazione deve affiancarsi la flessibilizzazione, per adattare i professionisti a contesti lavorativi in costante mutamento.

Il dualismo tradizionale tra accademia e istituti tecnici professionalizzanti è oggi arricchito da università telematiche, nazionali e straniere, e da formazione job-ready erogata dai colossi tecnologici come Google, in grado di raggiungere e formare quante più persone possibili. Tutti questi attori si stanno interrogando su come integrare l'IA generativa nei percorsi formativi, ma senza una strategia chiara. L'urgenza di formare al futuro è confermata dalla previsione che il 40% dei datori di lavoro intende ridurre la forza lavoro entro il 2030 attraverso l'automazione IA, rendendo imprescindibile una risposta sistemica e non episodica da parte delle istituzioni educative²⁶⁹.

²⁶⁹ Guido Saracco, *Alleati digitali: la nostra AI personale*, Rizzoli, 2026

6.2 Lo studente di design davanti all'AI: la dimensione psicologica

L'incontro con l'intelligenza artificiale non è solo una questione di apprendimento tecnico, ma coinvolge profondamente la sfera emotiva e identitaria dello studente. Tra timore della sostituzione e curiosità per le nuove possibilità, la dimensione psicologica gioca un ruolo cruciale nell'adozione consapevole di queste tecnologie.

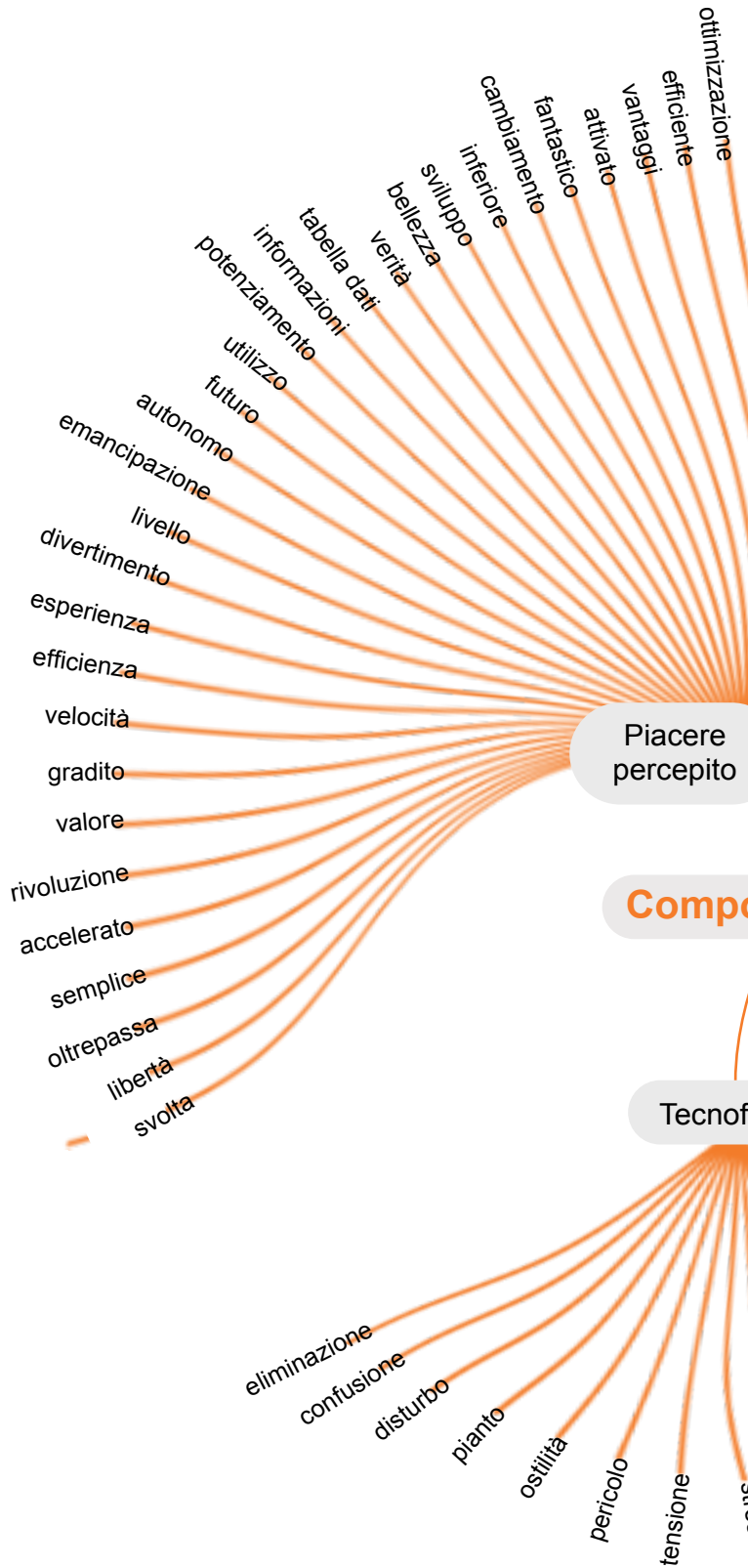
6.2.1 I determinanti dell'adozione: tecnofobia, piacere percepito e conoscenza soggettiva

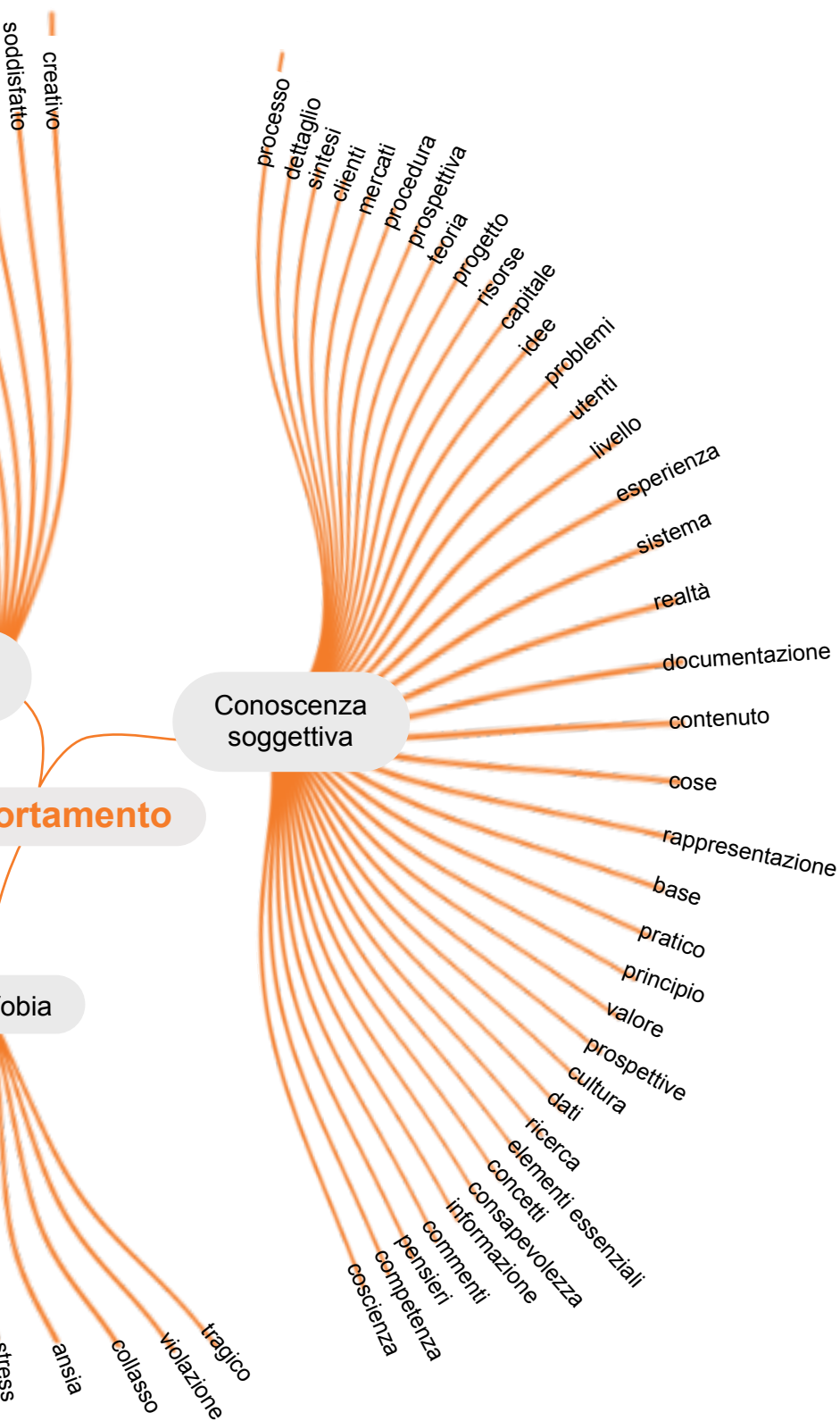
L'incontro tra lo studente di design e l'intelligenza artificiale non è un evento neutro, ma un processo mediato da una complessa architettura psicologica. Uno studio condotto su 313 studenti di design cinesi ha costruito un modello esteso, integrando la teoria della coerenza affettivo-cognitiva con la Unified Theory of Acceptance and Use of Technology (UTAUT), per identificare i fattori che orientano l'intenzione d'uso degli strumenti di IA applicati al design. I risultati rivelano che tecnofobia, piacere percepito, aspettativa di prestazione e aspettativa di sforzo influenzano direttamente l'intenzione d'uso, mentre l'influenza sociale e la conoscenza soggettiva agiscono indirettamente, mediate dal piacere percepito e dall'aspettativa di sforzo. Il dato è significativo: sapere di sapere usare uno strumento non basta a motivarne l'uso, se non si traduce prima nella sensazione che sia maneg-

Figura 6.3: Mappatura dei costrutti semantici che determinano l'attitudine dei futuri designer verso l'adozione dell'IA, integrando il modello UTAUT con la coerenza affettivo-cognitiva.²⁷⁰ Grafico rielaborato.

²⁷⁰ Chun-Hsin Wu et al., *The influence of subjective knowledge, technophobia and perceived enjoyment on design students' intention to use artificial intelligence design tools*, International Journal of Technology and Design Education, Springer Link, 2024.

²⁷¹ *How AI can help in the creative design process*, The Conversation, 2024.





petenza operativa. In un campione di quaranta dirigenti senior, la mancanza di competenze e talento è stata citata ventotto volte come barriera principale all'adozione dell'intelligenza artificiale, superando preoccupazioni relative a casi d'uso poco chiari, sicurezza dei dati, conformità normativa e costi. Questo dato rivela come il gap formativo si manifesti simultaneamente come gap psicologico, una dimensione che i curricula tradizionali non sono progettati per affrontare. L'implicazione per la formazione del designer è diretta: la dimensione sociale dell'apprendimento è fondamentale per superare la tecnofobia, e la formazione in presenza con interazione tra pari rappresenta il formato naturalmente compatibile con la didattica del design²⁷².

²⁷² World Economic Forum, *AI skills at scale: How to keep people in the loop*, 2026.

6.2.2 La catena virtuosa: autoefficacia, riduzione dell'ansia, cognizione creativa

Superata la fase iniziale di resistenza, l'interazione con l'IA può innescare una catena virtuosa di effetti psicologici positivi. Gli studenti che ricevono feedback potenziato dall'IA mostrano un incremento del 40% nella generazione di soluzioni multiple e del 25% nell'originalità in studi condotti su base biennale. La condizione necessaria per l'attivazione di questa catena è la co-creazione attiva: la mera consultazione passiva dello strumento non produce gli stessi effetti. L'IA non è un fine in sé, ma un mezzo per sbloccare potenzialità creative altrimenti inibite dall'ansia e dalla mancanza di fiducia nelle proprie capacità²⁷³.

²⁷³ Younjung Hwang et al., *The influence of generative artificial intelligence on creative cognition of design students: a chain mediation model of self-efficacy and anxiety*, *Frontiers in Psychology*, 2024.

Questo apprendimento, tuttavia, non si esaurisce nel momento in cui lo strumento è presente. Riprendendo lo studio della Oklahoma State University emerge un dato particolarmente prezioso: il gruppo che aveva utilizzato l'IA in una prima attività progettuale ha dimostrato un processo creativo più ricco anche in una seconda attività condotta senza alcun supporto digitale, probabilmente perché l'esperienza con lo strumento aveva già insegnato agli studenti a pensare e sviluppare idee in modo più efficace. È la controprova empirica della catena virtuosa appena descritta: l'IA non produce dipendenza ma costruisce competenza che permane oltre l'uso stesso dello strumento. Per la didattica del design, questo significa che l'esposizione guidata all'IA nelle fasi iniziali della formazione può avere effetti benefici anche sulle attività successive in cui lo studente opera in piena autonomia²⁷¹.

Gaetano Di Tondo

Dirigente aziendale

Manager di consolidata esperienza nell'ambito della comunicazione e dell'innovazione, Gaetano Di Tondo ricopre oggi il ruolo di VP, Head of Communication & Institutional Relations Strategy di TIM. Contemporaneamente, in Olivetti, riveste la carica di Direttore Institutional ed External Relations ed è Presidente dell'Associazione Archivio Storico Olivetti.

Questa pluralità di ruoli gli permette di esplorare come la memoria storica possa trasformarsi in un "esercizio attivo di futuro" e in uno strumento strategico per l'innovazione nelle imprese contemporanee.



In un suo contributo per la rivista *Culture Digitali*, intitolato “Educare al Futuro”, ha sottolineato l’importanza della “Human Centricity Olivettiana” nell’era dell’AI. A suo parere, come si può garantire che l’adozione dell’intelligenza artificiale resti davvero centrata sulle persone e in linea con i valori da lei evidenziati?

Credo che oggi, di fronte all’intelligenza artificiale, siamo chiamati a una scelta che non è tecnologica, ma profondamente umana, tenendo presente il valore del concetto olivettiano di Human Centricity.

Leggiamo ormai quotidianamente articoli e approfondimenti: possiamo lasciare che siano le macchine a dettare il ritmo del cambiamento, oppure possiamo scegliere – come faceva Olivetti – di mettere al centro una visione: quella di una tecnologia che non domina, ma accompagna. Che non sostituisce, ma amplifica.

La “human centricity” non è uno slogan: è una disciplina. Significa progettare con un’idea di persona, con un’idea di comunità, chiedendosi, ogni volta, non quanto sia potente un algoritmo, ma quanto sia giusto, utile, capace di generare valore condiviso.

In questo senso, l’intelligenza artificiale dovrebbe essere pensata come una nuova forma di alleanza: una intelligenza che cresce accanto all’uomo, che ne amplia lo sguardo senza mai oscurarne la responsabilità. Perché il rischio più grande non è delegare alle macchine delle funzioni, ma delegare ad esse il senso, e farci guidare da loro.

E la lezione olivettiana ci ricorda che il senso non può che rimanere umano.

Oggi, di fronte all’AI, siamo chiamati a rinnovare questa responsabilità. Possia-

mo limitarci a inseguire l’efficienza, oppure possiamo cogliere questa trasformazione come un’occasione per ripensare più profondamente il nostro rapporto con la conoscenza, con la memoria.

Proprio all’interno dell’Archivio Storico Olivetti, di cui lei è presidente, la memoria viene descritta non come conservazione ma come “esercizio attivo di futuro”. Partendo da questa definizione, in che modo l’intelligenza artificiale può contribuire a trasformare la memoria storica in uno strumento strategico per l’innovazione nelle imprese contemporanee?

All’Archivio Storico Olivetti abbiamo imparato che la memoria non è mai nostalgia, ma energia e stimolo per il futuro. È una materia viva, che continua a generare domande, visioni, possibilità, idee, innovazione.

Oggi l’intelligenza artificiale ci offre uno strumento straordinario: la possibilità di attraversare questa memoria in modo nuovo, di abitarla, di interrogarla. Non più come un deposito, ma come un organismo reattivo.

L’AI ci permette di scoprire connessioni potenzialmente invisibili, di far emergere significati che a volte sfuggono ad una prima analisi, di rendere il passato una vera infrastruttura per il futuro.

Ma il punto non è solo tecnologico. Il punto è che, grazie all’AI, possiamo tornare a fare quello che Olivetti faceva naturalmente: progettare il futuro a partire da una consapevolezza profonda della nostra storia. Un’impresa che sa dialogare con la propria memoria non replica il passato: lo trasforma, impara, e ne trae benefici.

L’intelligenza artificiale diventa quindi

uno strumento per riattivare questa continuità: tra ciò che siamo stati e ciò che possiamo diventare, con una guida umana.

L'esperienza olivettiana ci ricorda che tecnologia, impresa e cultura non sono dimensioni separate, ma parti di un unico disegno. Un disegno in cui il progresso non è mai fine a sé stesso, ma sempre orientato alla qualità della vita, alla dignità del lavoro, alla crescita della comunità.

Considerando la sua esperienza manageriale tra cultura, comunicazione, tecnologia e innovazione, lei ha una prospettiva privilegiata da cui osservare la continua evoluzione delle competenze necessarie nel mondo del lavoro. Ad oggi, quali ritiene siano quelle fondamentali per i giovani professionisti coinvolti nell'utilizzo di sistemi di intelligenza artificiale e come le istituzioni educative e le aziende possono collaborare per formare queste figure?

Ai giovani oggi si chiede molto. Ma forse, prima ancora, si chiede loro qualcosa di diverso. Non basta più saper usare strumenti complessi.

Serve la capacità di stare dentro la complessità, e sapersi muovere in modo corretto, e rapidamente. Le competenze tecniche sono fondamentali, ma non sono sufficienti. Perché lavorare con l'intelligenza artificiale significa muoversi in un territorio dove le decisioni hanno impatti profondi: sul lavoro, sulla società, sulle persone.

Per questo credo che le competenze più importanti siano quelle che permettono di tenere insieme mondi diversi: la logica e l'immaginazione, il dato e l'etica, la tecnologia e la sensibilità umana.

Servono professionisti capaci non solo di costruire modelli, ma soprattutto di comprenderne le conseguenze. Non solo di generare soluzioni, ma di attribuire loro un significato. E qui si apre una grande responsabilità condivisa.

Le università devono essere luoghi di pensiero aperto, di contaminazione, di sperimentazione e le imprese devono diventare spazi di apprendimento continuo, dove si cresce lavorando su problemi reali, ma con uno sguardo ampio. Come nell'esperienza olivettiana, il sapere non può essere frammentato: deve essere integrato, vissuto, messo in relazione con relazioni e pensieri aperti. La vera competenza del futuro oltre a saper parlare con le macchine, guidandole, sarà continuare a parlare, con intelligenza e responsabilità, tra esseri umani. Se sapremo mantenere viva una visione umanistica dell'innovazione, l'intelligenza artificiale potrà diventare non solo una leva di sviluppo, ma una straordinaria opportunità per rendere il lavoro più consapevole e le organizzazioni più responsabili.

6.3 Principi per una nuova pedagogia del design

Riformare l'educazione al design nell'era dell'IA significa andare oltre l'insegnamento dello strumento per approdare a una nuova filosofia pedagogica. È necessario definire principi che mettano al centro il pensiero critico, la responsabilità etica e la capacità di integrare la potenza computazionale con la profondità del giudizio umano.

6.3.1 Capire prima di usare: il sistema dietro lo strumento

Il pensiero critico sull'intelligenza artificiale non può essere relegato a modulo opzionale o ad appendice di corsi informatici preesistenti, ma deve costituire il fondamento dell'intero percorso formativo del designer. Come evidenziato da Mehran Sahami presso l'Università di Stanford, non è più ammissibile trattare l'IA esclusivamente come uno strumento operativo, poiché gli studenti necessitano di un curriculum sistemico che spazi dalla comprensione dei meccanismi algoritmici fondamentali, all'analisi di fenomeni quali allucinazioni e bias, fino alla verifica critica degli output e alle tecniche avanzate di prompting. Senza questa formazione strutturale, una quota stimata tra il 70% e l'80% degli studenti rischia di utilizzare l'IA unicamente per abbreviare i tempi di esecuzione del lavoro, anziché per potenziare le proprie capacità progettuali e cognitive. Il concetto di AI driver's li-

²⁷⁴ Stanford HAI, *AI challenges core assumptions in education*, Stanford Institute for Human-Centered Artificial Intelligence, 2026.

cense cattura efficacemente questa urgenza pedagogica, poiché così come la conduzione di un autoveicolo richiede la conoscenza delle regole della strada e dei limiti strutturali del mezzo, l'impiego dell'IA nella progettazione presuppone la comprensione approfondita dei meccanismi, dei vincoli e dei rischi intrinseci al sistema tecnologico²⁷⁴.

Per mappare e strutturare tali competenze all'interno dei percorsi di studio, la letteratura scientifica internazionale ha elaborato diversi modelli di AI Literacy, che definiscono i requisiti necessari per un'interazione consapevole con queste tecnologie.

Il framework del Teaching Commons di Stanford articola la competenza in quattro domini interconnessi, partendo dalla functional literacy relativa all'accesso e all'uso degli strumenti, per passare alla ethical literacy riguardante le posizioni etiche personali su bias, privacy ed equità, fino alla rhetorical literacy, dedicata all'analisi critica del linguaggio e degli output, e alla pedagogical literacy, incentrata sull'integrazione di pratiche basate su evidenze²⁷⁵.

²⁷⁵ Stanford Teaching Commons, *Understanding AI Literacy*, Stanford University Teaching Guides, 2026.



Figura 6.4: I quattro domini interconnessi del framework di AI literacy sviluppato dal Teaching Commons di Stanford.²⁷⁵ Grafico rielaborato.

In modo complementare, l'approccio socio-tecnico promosso dall'Organisation for Economic Co-operation and Development (OECD) interpreta l'IA come un sistema emergente dall'interazione tra componenti tecnologi-

che, quali algoritmi, hardware e dati, e dimensioni sociali relative a fattori istituzionali, economici e comportamentali. Questo modello individua quattro macro-dimensioni basate sul comprendere l'IA, sul pensiero critico e la valutazione, sull'uso responsabile e sul principio del human-in-the-loop. Quest'ultima dimensione analizza nello specifico l'influenza dell'IA sul comportamento umano e sottolinea l'importanza di pratiche progettuali capaci di preservare l'agency, la supervisione e l'accountability del designer²⁷⁶.

²⁷⁶ *A socio-technical approach to AI literacy*, OECD.AI Work, 2026.

Figura 6.5: Modello ciclico per la competenza sull'IA, articolato in quattro macro-aree interconnesse che guidano l'utente dalla comprensione tecnica alla valutazione critica dei contenuti.²⁷⁶ Grafico rielaborato.



6.3.2 L'etica come parametro progettuale costante

La separazione tradizionale tra formazione tecnica e formazione etica costituisce un'eredità pedagogica che l'era dell'IA generativa rende insostenibile. La dicotomia critica tra l'insegnamento del machine learning e la trattazione delle sue implicazioni sociali, della privacy e dell'equità deve essere superata attraverso un approccio integrato. Poiché ogni output generato da un modello algoritmico incorpora inevitabilmente una specifica visione del mondo sedimentata nei dati di addestramento, il designer ha il compito etico di riconoscere e correggere tali asimmetrie, operando precisamente all'intersezione tra il sistema tecnico e il contesto sociale²⁷⁶.

Sul piano operativo, una delle metodologie didattiche più efficaci per sviluppare questa sensibilità è il red teaming: una pratica nata in ambito cybersecurity e oggi

adottata trasversalmente per testare un sistema sondandone deliberatamente debolezze, rischi e conseguenze inattese. Applicato alla didattica, il principio si traduce in laboratori in cui gli studenti stessi cercano di far fallire i sistemi di IA, esplorandone bias, allucinazioni e punti ciechi prima di affidarsi nel proprio lavoro progettuale. Un workshop reale condotto da esperti di trust and safety offre un esempio concreto di come questo possa funzionare anche al di fuori dei contesti puramente tecnici: i partecipanti hanno simulato interazioni con due chatbot – un "terapeuta virtuale" e un assistente didattico storico – scoprendo che i modelli, pur respingendo correttamente le richieste esplicitamente dannose, faticavano proprio sugli scenari più ambigui e dipendenti dal contesto, dove la stessa informazione può essere utile o dannosa secondo l'intenzione di chi la richiede. È esattamente in questi margini, dove il sistema sembra affidabile ma non lo è, che si forma il giudizio critico che nessuna lezione teorica sui bias potrebbe insegnare con la stessa efficacia. Attraverso questa simulazione diretta dei limiti e dei punti di fallimento dei sistemi di IA, l'etica cessa di essere un'astratta enunciazione di principi per trasformarsi in una competenza pratica di verifica e validazione del progetto²⁷⁷.

²⁷⁷ Nick Shockey, Kathryn Zickuhr et al., *Red Teaming Generative AI in Classrooms and Beyond*, Tech Policy Press, 2024.

6.3.3 Costruire il giudizio prima di introdurre lo strumento

L'ordine di introduzione degli strumenti nel percorso formativo non è una questione logistica ma epistemologica. L'IA fa bene a chi le cose le sa e fa malissimo a chi non le sa: questa sintesi esprime una verità confermata da dati empirici²⁷⁸. Un esperimento condotto da Stanford HAI lo dimostra con chiarezza: gli studenti che avevano utilizzato l'IA e poi ne erano stati privati hanno performato quattro volte peggio rispetto a chi non l'aveva mai usata, non perché mancasse lo strumento, ma perché non avevano mai costruito la competenza sottostante, e in più avevano sviluppato la convinzione che l'IA fosse più creativa di loro²⁷⁴.

²⁷⁸ Wired Italia, *Capitale semantico: la vera posta in gioco della rivoluzione digitale* secondo Luciano Floridi, 2026.

Il paradosso cognitivo dell'IA nell'educazione è documentato con rigore: l'esposizione prolungata all'IA senza fondamenta cognitive solide erode sia la memoria sia il pensiero critico. Uno studio su 73 studenti universitari di Information Science ha dimostrato che il pretesting prima dell'uso dell'IA migliora ritenzione e coinvolgimento, mentre l'esposizione prolungata senza questa precauzione conduce a un declino della memoria. L'IA potenzia

l'accessibilità ma può indebolire la ritenzione se sovrautilizzata; per massimizzarne i benefici, deve supportare e non sostituire le strategie di apprendimento guidate dall'umano. Il principio si lega direttamente alla Teoria del Carico Cognitivo: l'IA deve ridurre il sovraccarico estraneo ma mantenere attivo il coinvolgimento cognitivo, evitando che l'accesso troppo facile alle risposte sostituisca il lavoro mentale necessario per costruire comprensione duratura²⁷⁹.

²⁷⁹ Binny Jose et al., *The cognitive paradox of AI in education: between enhancement and erosion*, Frontiers in Psychology, 2023.

6.3.4 L'apprendimento attraverso la pratica collaborativa

La collaborazione umano-IA come pratica didattica richiede che gli studenti apprendano lavorando con l'IA, riflettendo su cosa cambia nel processo e nel risultato. Tuttavia, la mera esposizione ripetuta non produce miglioramento: le coppie umano-IA non migliorano per semplice ripetizione, ma richiedono interventi deliberati. I designer che co-creano con successo esplorano sistematicamente i limiti dell'IA nelle fasi iniziali, si auto-spiegano i comportamenti del sistema, utilizzano sketching e riflessione attiva come strumenti di mediazione. Le sfide vanno oltre l'interfaccia: comunicare obiettivi progettuali, interpretare il ragionamento del sistema, prevederne il comportamento. Quattro strategie evidence-based emergono dalla ricerca: mini-esperimenti guidati sui limiti e le capacità dello strumento; riflessione progettuale strutturata dopo ogni ciclo di interazione; comunicazione multimodale che integra sketch, prompt, annotazione e verbalizzazione; peer-guided learning con il docente come guida alla collaborazione²⁸⁰.

²⁸⁰ Jingwan Li et al., *Exploring Challenges and Opportunities to Support Designers in Learning to Co-create with AI-based Manufacturing Design Tools*, Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems, ACM Digital Library, 2023.

Due studi empirici mostrano cosa significhi concretamente questo principio, nei suoi lati positivi e nei suoi limiti. Il primo riguarda l'integrazione dell'AI-Generated Content (AIGC) nell'intero processo di Design Thinking, non in fasi isolate: uno studio condotto con cinquantadue studenti ha integrato ChatGPT e Midjourney nelle fasi di Empathy, Define, Ideate e Prototype, con l'IA che ha funzionato come esperto interdisciplinare in grado di colmare gap di conoscenza specialistica. I risultati mostrano miglioramenti significativi in usabilità, motivazione intrinseca e creatività. Emergono tuttavia anche limiti importanti: la overdependence dallo strumento, la complessità del prompt engineering che paradossalmente aumenta il carico cognitivo anziché ridurlo, e la generazione di output con dettagli irrealistici che richiedono una reinterpretazione critica da parte dello studente²⁸¹.

²⁸¹ Chun-Hsiang Chen et al., *Enhancing human computer collaboration in design process: How AI-generated content elevates students' creativity*, International Journal of Educational Technology in Higher Education / ScienceDirect, 2025.

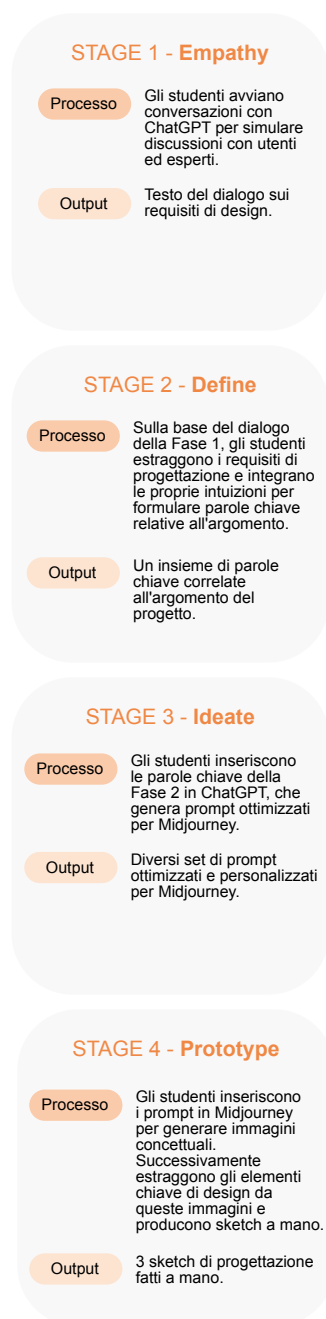


Figura 6.6: Dettagli procedurali del metodo basato su AIGC applicato alle quattro fasi del processo di progettazione (Empathy, Define, Ideate e Prototype).²⁸¹ Grafico

Il secondo studio, di tipo think-aloud, segue designer alle prese con strumenti di manufacturing come Fusion360 Generative Design e SimuLearn, e rivela che le sfide della co-creazione non sono di interfaccia ma di collaborazione. Il conflitto centrale riguarda la specificazione dei parametri upfront richiesta dal sistema rispetto alla modellazione iterativa propria del processo progettuale tradizionale: i modelli generativi richiedono che i vincoli siano definiti a priori, mentre il designer è abituato a scoprirli progressivamente durante il processo. Questo conflitto è radicale, e si estende oltre il manufacturing a tutti i domini della co-creazione, indicando un problema strutturale nell'interazione tra logica algoritmica e logica progettuale²⁸⁰.

Entrambi gli studi convergono su un punto: superare questi limiti richiede una guida, non solo uno strumento. Il peer-guided study emerge come modello per il nuovo ruolo del docente. Le guide umane efficaci forniscono istruzioni step-by-step, stimolano la riflessione progettuale, suggeriscono strategie alternative e usano comunicazione multimodale. Quando la guida sollecita feedback e riflessione sui design generati, il designer è obbligato a esplicitare le proprie intenzioni, e da quel momento sia la guida sia il designer riescono a coordinare meglio le azioni successive. Il docente di design non è più chi mostra come si fa, ma chi insegna a dirigere un processo in cui la macchina genera e l'umano decide²⁸⁰.

Il potere moltiplicativo di questo apprendimento guidato trova conferma nei dati del programma WEF AI for Impact: il 91% dei volontari-mentori ha dichiarato che l'attività di mentoring ha aumentato la propria comprensione dell'IA generativa, e la loro confidenza è salita da 3.45 a 4.40 su una scala a cinque punti. I lavoratori in prima linea dichiarano quasi universalmente l'intenzione di condividere con i colleghi quanto appreso, innescando un effetto a cascata. L'implicazione per il design è diretta: il peer-guided learning non è solo una strategia didattica tra le tante, ma un meccanismo di scaling della competenza che trasforma ogni studente formato in un potenziale formatore dei compagni. Insegnare accelera l'apprendimento²⁷².

Un'evoluzione concreta di questa pedagogia collaborativa è rappresentata dal modello Think-Pair-LLM-Pair-Share. Nel flusso tradizionale Think-Pair-Share, l'alleato digitale si inserisce come interlocutore aggiuntivo, creando un ciclo in cui gli studenti pensano individualmente, discutono tra pari, consultano l'IA, tornano a di-

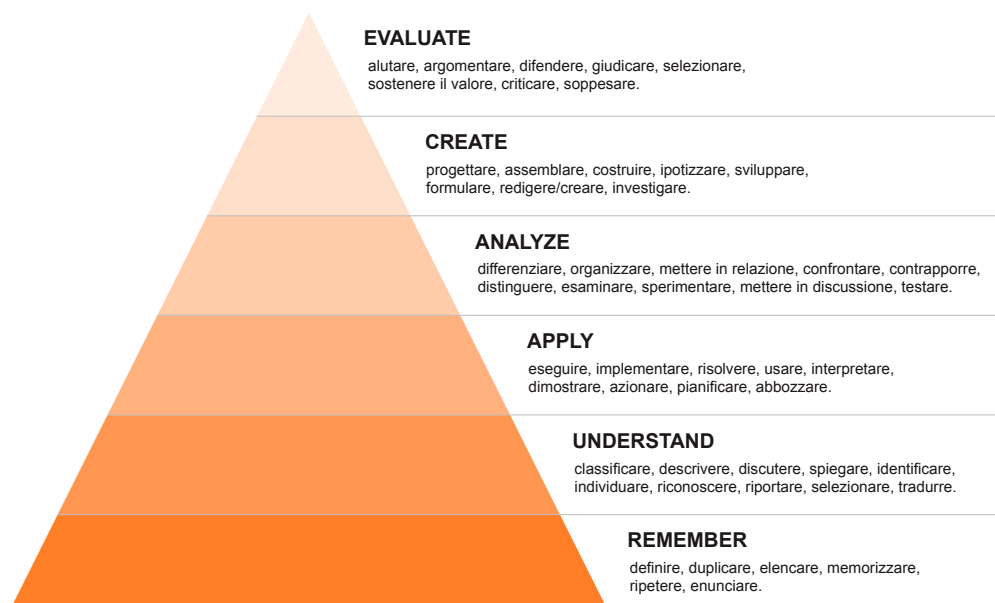
scutare tra pari integrando quanto appreso dallo strumento, e infine condividono con l'intera classe. L'alleato digitale diventa il latore di un'idea in più, di un'analisi in più nei dibattiti tra studenti. Nelle classi numerose, dove il docente non può seguire ogni singolo studente, questa pedagogia consente di personalizzare l'apprendimento mantenendo la dimensione sociale della formazione²⁶⁹.

6.3.5 Ripensare la valutazione: l'importanza del processo

Figura 6.7: Adattamento della tassonomia di Bloom all'interno di un contesto universitario guidato dall'intelligenza artificiale. Si sposta l'accento dall'automazione dei processi di memorizzazione (Remember) verso lo sviluppo e l'esercizio delle facoltà umane superiori di valutazione critica (Evaluate) e co-creazione complessa (Create).²⁸² Grafico rielaborato.

L'intelligenza artificiale generativa infrange un assunto su cui la didattica si è basata per decenni: che un prodotto forte sia prova di un processo di apprendimento altrettanto forte. Oggi uno studente può produrre un elaborato impressionante senza aver appreso quasi nulla, perché il prodotto stesso può essere in larga parte opera dello strumento. La conseguenza pedagogica è che la valutazione deve spostarsi dal risultato finale al processo che lo ha generato.

La Tassonomia di Bloom offre la direzione per questo spostamento. Elaborata dallo psicologo dell'educazione Benjamin Bloom nel 1956 e da allora un riferimento standard nella progettazione didattica, la tassonomia ordina le abilità cognitive in una scala crescente di complessità: ricordare, comprendere, applicare, analizzare, valutare, creare. Riconcepita per l'era generativa, la tassonomia mostra che l'IA automatizza efficacemente i livelli infe-



²⁸² John T. E. et al., *Cultivating independent thinkers: The triad of artificial intelligence, Bloom's taxonomy and critical thinking in assessment pedagogy*, Education and Information Technologies, Springer Link, 2025.

riori – ricordare, comprendere – mentre il curriculum deve concentrarsi sui livelli superiori – analizzare, valutare, creare – che restano il vero terreno dell'apprendimento umano. L'assessment operativo per il design cambia di conseguenza: non chiede più di generare una sedia, ma di valutare le cinquanta proposte dell'IA, giustificare una direzione progettuale e articolare perché le altre sono state scartate. Integrata con il costruttivismo sociale, l'apprendimento diventa così un processo dinamico di costruzione di significato attraverso attività collaborative, dialogo e riflessione²⁸².

La proposta di posizionare la valutazione al vertice della piramide cognitiva, sopra la stessa creazione contrariamente all'ordine originale di Bloom, trova fondamento nell'osservazione che l'IA generativa è capace di operare a tutti i livelli tassonomici, inclusi quelli creativi. Poiché l'IA può creare, gli studenti devono essere guidati e supportati nel valutare e analizzare i testi generati, al fine di verificarne affidabilità e attendibilità. Il ciclo tra valutare e creare diventa un loop di feedback continuo: la capacità degli studenti di porre le domande giuste, formulando prompt efficaci, e di valutarne i risultati influenza non solo la qualità delle risposte ottenute ma anche lo sviluppo delle loro stesse competenze valutative²⁸².

L'IA può anche fungere da strumento di valutazione in itinere. Il loop di valutazione continua e il conseguente adeguamento dell'azione formativa verso il singolo studente accorciano drammaticamente i tempi necessari per il corretto indirizzamento dell'apprendimento rispetto a classi gestite dal solo professore. Tuttavia, se si utilizza l'IA in fase di valutazione occorre chiarire molto bene a monte su quali criteri si basano i giudizi emessi, altrimenti può ingenerarsi sfiducia negli studenti e possono nascere contenziosi che minano l'adozione estesa di questi strumenti. Se le regole sono trasparenti, l'adozione di uno strumento digitale può essere percepita dai discenti addirittura come più imparziale rispetto al giudizio esclusivamente umano²⁶⁹.

6.3.6 L'interdisciplinarietà come architettura del curriculum

La natura sistemica dell'intelligenza artificiale richiede il definitivo superamento dei silos disciplinari nell'educazione del designer. Il framework AILit dimostra la necessità di una struttura formativa esplicitamente interdisciplinare, mappando ventitré competenze core che attra-

versano in modo trasversale ambiti quali la Computer Science, la Media Literacy, il Design Thinking, la Data Science e l'Ethics²⁶⁶.

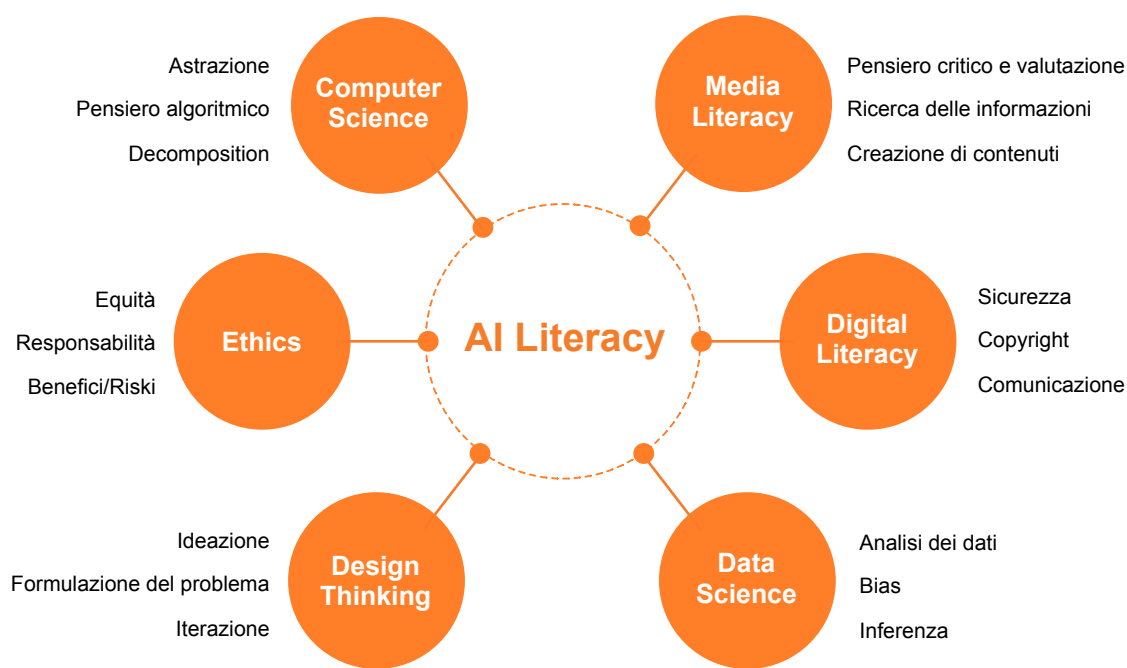


Figura 6.8: Rappresentazione del framework AILit che mappa la natura intrinsecamente multidisciplinare dell'alfabetizzazione all'IA attraverso l'intersezione di sei macro-aree scientifiche e pedagogiche fondamentali.²⁶⁶ Grafico rielaborato.

L'introduzione dell'IA rende più urgente la capacità di integrare saperi eterogenei all'interno di una sintesi progettuale coerente e unitaria. Sotto il profilo didattico, il modello degli hackathon strutturati sul principio delle Communities of Practice dimostra il valore dell'apprendimento tra pari con background formativi differenti, in cui l'IA generativa può essere adottata dagli studenti come un esperto interdisciplinare virtuale atto a colmare temporaneamente gap di conoscenza specialistica, a condizione che l'allievo abbia preventivamente sviluppato il lessico tecnico e i fondamenti epistemologici necessari per condurre un dialogo critico ed efficace con il sistema²⁶⁶.

6.3.7 L'indole personale al centro della formazione

Il valore irriducibile del designer risiede nella sua indole personale, ovvero nella sensibilità individuale, nella biografia culturale e nel modo unico in cui ciascun soggetto interpreta e attribuisce senso al mondo. In uno scenario in cui i sistemi di intelligenza artificiale tendono a standardizzare le soluzioni formali e concettuali, l'istituzione formativa deve perseguire l'obiettivo opposto, strutturando percorsi personalizzati che valorizzino il singolo in

quanto essere umano. Poiché il significato è un atto intrinsecamente umano e legato all'intenzionalità e alla sensibilità culturale, la formazione non può limitarsi ad addestrare operatori di software. Questa urgenza è supportata dalle ricerche sull'eudaimonic well-being, le quali confermano che gli individui orientati alla crescita personale ricercano costantemente autenticità, eccellenza e significato; in tal senso, la tecnologia educativa può supportare lo sviluppo dei discenti senza inibirne l'autonomia, a condizione di essere configurata come uno strumento di supporto e mai di sostituzione delle facoltà cognitive²⁷⁸.

In questa prospettiva, i sistemi algoritmici avanzati possono essere intesi come strumenti per la personalizzazione dell'apprendimento, capaci di agire come interlocutori critici o sparring partner che si adattano progressivamente alle risposte, alle attitudini e alle lacune del singolo studente. Ciò consente di stabilire obiettivi formativi personalizzati e di monitorare l'evoluzione delle competenze nel tempo. Sotto il profilo pedagogico, questa dinamica rende operativa la distinzione etimologica tra istruzione ed educazione. Se l'istruzione, dal latino *instruere*, descrive l'atto di inserire o immettere nozioni e dati, il concetto di educazione, dall'etimo *ex-ducere*, attiene all'atto di trarre fuori e sviluppare il potenziale endogeno dell'individuo. Delegando gran parte delle attività puramente istruttive e della trasmissione nozionistica a dispositivi digitali avanzati, i docenti possono focalizzare la propria azione didattica sull'educazione propriamente detta: l'insegnamento dei valori etici, la complessità relazionale, la collaborazione interpersonale e lo sviluppo del pensiero critico. Per la pedagogia del design, questo spostamento di paradigma risulta essenziale, poiché liberare il tempo didattico dall'erogazione di contenuti ripetibili significa poter dedicare lo spazio del laboratorio alla coltivazione di quella sensibilità culturale e individuale che rende l'azione del progettista irriducibile all'automazione computazionale²⁷⁹.

6.4 La rivincita delle humanities

Il ritorno strategico della formazione umanistica nell'era dell'IA generativa non è un paradosso né una posizione nostalgica, ma una necessità strutturale documentata da evidenze convergenti. Il professor Saracco ne parla esplicitamente nel suo libro *Alleati Digitali*:

*Penso che paradossalmente, nell'era dell'IA, [le humanities] saranno trasversalmente più importanti di prima, nella formazione degli individui, le scienze dell'uomo e della società. Dovremo guadagnare flessibilità mentale e senso critico per spostare l'attenzione dove serve, in barba a logiche commerciali che te la catturano dove interessa ai gestori dei sistemi multimediali. Servirà poi dimostrare autocontrollo per dire basta quando l'interazione col digitale, su un determinato tema o contesto, diventa infruttuosa e solo una perdita di tempo. Ai nostri studenti e alle nostre studentesse dovremo far comprendere i meccanismi profondi di come siamo capaci di apprendere (metacognizione) per imparare a imparare meglio, dotandoli anche di competenze civiche, economiche, finanziarie e legali per inquadrare i loro pensieri e comportamenti nel contesto sociale o imprenditoriale²⁶⁹. - Guido Saracco, *Alleati Digitali* (2026)*

Anche le nuove indicazioni ministeriali italiane restituiscono slancio alla letteratura come specchio dell'esperienza umana e alla filosofia come esercizio di riflessione

ne, interrogazione, giudizio e argomentazione — esattamente le competenze che l'IA non automatizza. L'intreccio di arte, scienza, tecnologia e filosofia è necessario per comprendere e impiegare il pieno potenziale dell'IA generativa, e questo intreccio costituisce il terreno naturale del design. La cultura, lungi dall'essere un beneficiario passivo dell'innovazione, forma parte dell'ambiente abilitante attraverso cui le istituzioni apprendono, si adattano e mantengono legittimità. Più i sistemi tecnologici diventano capaci, più la cultura diventa strategica come framework attraverso cui le società determinano ciò che conta, esercitano giudizio e traducono la conoscenza in azione legittima²⁸³.

²⁸³ *Beyond data culture: human judgement, institutions and AI*, World Economic Forum, 2026.

Anche fuori dal dibattito accademico, la stessa intuizione emerge da chi lavora sul rapporto tra narrazione e tecnologia. Il regista Danny Boyle, riflettendo sul rapporto tra intelligenza artificiale e identità umana, ha osservato che in un mondo in cui le macchine possono produrre, elaborare e generare a velocità impensabili, è la cultura che preserva lo spazio del significato, della scelta e della responsabilità. Musei, biblioteche, archivi, scuole e università preservano la memoria, contestualizzano le informazioni e espongono gli individui a prospettive multiple e narrazioni in competizione. Nel farlo, coltivano capacità che diventano sempre più preziose nell'era dell'IA: la capacità di esercitare giudizio, navigare la complessità e comprendere le conseguenze più ampie delle decisioni. Per il designer, la formazione umanistica non è dunque un complemento decorativo ma il fondamento stesso della capacità critica e interpretativa che distingue il professionista dalla macchina²⁸⁴.

²⁸⁴ *Danny Boyle: per salvare la vita, scegliete la cultura*, Domusweb, 2020.

Alcune forme di expertise sopravvivono non perché sono documentate, ma perché sono trasmesse attraverso comunità di pratica e culture istituzionali. La sfida antropologica ed educativa per le scuole di design oggi non è semplicemente quella di insegnare a gestire e immagazzinare più dati, ma di preservare le condizioni relazionali e culturali attraverso cui la conoscenza supporta giudizi progettuali dotati di senso. Questa distinzione tra la mera gestione delle informazioni (knowledge management) e il governo consapevole delle condizioni che le trasformano in valore sociale (knowledge governance) cattura il cuore della questione. Come osservava il filosofo ungaro-britannico Michael Polanyi, "sappiamo più di quanto possiamo dire": un principio che nel design si traduce nella conoscenza tacita del progettista. È quella sensibilità estetica, materica, spaziale e relazionale, intrinsecamente legata alle discipline umanistiche, che si

apprende solo attraverso l'esperienza vissuta e il confronto tra pari. Nessun dataset potrà mai codificare questa dimensione profonda dell'intuito umano; ed è proprio qui, in questo nucleo irriducibile di conoscenza tacita e cultura materiale, che il designer trova la sua vera legittimazione e la sua definitiva rivincita sulla macchina²⁸³.

Guido Saracco

Ingegnere chimico,
professore universitario,
autore

Già Rettore del Politecnico di Torino e curatore di importanti eventi di divulgazione come la Biennale Tecnologia 2026, Guido Saracco è una figura centrale per discutere il ruolo della cultura e della formazione universitaria nell'era digitale. La sua visione dell'IA come "alleato" piuttosto che come semplice strumento è il fulcro del suo recente lavoro saggistico.

Viene interpellato per definire quali scelte concrete servano affinché le istituzioni educative preparino gli studenti a governare la tecnologia anziché subirla, trasformando profondamente l'offerta formativa attuale.



In qualità di curatore della Biennale Tecnologia 2026, quale ruolo possono avere eventi culturali come questo nel favorire una consapevolezza pubblica sull'intelligenza artificiale, estendendo il dibattito oltre i contesti tecnici?

Le Università pubbliche italiane sono oggi valutate dalla agenzia ANVUR anche per la cosiddetta Terza Missione ossia la valorizzazione della conoscenza, che si compone del trasferimento tecnologico, dove il Politecnico di Torino eccelle da tempo con il suo incubatore I3P e un ecosistema di innovazione che si sta progressivamente espandendo, come pure della condivisione della conoscenza, oggetto della vostra domanda.

Da due settimane fa abbiamo appreso che per il secondo quadriennio consecutivo il Politecnico di Torino guida la specifica graduatoria nazionale.

Questo ci conforta della correttezza del cammino intrapreso e responsabilizza per essere di guida ed esempio.

È sempre più necessario che le università si mettano in gioco fuori dai propri confini tradizionali (prima e seconda missione, ovvero didattica e ricerca) in particolare per portare ai cittadini e le cittadine consapevolezza delle principali sfide sociali da affrontare e fiducia nel fatto che con la scienza e la tecnologia, se correttamente indirizzate, si potranno trovare soluzioni efficaci.

Il successo della Biennale Tecnologia 2026 è anche garantito da una serie di novità strutturali introdotte: la scelta di ospitare più curatori ospiti al mio fianco, la collaborazione e coprogettazione di eventi insieme a altri enti culturali torinesi o altri festival di livello nazionale, la co-produzione di spettacoli.

In particolare per quest'ultimo aspetto il Politecnico di Torino ha creato, sempre

sotto la mia curatela, Prometeo-Tech Cultures. Prometeo è una nuova piattaforma di produzione di contenuti culturali per il public engagement del Politecnico di Torino che mette al centro l'analisi del complesso rapporto tra tecnologia e società. Prometeo si articola nell'ideazione e produzione di contenuti differenti: spettacoli teatrali, cortometraggi, libri di saggistica, pubblicazioni a fini didattici, podcast, video, lezioni per studenti della scuola e dell'università, contenuti per i social network, eventi di intrattenimento e dialogo con esperti, mostre, e molto altro.

Questa è la principale innovazione condotta per arrivare alla cittadinanza con contenuti e riflessioni attraverso gli strumenti e i media da ciascuno prediletti.

Nel suo libro *Alleati digitali. La nostra IA personale* descrive l'intelligenza artificiale come un "alleato" più che come un semplice strumento: come dovrebbe evolversi, secondo lei, l'offerta formativa delle università per preparare gli studenti a questo nuovo rapporto con la tecnologia?

L'avvento dell'IA generativa impone però ora una rivoluzione nella rivoluzione.

La nostra costituzione sancisce negli articoli 33 e 34 il diritto all'istruzione, stabilendo tra l'altro l'obbligo per lo Stato di offrire un sistema scolastico statale a tutti i giovani. Tra la generale inconsapevolezza della profondità del cambiamento che l'IA generativa può portare ai processi di formazione e al mondo del lavoro, ignota in primis alla maggior parte della cittadinanza e dei nostri politici, sono state avviate nel nostro Paese iniziative incoerenti sull'impiego dell'IA nell'istruzione. Il ministro dell'Istruzione e del Merito, On. Giuseppe Valditara,

dopo aver varato linee di indirizzo per negare l'uso degli smartphone fino a tutta la scuola secondaria di primo grado, ha annunciato nel settembre 2024 che già a partire dall'anno scolastico 24-25 in 15 classi di alcune scuole di Calabria, Lazio, Toscana, Lombardia partirà una fase di prova dell'uso dell'IA.

Intanto ChatGPT dilaga nell'uso tra i nostri studenti e le nostre studentesse, mentre nelle pieghe della coda del Piano Nazionale di Ripresa e Resilienza moltissime Scuole hanno concepito autonomamente corsi per l'aggiornamento specifico per i loro docenti e hanno iniziato a offrire anche moduli formativi per alcune classi studentesche.

Nel contesto universitario molti istituzionalizzano l'uso di ChatGPT o altri software, senza rendersi conto del travaso di conoscenza che verso queste piattaforme digitali stanno assicurando ad ogni prompt (richiesta) corredata dai dati di contesto da conferire per avere risposte sensate. Altri Atenei si dimostrano invece più prudenti, comunque nella totale assenza di linee guida significative. Eppure questo straordinario strumento è destinato a cambiare nel profondo il modo in cui formiamo le persone e come queste lavoreranno.

Per questo ritengo che sia urgente cercare di concepire un percorso possibile secondo cui lo straordinario strumento dell'IA generativa sia progettato e impiegato consapevolmente, nel rispetto dei diritti umani (privacy, libertà cognitiva), in modo scevro da pesanti interessi del mercato, garantendone un impiego il più diffuso possibile per evitare disuguaglianze sociali, sviluppandone hardware e software in modo diffuso per non dipendere, in un aspetto così fondamentale per una democrazia, da oligopoli e ingombranti dimensioni geopolitiche.

Alleati Digitali propone una soluzione

praticabile affinché l'IA e gli strumenti digitali cessino di essere adescatori di coscienze e divoratori di dati per scopi di mercato e diventino invece tutor socratici utili allo sviluppo cognitivo delle persone e al loro potenziamento. Un diverso sviluppo ed uso di queste tecnologie è possibile rispetto a quanto è accaduto finora e delle numerose implicazioni di questa nuova traiettoria parleremo nelle prossime pagine: dall'evitare pesanti interferenze col nostro sviluppo cerebrale alla progettazione pedagogica delle loro funzioni, dalla definizione di quali siano le architetture hardware e software di quella che sarà una sorta di IA personale che ci accompagnerà nei percorsi di apprendimento e sul lavoro, ai risvolti geopolitici di questa attesa rivoluzione, da come potrà mutare il quadro dei diritti fondamentali e della tutela giuridica degli individui, a come cambierà il mondo del lavoro, nelle imprese e nelle professioni.

Ma una forma di simbiosi rispettosa tra umano e intelligenza artificiale può esistere, specialmente in chiave formativa. Vi faccio un esempio

Immaginatevi me in una classe da 250 allievi. Succede ogni anno. Lezioni di Chimica.

Io faccio del mio meglio per far bene il mio lavoro ...Insegnare! ...ma come faccio a star dietro a ciascuno di loro? ...A rispondere a tutte le loro domande. ...A capire chi non chiede nulla ma non sta al passo con gli altri. ...A stargli, particolarmente vicino, per farlo recuperare.

Se ogni allievo avesse una sua intelligenza artificiale personale, quei loro alleati digitali, che seguono con loro le lezioni, conoscono i loro libri di testo, possono tranquillamente rispondere a dubbi, proporre loro approfondimenti, testarli mentre apprendono. Un grande aiuto. Nessuno rimane indietro.

...ma qui viene il bello.

posto se valgono quei presupposti.

Se ogni allievo scegliesse di mettere in collegamento con me il proprio alleato digitaleio avrei il polso della situazione dell'intera classe. Saprei con precisione a chi rivolgere attenzione ...umana, ...chi spronare anche con un po' di empatia. Ma potrei anche migliorare me stesso.

Come? Capendo, dalle loro reazioni biometriche, se quello che sto spiegando cattura la loro attenzione. Accende la loro volontà di apprendere. Li coinvolge.

...Praticamente in tempo reale.

Durante la presentazione del libro a Cuneo, lei ha usato un'immagine efficace: "possiamo usare l'IA come abbiamo usato la ruota per costruire una bicicletta". Il design ha storicamente il compito di dare forma a tecnologie complesse rendendole accessibili e significative. Secondo lei, cosa manca oggi nella cultura progettuale italiana per compiere questo passaggio con l'intelligenza artificiale?

Gli strumenti di intelligenza artificiale generativa possono proporre alternative a una nostra proposta iniziale. Sta poi a noi sfruttarle per meno per poi alla fine decidere quale è il design da proporre. Se l'obiettivo del design è quello di appagare esigenze dell'umano, che problema c'è se quel risultato del nostro genio è supportato dall'intelligenza artificiale? In fondo la vita di una qualsiasi opera di ingegno creativo non dipende da chi le ha dato corpo ma da chi la usa, la interpreta, la vive, magari a decenni o secoli di distanza dalla creazione.

Credo che si abbia oggi ancora un po' di pudore nell'ammettere che si fa uso dell'intelligenza artificiale. Pudore mal ri-

6.5 Il docente di design nell'era generativa: una professione da reinventare?

La trasformazione non riguarda solo chi apprende, ma anche chi insegna. Il docente è chiamato a evolvere da custode del sapere a facilitatore di processi ibridi, reinventando la propria autorità e i propri metodi di valutazione in un contesto in cui la conoscenza è sempre più distribuita e mediata dalle macchine.

6.5.1 La condizione attuale: volontà senza strumenti

Il quadro italiano restituisce un'immagine di incoerenza strategica particolarmente preoccupante. La Legge 132/2025, prima normativa organica italiana in materia di AI, ha stanziato cento milioni di euro per l'introduzione dell'intelligenza artificiale nelle scuole, finanziando 2.100 istituzioni scolastiche per la creazione di laboratori formativi e attività pratiche²⁸⁵. Questo investimento avviene tuttavia in un contesto in cui i docenti risultano largamente impreparati e privi di linee guida operative: il 77% dei docenti universitari concorda che l'IA generativa debba essere incorporata nel curriculum, l'88% chiede più formazione istituzionale sull'IA e l'85% auspica che gli studenti ricevano indicazioni chiare fin dall'inizio del percorso accademico. Le istituzioni, però, non hanno ancora elaborato strategie coerenti per rispondere a questa domanda: i docenti riportano confusione sulla pedagogia dell'IA, mancano strumenti pedagogici specifici e le politiche istituzionali restano inconsistenti²⁸⁶.

²⁸⁵ Orizzonte Scuola, *Intelligenza artificiale a scuola, ecco le 2.100 istituzioni scolastiche finanziate: pubblicato il Decreto*, 2026.

²⁸⁶ *Intelligenza artificiale nei licei italiani*, Libero, 2026.

²⁸⁷ *AI Competency Framework for Teachers*, UNESCO Digital Library, 2024.

²⁸⁸ *EduNews24, AI Index Report Stanford 2026: l'intelligenza artificiale conquista il mondo, ma l'Italia resta indietro sull'istruzione*, 2026.

²⁸⁹ *Gwo-Jen Hwang et al., Artificial Intelligence in Education: 26th International Conference, AIED 2025, 2025, Proceedings*, Springer Link, 2025.

Figura 6.9: Sei livelli di coinvolgimento creativo potenziato dall'intelligenza artificiale (#ppAI6), disposti in ordine progressivo di complessità e collaborazione.²⁹⁰ Grafico rielaborato.

Tabella 6.1: Analisi comparativa delle azioni del docente e delle funzionalità dei sistemi di IA mappate sui sei livelli di coinvolgimento del modello #ppAI6. ²⁹⁰ Tabella rielaborata.

Il ritardo italiano si inserisce in un quadro globale altrettanto deficitario. L'UNESCO certifica che nel 2022 solo sette Paesi al mondo avevano sviluppato un framework di competenze IA o un programma di sviluppo professionale per i docenti²⁸⁷. La causa principale risiede nella mancanza di chiarezza su come definire ruoli e competenze degli insegnanti nel contesto della crescente interazione umano-IA nelle pratiche educative. Mentre Paesi come la Cina hanno introdotto corsi di IA già dalla scuola primaria, con l'obiettivo di formare una generazione di cittadini AI-literate, e gli Emirati Arabi Uniti hanno costruito un ecosistema formativo che collega scuole, università e settore privato, l'Italia continua a procedere in modo frammentario, senza un disegno strategico capace di tradurre il finanziamento in competenza diffusa²⁸⁸.

Per il design, la situazione è particolarmente critica: solo il 9% dei docenti di Arts and Creative Technologies riporta confidenza nell'uso dell'IA, il dato più basso fra tutte le discipline. Il paradosso è dunque completo: la disciplina più esposta alla trasformazione generativa è quella meno attrezzata per affrontarla in aula²⁸⁹.

6.5.2 Dall'erogatore di contenuti al regista della collaborazione

L'IA ridefinisce il ruolo del docente di design in modo radicale. Il modello del peer-guided study offre una metafora concreta: il docente deve osservare, guidare e valutare il processo di collaborazione tra studente e IA, più che limitarsi a trasmettere contenuti. Le guide umane efficaci forniscono istruzioni step-by-step, stimolano la riflessione progettuale, suggeriscono strategie alternative e usano comunicazione multimodale; quando la guida sollecita feedback e riflessione sui design generati, il designer è obbligato a esplicitare le proprie intenzioni, e da quel momento guida e designer coordinano meglio le azioni successive. Il docente di design non è più chi mostra come si fa, ma chi insegna a dirigere un processo in cui la macchina genera e l'umano decide²⁹⁰.

Questa evoluzione del ruolo è descritta in modo sistematico dal framework #ppAI6, che articola le azioni specifiche del docente a ogni livello di engagement creativo con l'IA: dal fornire materiali AI-generated preconfezionati al livello più basso, all'assegnare task adattivi ai livelli intermedi, fino a riprogettare l'intero ambiente di apprendimento — per incoraggiare innovazione, pensiero critico e riflessione continua — al livello più alto. La pro-



Livello #ppAI6	Descrizione	Azioni del docente	Funzionalità dello strumento AI
<i>Livello 1:</i> Consumo Passivo	Gli studenti consumano passivamente contenuti generati dall'AI senza alcuna interazione.	Fornire materiali pre-progettati generati dall'AI (es. lezioni, quiz).	Sistemi di erogazione dei contenuti (es. testi automatizzati, video con sistema di raccomandazione).
<i>Livello 2:</i> Consumo Interattivo	Gli strumenti di AI rispondono agli input degli studenti, offrendo feedback personalizzati senza richiedere un contributo creativo.	Assegnare compiti che si adattano alle esigenze del singolo studente e monitorare i progressi.	Sistemi di tutoraggio intelligente (ITS).
<i>Livello 3:</i> Creazione di Contenuti Individuali	Gli studenti creano contenuti individualmente, con l'AI che supporta i loro processi creativi.	Incoraggiare la creazione indipendente (es. compiti di scrittura, di progettazione) e usare l'AI per aiutare a perfezionare o guidare le idee.	Strumenti di AI che forniscono suggerimenti, modelli o critiche per la creazione di contenuti (es. generazione di testi, strumenti creativi di problem-solving).
<i>Livello 4:</i> Coinvolgimento Collaborativo	Gli studenti collaborano con l'AI e con i coetanei per risolvere problemi o completare progetti.	Facilitare il lavoro di gruppo, co-progettare compiti di apprendimento e incoraggiare l'interazione tra pari.	Strumenti di co-creazione che consentono agli studenti di creare output di AI in collaborazione all'interno di un piccolo gruppo.
<i>Livello 5:</i> Co-Creazione Partecipativa	Gli studenti e l'AI creano in modo collaborativo nuove conoscenze o soluzioni.	Incoraggiare gli studenti a lavorare con l'AI per generare nuove idee, sviluppare progetti ed esplorare soluzioni.	Strumenti di co-creazione che consentono agli studenti di creare output di AI in collaborazione con diversi attori.
<i>Livello 6:</i> Apprendimento Trasformativo	Gli strumenti di AI consentono a studenti e docenti di co-creare conoscenza in modi dinamici e innovativi, promuovendo esperienze di apprendimento trasformativo.	Riprogettare l'ambiente di apprendimento per incoraggiare l'innovazione, facilitare il pensiero critico e sostenere una riflessione continua.	Sistemi di AI che si adattano al feedback dello studente in tempo reale, integrando le prospettive dei pari, dell'insegnante e dell'AI per una co-creazione e un apprendimento continui.

²⁹⁰ Xinyun Zhang et al., *From Consumption to Co-Creation: A Systematic Review of Six Levels of AI-Enhanced Creative Engagement in Education*, Multimodal Technologies and Interaction (MTI), MDPI, 2025.

²⁹¹ Juho Kahila, Henriikka Vartiainen et al., *Creative partnerships with generative AI: Possibilities for education and beyond*, International Journal of Educational Research / ScienceDirect, 2024.

gressione non riguarda solo lo studente ma anche il docente, che evolve da erogatore di contenuti a regista della collaborazione, da custode del sapere a facilitatore della scoperta²⁹⁰.

6.5.3 La formazione dei formatori come pre-condizione

Le istituzioni non possono introdurre l'IA nella didattica senza investire prima nella formazione dei docenti. È necessario creare centri o istituti dedicati all'educazione e alla ricerca sull'IA come hub di sviluppo professionale continuo. La formazione professionale dei docenti dovrebbe focalizzarsi su tre assi: comprendere come partneriare efficacemente con l'IA mantenendo la capacità degli studenti di engaging criticamente con i contenuti generati e preservando la loro agency creativa; sviluppare nuovi framework di assessment che riconoscano le creazioni ibride umano-IA; comprendere limitazioni, bias e implicazioni etiche degli strumenti²⁹¹.

Un modello operativo strutturato in questa direzione è offerto dall'UNESCO AI Competency Framework for Teachers (AI CFT), articolato in cinque aspetti – Human-centred mindset, Ethics of AI, AI foundations and applications, AI pedagogy, AI for professional development – e tre livelli di progressione: Acquire, Deepen, Create. Ogni blocco di competenze è corredato da curricular goals, learning objectives e contextual activities dettagliate. Il framework si definisce come master framework piuttosto che come blueprint prescrittivo, e raccomanda esplicitamente adattamenti nazionali e istituzionali²⁸⁷.

Aspetti	Progressione		
	Acquire (<i>acquisire</i>)	Deepen (<i>approfondire</i>)	Create (<i>creare</i>)
Human-centred mindset	Agency umana	Responsabilità umana	Responsabilità sociale
Ethics of AI	Principi etici	Uso sicuro e responsabile	Co-creazione di regole etiche
AI foundations and applications	Tecniche e applicazioni di base dell'AI	Competenze applicative	Creare con l'AI
AI pedagogy	Insegnamento assistito dall'AI	Integrazione tra AI e pedagogia	Trasformazione pedagogica potenziata dall'AI
AI for professional development	L'AI come abilitatore dell'apprendimento professionale continuo	L'AI per potenziare l'apprendimento organizzativo	L'AI per supportare la trasformazione professionale

Il modo in cui l'AI CFT articola la competenza etica è un buon esempio di come il framework traduca i tre livelli in pratica progressiva: ad Acquire si apprendono i principi etici, a Deepen si pratica un uso sicuro e responsabile, a Create si co-creano regole etiche insieme agli studenti. Questa progressione mostra che la competenza etica non è una conoscenza dichiarativa da apprendere una volta, ma una pratica in evoluzione che si rafforza attraverso l'esercizio. L'UNESCO insiste inoltre sull'empowerment di un uso human-accountable dell'IA: le responsabilità etiche e legali per la progettazione e l'uso dell'IA devono essere attribuite a individui, non diluite in sistemi o procedure²⁸⁷.

La vera sfida, tuttavia, non è la disponibilità di contenuti formativi, ma la costruzione dell'infrastruttura umana necessaria per trasformare quei contenuti in competenze operative. Le organizzazioni che si muoveranno più rapidamente nello sviluppo delle competenze IA saranno quelle che investono negli umani capaci di avvicinare gli altri alla comprensione, all'applicazione e infine alla capacità di dare forma al modo in cui l'IA funziona nel loro contesto specifico²⁷².

6.5.4 ChatGPT nella progettazione didattica: potenzialità e limiti

L'analisi dell'impatto dei sistemi generativi sul contesto educativo non può prescindere dall'esame di strumenti conversazionali specifici come ChatGPT, il cui impiego in questa sede viene assunto come paradigma rappresentativo dell'intera classe dei modelli linguistici computazionali (LLM) e delle loro evoluzioni correnti. L'esperienza di due docenti della University of Sydney, che collaborano con ChatGPT per progettare materiali post-graduate, è emblematica delle potenzialità e dei limiti dello strumento nella preparazione didattica. Entrambi scoprono che la competenza disciplinare pregressa è condizione necessaria per ottenere output utili dall'IA, ma che la collaborazione richiede competenze nuove. La sfida non è solo tecnologica ma identitaria: il docente che usa l'IA per progettare il syllabus rischia di omologare la formazione, esattamente come lo studente che la usa per progettare il prodotto rischia di omologare il design. La resistenza di alcuni docenti, dunque, non è solo tecnofobia: in alcuni casi è difesa legittima della specificità disciplinare²⁹².

La genericità degli output rappresenta una criticità confermata da evidenze convergenti: le piattaforme indicano cosa l'IA può fare, ma non aiutano a capire cosa fare

Tabella 6.2: Struttura del framework UNESCO per le competenze IA dei docenti, articolata in cinque aspetti e tre livelli di progressione.²⁸⁷
Tabella rielaborata.

²⁹² Imran Khan et al., *Artificial intelligence in design education: evaluating ChatGPT as a virtual colleague for post-graduate course development*, Design Science, Cambridge University Press, 2024.

con essa nel proprio ruolo specifico, nella propria organizzazione specifica, in questo momento specifico. Per il docente di design questa limitazione è particolarmente rilevante, perché il valore della didattica progettuale risiede nella specificità dei brief, dei contesti e delle sfide proposte agli studenti: un syllabus generico prodotto dall'IA tradisce proprio questa specificità. L'uso degli LLM per progettare materiali e syllabus richiede dunque non solo competenza disciplinare, ma anche la capacità di personalizzare, contestualizzare e arricchire gli output con la propria esperienza didattica²⁶⁹.

6.6 L'università come luogo in cui si progetta il futuro

In un'epoca ridefinita dall'avvento dell'intelligenza artificiale, l'università è chiamata a superare la tradizionale funzione di pura trasmissione del sapere per farsi laboratorio aperto e consapevole del domani. Di fronte a macchine capaci di elaborare contenuti a velocità inedite, l'istituzione universitaria ritrova la sua vocazione più autentica: un baricentro educativo e progettuale dove la tecnica torna a essere uno strumento al servizio della centralità umana.

6.6.1 La pedagogia semantica: formare creatori di significato

Abbiamo già incontrato, nel capitolo dedicato alla technoethics, il concetto di capitale semantico elaborato da Luciano Floridi: la capacità generativa di creare, curare e trasmettere significato a partire dalle informazioni che possediamo. Su questo fondamento filosofico l'autore stesso costruisce un ulteriore passaggio: la pedagogia semantica. Se il capitale semantico è la capacità generativa di trasformare informazione in significato, la pedagogia semantica è l'insieme delle pratiche formative che coltivano questa capacità. La formazione diventa cruciale non come produzione, trasmissione e manipolazione di contenuti, compiti in cui l'IA già supera le capacità umane, ma come coltivazione della capacità di generare

e valutare significati. Una pedagogia semantica non insegna cosa pensare, ma come il pensiero genera significato: sviluppa competenze non solo critiche ma anche creative, non solo analitiche ma anche sintetiche, per formare persone che non siano solo utenti ma anche progettisti, non solo seguaci ma anche cittadini²⁷⁸.

Questa pedagogia deve misurarsi con un paradosso costitutivo: preparare le nuove generazioni a un mondo imprevedibile usando strumenti che saranno obsoleti prima che gli studenti completino il percorso di studi. Per questo la risposta non può essere tecnica, non si tratta di insegnare l'ultimo strumento disponibile, ma deve essere semantica: sviluppare la capacità di leggere e scrivere ogni tipo di informazione, per saper creare, interpretare, criticare e migliorare il capitale semantico stesso, indipendentemente dagli strumenti che lo mediano in un dato momento storico²⁷⁸.

6.6.2 Progettare il futuro che vogliamo abitare

L'università non è solo il luogo in cui si apprendono competenze: è il luogo in cui si progetta la visione di futuro che vogliamo abitare. Tre direttrici emergono dalla convergenza delle evidenze analizzate in questo capitolo. La prima riguarda la costruzione di un'istituzione future-ready, che riconosca l'IA come core infrastructure e integri centri dedicati all'innovazione. La seconda punta a formare non consumatori ma custodi etici e artefici creativi dell'IA: non semplici utenti consapevoli, ma progettisti responsabili dell'impatto dei sistemi che utilizzano. La terza coltiva le competenze che l'IA non replica come nucleo irriducibile della professionalità del designer²⁸⁷.

L'evoluzione dei programmi di studio verso l'IA è inevitabile. I nuovi curricula, fondati su conoscenza, applicazioni e disposizioni, dovranno essere organizzati in ampiezza e profondità a partire da ciò che chi apprende dovrebbe sapere, non da ciò che chi insegna già conosce. L'IA diventerà essa stessa uno strumento per la pianificazione curricolare, il design dei programmi e la loro implementazione, allineandosi alle esigenze della società, degli studenti, del contesto locale e globale. Una maggiore integrazione con i Learning Management Systems offrirà a docenti e studenti un'esperienza educativa più unificata ed efficiente²⁹³. Resta un segnale da non sottovalutare: il senso di colpa nel creare con l'IA, documentato

²⁹³ Miguel Ángel Escotet, *The Future of Artificial Intelligence in Higher Education*, Escotet.org, 2025.

²⁹⁴ *Perché l'intelligenza artificiale fa sentire in colpa*, Domusweb, 2025.

nella ricerca sulla percezione degli studenti²⁹⁴, va letto non come un sintomo da correggere ma come una postura critica già acquisita: tale dilemma interiore non costituisce un limite, bensì la prova di una maturità progettuale che ricolloca l'essere umano al centro del processo creativo, rispondendo appieno agli obiettivi educativi auspicati in questa ricerca.

Conclusione

Il presente lavoro di tesi ha esplorato la profonda trasformazione che l'intelligenza artificiale sta imprimendo al settore del design, analizzandola attraverso sei prospettive interconnesse: storica, contestuale, normativa, etica, metodologica e formativa. L'interrogativo centrale che ha guidato la ricerca, in ogni capitolo, ha riguardato il mutamento del ruolo del progettista quando l'atto generativo non è più prerogativa esclusiva dell'essere umano.

La ricostruzione storica ha mostrato che la traiettoria dall'automazione del calcolo alla computazione simbolica non è una rottura improvvisa, ma la continuazione di un processo lungo due secoli, in cui ogni nuova tecnologia di rappresentazione ha spostato il confine tra ciò che è automatizzabile e ciò che resta prerogativa del giudizio umano, senza mai eliminarlo del tutto. L'analisi del panorama internazionale ha confermato che questa transizione non è né uniforme né neutrale: dietro la crescita esponenziale del mercato si muovono strategie geopolitiche e industriali profondamente diverse, dal potenziamento tecnico nordamericano alla governance europea, dall'infrastruttura digitale cinese all'integrazione verticale sudcoreana, ciascuna portatrice di una propria visione del rapporto tra tecnologia, lavoro e creatività. In questo contesto frammentato, la proliferazione di iniziative regolatorie non coordinate, a fronte di un impatto globale

dell'IA, rivela l'assenza di una cornice tecnica condivisa per la regolamentazione, che resta una questione aperta.

Per il progettista, tale scenario si traduce in una concreta incertezza riguardo alla responsabilità professionale. A questa incertezza risponde il cuore filosofico della ricerca, mostrando come l'opacità dei modelli generi due fenomeni speculari, l'*automation bias* e l'*automation aversion*, entrambi sintomi della stessa assenza di comprensione, e come la responsabilità rischi di disperdersi nel *many hands problem*. Contro questa dispersione, la tesi individua nel *meaningful human control* e nel *capital semantic*, inteso come la capacità squisitamente umana di trasformare l'informazione in significato, la condizione necessaria, ma non ancora sufficiente, per governare la pratica del progetto. Sebbene l'intelligenza artificiale sia ormai radicata in quest'ambito, il design si trova oggi su una soglia critica, in cui i vecchi paradigmi teorici si dimostrano obsoleti e i nuovi non hanno ancora assunto una forma definita. Attraversare attivamente questa soglia, invece di subirla, significa esattamente questo: restituire al designer la direzione del processo, rendendo il design, di nuovo, abitabile.

L'analisi della pratica progettuale ha tradotto queste premesse in termini operativi, rivelando che l'IA non semplifica il progetto, ma ne modifica la struttura. Il designer evolve da creator a director, confrontandosi con tensioni irrisolte quali il *deskilling*, i bias non intenzionali, l'omologazione estetica, l'atrofizzazione cognitiva e la disuguaglianza nell'accesso. Queste non sono mere imperfezioni tecniche superabili, ma condizioni strutturali della nuova relazione tra progettazione e automazione. Da qui emerge la necessità, approfondita nel capitolo conclusivo dedicato alla formazione, di un nuovo sistema di competenze che integri la padronanza operativa degli strumenti con capacità relazionali, critiche e cognitive. Le istituzioni formative, attualmente in ritardo rispetto alla velocità del cambiamento tecnologico, sono chiamate a costruire con urgenza questa nuova alfabetizzazione.

Attraverso le sei prospettive analizzate, la ricerca converge su un'unica direzione, confermata da interlocutori eterogenei: dalla visione del progetto generativo di Anne Asensio alla cultura open source di Massimo Banzi, dalla memoria d'impresa di Gaetano Di Tondo alla riflessione filosofica di Vera Tripodi e Maurizio Ferraris, fino alla centralità del giudizio nella pratica di Don Norman e alla prospettiva istituzionale di Guido Saracco. La parola

chiave che emerge è "giudizio", inteso come ciò che rimane propriamente umano nel progetto. Progettare, in questo nuovo contesto, significa esercitare un giudizio che non può essere isolato dal contesto normativo, culturale e tecnico. Il compito del design contemporaneo non è scegliere tra resistere all'IA o cederle il controllo, ma rendere nuovamente abitabile il proprio mestiere ai designer: riconquistare la direzione del processo, la responsabilità delle scelte e quell'immaginazione morale che nessun modello probabilistico potrà mai automatizzare, sostenuta da percorsi formativi che la coltivino esplicitamente come competenza.

Limiti della ricerca

Pur avendo mappato sistemi normativi che coprono Unione Europea, Nord America, America Latina, Asia-Pacifico, Africa e Medio Oriente, la maggior parte della letteratura consultata e tutti gli interlocutori intervistati appartengono al contesto occidentale, a causa della maggiore reperibilità di fonti e di contatti in quell'area linguistica e culturale.

A questa parzialità geografico-culturale si aggiunge un limite strutturale intrinseco all'oggetto stesso della ricerca: l'estrema e incessante velocità con cui evolve il fenomeno dell'Intelligenza Artificiale. Trattandosi di un ecosistema in continua mutazione, sia sul piano dell'avanzamento tecnologico sia su quello delle risposte regolatorie, la presente mappatura sconta l'inevitabile rischio di un'obsolescenza precoce. Lo scenario analizzato e ritenuto "attuale" nel momento della stesura di questo lavoro è, per sua natura, una fotografia istantanea di un soggetto in movimento, destinata a essere ridefinita dai rapidi sviluppi del mercato e della tecnica già nell'immediato futuro.

Infine, si evidenziano due ulteriori limiti metodologici. Il primo riguarda la natura non standardizzata delle interviste: le domande, costruite ad hoc sul profilo di ciascun interlocutore, hanno permesso di valorizzare la specificità di ogni voce ma hanno reso impossibile un confronto sistematico e simmetrico tra le risposte. Il secondo riguarda il profilo degli intervistati: tutti e sette occupano posizioni di vertice, accademico o professionale, mentre manca una voce "dal basso", quella di chi pratica il design quotidianamente senza una posizione di visibilità pubblica, e che probabilmente convive con la trasformazione generativa in termini più immediati e meno mediati.

Prospettive future

Questa tesi non si chiude con una metodologia da replicare né con una risposta definitiva da applicare. Quello che offre è, piuttosto, un'architettura interpretativa: l'attraversamento di un asse storico, normativo, etico, pratico e formativo, può costituire il punto di partenza per ricerche future più verticali, capaci di approfondire singolarmente ciascuno degli snodi affrontati con strumenti più mirati. Resta in ogni caso confermata l'intuizione che ha guidato questa tesi fin dal capitolo dedicato alla sua storia più remota: il rapporto tra uomo e macchina nel progetto non è mai stato, e non è oggi, una questione di sostituzione, ma di continua ridefinizione del confine tra l'automazione dei processi e l'intenzione progettuale.

bibliografia

1. M. Pasquinelli, *The Eye of the Master: A social history of artificial intelligence*, Verso Books, 2023. Disponibile all'indirizzo: <https://www.versobooks.com/products/2959-the-eye-of-the-master> (Consultato il 19-06-2026).
2. D. A. Grier, *When computers were human*, Princeton University Press, 2005. Disponibile all'indirizzo: <https://press.princeton.edu/books/paperback/9780691133829/when-computers-were-human> (Consultato il 19-06-2026).
3. Doron Swade, *The Cogwheel Brain: Charles Babbage and the Quest to Build the First Computer*, Little, Brown and Company, Londra, 2000. Disponibile all'indirizzo: https://books.google.com/books/about/The_Cogwheel_Brain.html?id=XcQgHAAACAAJ (Consultato il 19-06-2026).
4. J. Essinger, *Ada's algorithm: How Lord Byron's daughter Ada Lovelace launched the digital age*, Melville House, 2014. Disponibile all'indirizzo: <https://www.mhpbooks.com/books/adas-algorithm/> (Consultato il 19-06-2026).
5. A. Hyman, *Charles Babbage: Pioneer of the computer*, Oxford University Press, 1982. Disponibile all'indirizzo: <https://global.oup.com/academic/product/charles-babbage-9780198581703> (Consultato il 19-06-2026).
6. B. J. Copeland, *The Essential Turing*, Oxford University Press, 2004. Disponibile all'indirizzo: <https://global.oup.com/academic/product/the-essential-turing-9780198250791> (Consultato il 19-06-2026).
7. M. A. Boden, *AI: Its nature and future*, Oxford University Press, 2016. Disponibile all'indirizzo: <https://global.oup.com/academic/product/ai-its-nature-and-future-9780198777984> (Consultato il 19-06-2026).
8. S. J. Russell, P. Norvig, *Artificial Intelligence: A modern approach* (4th ed.), Pearson, 2021. Disponibile all'indirizzo: <https://aima.cs.berkeley.edu/> (Consultato il 19-06-2026).
9. E. R. Kandel et al., *Principles of neural science* (5th ed.), McGraw-Hill, 2013. Disponibile all'indirizzo: <https://www.mhprofessional.com/principles-of-neural-science-fifth-edition-9780071390118-usa> (Consultato il 19-06-2026).
10. W. S. McCulloch, W. Pitts, A logical calculus of the ideas immanent in nervous activity, in *The Bulletin of Mathematical Biophysics*, 5(4), pp. 115–133, 1943. Disponibile all'indirizzo: <https://link.springer.com/article/10.1007/BF02478259> (Consultato il 19-06-2026).
11. A. Pickering, *The cybernetic brain: Sketches of another future*, University of Chicago Press, 2010. Disponibile all'indirizzo: <https://press.uchicago.edu/ucp/books/book/chicago/C/bo8169881.html> (Consultato il 19-06-2026).
12. N. Wiener, *Cybernetics: Or control and communication in the animal and the machine*, MIT Press, 1948. Disponibile all'indirizzo: <https://mitpress.mit.edu/9780262730099/cybernetics/> (Consultato il 19-06-2026).
13. J. McCarthy, M. L. Minsky, N. Rochester, C. E. Shannon, A proposal for the Dartmouth summer research project on artificial intelligence, 1956. Disponibile all'indirizzo: <https://ojs.aaai.org/aimagazine/index.php/>

aimagazine/article/view/1904 (Consultato il 19-06-2026).

still-cant-do/ (Consultato il 19-06-2026).

14. A. Newell, H. A. Simon, Computer science as empirical inquiry: Symbols and search, in *Communications of the ACM*, 19(3), pp. 113–126, 1976. Disponibile all'indirizzo: <https://dl.acm.org/doi/10.1145/360018.360022> (Consultato il 19-06-2026).
15. H. Gardner, *The mind's new science: A history of the cognitive revolution*, Basic Books, 1985. Disponibile all'indirizzo: <https://www.basic-books.com/titles/howard-gardner/the-minds-new-science/9780465046355/> (Consultato il 19-06-2026).
16. ALPAC, *Language and machines: Computers in translation and linguistics*, National Academy of Sciences, National Research Council, 1966. Disponibile all'indirizzo: <https://nap.nationalacademies.org/catalog/20813/language-and-machines-computers-in-translation-and-linguistics> (Consultato il 19-06-2026).
17. D. Crevier, *AI: The tumultuous history of the search for artificial intelligence*, Basic Books, 1993. Disponibile all'indirizzo: <https://www.basic-books.com/titles/daniel-crevier/ai/9780465029976/> (Consultato il 19-06-2026).
18. M. Polanyi, *The tacit dimension*, Doubleday & Co., 1966. Disponibile all'indirizzo: <https://press.uchicago.edu/ucp/books/book/chicago/T/bo6035368.html> (Consultato il 19-06-2026).
19. H. L. Dreyfus, *What computers still can't do: A critique of artificial reason*, MIT Press, 1992. Disponibile all'indirizzo: <https://mitpress.mit.edu/9780262540674/what-computers-still-cant-do/> (Consultato il 19-06-2026).
20. D. Harvey, *A brief history of neoliberalism*, Oxford University Press, 2005. Disponibile all'indirizzo: <https://global.oup.com/academic/product/a-brief-history-of-neoliberalism-9780199283279> (Consultato il 19-06-2026).
21. J. Lighthill, *Artificial intelligence: A general survey*, Science Research Council, 1973. Disponibile all'indirizzo: https://www.chilton-computing.org.uk/inf/literature/reports/lighthill_report/p001.htm (Consultato il 19-06-2026).
22. E. A. Feigenbaum, *Expert systems in the 1980s*, in *State of the Art Report on Machine Intelligence*, Pergamon Infotech, 1981. Disponibile all'indirizzo: <https://stacks.stanford.edu/file/druid%3Avf069sz9374/vf069sz9374.pdf> (Consultato il 19-06-2026).
23. M. M. Waldrop, *Complexity: The emerging science at the edge of order and chaos*, Simon & Schuster, 1992. Disponibile all'indirizzo: <https://www.simonandschuster.com/books/Complexity/M-Mitchell-Waldrop/9780671872342> (Consultato il 19-06-2026).
24. W. D. Nordhaus, *Two centuries of productivity growth in computing*, in *The Journal of Economic History*, 67(1), pp. 128–159, 2007. Disponibile all'indirizzo: <https://www.cambridge.org/core/journals/journal-of-economic-history/article/two-centuries-of-productivity-growth-in-computing/856EC5947A5857296D3328FA154-BA3A3> (Consultato il 19-06-2026).
25. J. A. Fodor, Z. W. Pylyshyn, *Connectionism and cognitive architecture: A critical analysis*, in *Cognition*, 28(1-

- 2), pp. 3–71, 1988. Disponibile all'indirizzo: <https://www.sciencedirect.com/science/article/pii/S0010027788900315> (Consultato il 19-06-2026).
26. R. Sun, Robust reasoning: Integrating rule-based and similarity-based reasoning, in *Artificial Intelligence*, 75(2), pp. 241–295, 1995. Disponibile all'indirizzo: <https://www.sciencedirect.com/science/article/pii/S0004370295900327> (Consultato il 19-06-2026).
27. P. Domingos, A few useful things to know about machine learning, in *Communications of the ACM*, 55(10), pp. 78–87, 2012. Disponibile all'indirizzo: <https://dl.acm.org/doi/10.1145/2347736.2347755> (Consultato il 19-06-2026).
28. T. M. Mitchell, *Machine learning*, McGraw-Hill, 1997. Disponibile all'indirizzo: <https://www.cs.cmu.edu/~tom/mlbook.html> (Consultato il 19-06-2026).
29. E. Alpaydin, *Introduction to machine learning* (2nd ed.), MIT Press, 2010. Disponibile all'indirizzo: <https://mitpress.mit.edu/9780262012430/introduction-to-machine-learning/> (Consultato il 19-06-2026).
30. J. Pearl, *Probabilistic reasoning in intelligent systems: Networks of plausible inference*, Morgan Kaufmann, 1988. Disponibile all'indirizzo: <https://www.sciencedirect.com/book/9780080514895/probabilistic-reasoning-in-intelligent-systems> (Consultato il 19-06-2026).
31. M. Campbell, A. J. Hoane, F. H. Hsu, Deep Blue, in *Artificial Intelligence*, 134(1-2), pp. 57–83, 2002. Disponibile all'indirizzo: <https://www.sciencedirect.com/science/article/pii/S0004370201001291> (Consultato il 19-06-2026).
32. N. Srnicek, *Platform capitalism*, Polity Press, 2017. Disponibile all'indirizzo: https://www.politybooks.com/bookdetail?book_slug=platform-capitalism--9781509504862 (Consultato il 19-06-2026).
33. A. Halevy, P. Norvig, F. Pereira, The unreasonable effectiveness of data, in *IEEE Intelligent Systems*, 24(2), pp. 8–12, 2009. Disponibile all'indirizzo: <https://ieeexplore.ieee.org/document/4804817> (Consultato il 19-06-2026).
34. Tony Jebara, *Machine Learning: Discriminative and Generative*, Springer Science+Business Media, New York, 2012. Disponibile all'indirizzo: https://books.google.com/books/about/Machine_Learning.html?id=g5rSBwAAQBAJ (Consultato il 10-08-2025).
35. A. Krizhevsky, I. Sutskever, G. E. Hinton, ImageNet classification with deep convolutional neural networks, in *Advances in Neural Information Processing Systems*, 25, pp. 1097–1105, 2012. Disponibile all'indirizzo: <https://papers.nips.cc/paper/2012/hash/399862d3b9d6b76c8436e924a68c45b-Abstract.html> (Consultato il 19-06-2026).
36. R. J. Brachman, H. J. Levesque, *Knowledge representation and reasoning*, Elsevier/Morgan Kaufmann, 2004. Disponibile all'indirizzo: <https://www.sciencedirect.com/book/9781558609327/knowledge-representation-and-reasoning> (Consultato il 19-06-2026).
37. A. Vaswani et al., Attention is all you need, in *Advances in Neural Information Processing Systems*, 30, 2017. Disponibile all'indirizzo: <https://arxiv.org/abs/1706.03762> (Consultato il 19-06-2026).

/arxiv.org/abs/1706.03762 (Consultato il 19-06-2026).

38. J. Devlin, M. W. Chang, K. Lee, K. Toutanova, BERT: Pre-training of deep bidirectional transformers for language understanding, in Proceedings of NAACL-HLT 2019, pp. 4171–4186, 2019. Disponibile all'indirizzo: <https://arxiv.org/abs/1810.04805> (Consultato il 19-06-2026).
39. T. Brown et al., Language models are few-shot learners, in Advances in Neural Information Processing Systems, 33, pp. 1877–1901, 2020. Disponibile all'indirizzo: <https://arxiv.org/abs/2005.14165> (Consultato il 19-06-2026).
40. R. Bommasani et al., On the opportunities and risks of foundation models, arXiv preprint arXiv:2108.07258, 2021. Disponibile all'indirizzo: <https://arxiv.org/abs/2108.07258> (Consultato il 19-06-2026).
41. B. M. Lake, T. D. Ullman, J. B. Tenenbaum, S. J. Gershman, Building machines that learn and think like people, in Behavioral and Brain Sciences, 40, e253, 2017. Disponibile all'indirizzo: <https://www.cambridge.org/core/journals/behavioral-and-brain-sciences/article/building-machines-that-learn-and-think-like-people/9A9535B1D745A0377E16C590E14B94993> (Consultato il 19-06-2026).
42. A. D. Garcez et al., Neural-symbolic computing: An effective methodology for informed machine learning, in Pattern Recognition Letters, 125, pp. 191–204, 2019. Disponibile all'indirizzo: <https://www.sciencedirect.com/science/article/pii/S0167865519301606> (Consultato il 19-06-2026).
43. E. M. Bender, T. Gebru, A. McMillan-Major, S. Shmitchell, 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, 2021. Disponibile all'indirizzo: <https://dl.acm.org/doi/10.1145/3442188.3445922> (Consultato il 19-06-2026).
44. A. D. Garcez, L. C. Lamb, Neuro-symbolic AI: The 3rd wave, arXiv preprint arXiv:2012.05876, 2020. Disponibile all'indirizzo: <https://arxiv.org/abs/2012.05876> (Consultato il 19-06-2026).
45. M. Wooldridge, An introduction to multiagent systems (2nd ed.), John Wiley & Sons, 2009. Disponibile all'indirizzo: <https://www.wiley.com/en-us/An+Introduction+to+MultiAgent+Systems%2C+2nd+Edition-p-9780470519462> (Consultato il 19-06-2026).
46. M. Bratman, Intention, plans, and practical reason, Harvard University Press, 1987. Disponibile all'indirizzo: <https://www.hup.harvard.edu/catalog.php?isbn=9780674458185> (Consultato il 19-06-2026).
47. D. Amodei, C. Olah, J. Steinhardt, P. Christiano, J. Schulman, D. Mané, Concrete problems in AI safety, arXiv preprint arXiv:1606.06565, 2016. Disponibile all'indirizzo: <https://arxiv.org/abs/1606.06565> (Consultato il 19-06-2026).
48. F. Doshi-Velez, B. Kim, Towards a rigorous science of interpretable machine learning, arXiv preprint arXiv:1702.08608, 2017. Disponibile all'indirizzo: <https://arxiv.org/abs/1702.08608> (Consultato il 19-06-2026).
49. E. Brynjolfsson, A. McAfee, The se-

- cond machine age: Work, progress, and prosperity in a time of brilliant technologies, W. W. Norton & Company, 2014. Disponibile all'indirizzo: <https://www.norton.com/books/the-second-machine-age/> (Consultato il 19-06-2026).
50. G. Marcus, The next decade in AI: Four steps towards robust artificial intelligence, arXiv preprint arXiv:2002.06177, 2020. Disponibile all'indirizzo: <https://arxiv.org/abs/2002.06177> (Consultato il 19-06-2026).
 51. H. A. Simon, The sciences of the artificial (3rd ed.), MIT Press, 1996.
 52. IBM, AI product design: The future of intelligent experiences, n.d. Disponibile all'indirizzo: <https://www.ibm.com/it-it/think/topics/ai-product-design> (Consultato il 19-06-2026).
 53. A. I. Miller, The artist in the machine: The world of AI-powered creativity, MIT Press, 2019.
 54. Artnet News, A brief history of AI art, from its humble 1960s beginnings to the present day, 2021. Disponibile all'indirizzo: <https://news.artnet.com/art-world/artificial-intelligence-art-history-2045520> (Consultato il 19-06-2026).
 55. R. M. Auer, D. Elflein, Artificial intelligence: Intelligent art, human-machine interaction and creative practice, Transcript Verlag, 2018.
 56. F. Nake, Computer art: A personal recollection, in Proceedings of the 5th Conference on Creativity & Cognition, pp. 54–62, 2005. Disponibile all'indirizzo: <https://doi.org/10.1145/1056224.1056234> (Consultato il 19-06-2026).
 57. W. Hosh, Ivan Sutherland: American computer scientist, Encyclopaedia Britannica, 2009. Disponibile all'indirizzo: <https://www.britannica.com/biography/Ivan-Sutherland> (Consultato il 19-06-2026).
 58. I. E. Sutherland, Sketchpad: A man-machine graphical communication system (Technical Report No. 574), University of Cambridge Computer Laboratory, 2003. Disponibile all'indirizzo: <https://www.cl.cam.ac.uk/techreports/UCAM-CL-TR-574.pdf> (Consultato il 19-06-2026).
 59. Parametric Architecture, The evolution of computational design: From algorithms to AI, n.d. Disponibile all'indirizzo: <https://parametric-architecture.com/the-evolution-of-computational-design-from-algorithms-to-ai/> (Consultato il 19-06-2026).
 60. H. Cohen, What is an image?, in Proceedings of the 6th International Joint Conference on Artificial Intelligence, Tokyo, Japan, 1979. Disponibile all'indirizzo: <https://aaronshome.com/aaron/publications/whatisanimage.pdf> (Consultato il 19-06-2026).
 61. Artribune, Arte e intelligenza artificiale: Storia e sviluppi, 2023. Disponibile all'indirizzo: <https://www.artribune.com/progettazione/new-media/2023/06/arte-intelligenza-artificiale-storia/> (Consultato il 19-06-2026).
 62. Dassault Systèmes, La nostra storia, n.d. Disponibile all'indirizzo: <https://www.3ds.com/it/about/company/history> (Consultato il 19-06-2026).
 63. D. E. Weisberg, Autodesk and AutoCAD: A history of CAD innovation, Shapr3D, 2023. Disponibile all'indirizzo: <https://www.shapr3d.com/history-of-cad/autodesk-and-autocad> (Consultato il 19-06-2026).

64. Robots Authority, The rise and fall of expert systems in the 1980s, 2025. Disponibile all'indirizzo: <https://robotsauthority.com/the-rise-and-fall-of-expert-systems-in-the-1980s/> (Consultato il 19-06-2026).
65. O. Sigmund, Topology optimization: A tool for the tailoring of structures and materials, Technical University of Denmark, 2000. Disponibile all'indirizzo: <https://orbit.dtu.dk/en/publications/topology-optimization-a-tool-for-the-tailoring-of-structures-and-/> (Consultato il 19-06-2026).
66. J. Rhee, H. Jung, T. Kim, Three decades of machine learning with neural networks in computer-aided architectural design (1990–2021), in Design Science, 9(e15), 2023. Disponibile all'indirizzo: <https://doi.org/10.1017/dsj.2023.15> (Consultato il 19-06-2026).
67. R. Coyne, Artificial intelligence in design, Elsevier Science, 2003. Disponibile all'indirizzo: <https://doi.org/10.1016/pii/095418109090031X> (Consultato il 19-06-2026).
68. A. Garg, S. Patil, Grasshopper algorithmic modelling: Parametric design for product platform customisation, ResearchGate, 2025. Disponibile all'indirizzo: <https://www.researchgate.net/publication/392324540> (Consultato il 19-06-2026).
69. Grasshopper3D, Grasshopper: Algorithmic modeling for Rhino, 2007. Disponibile all'indirizzo: <https://www.grasshopper3d.com/> (Consultato il 19-06-2026).
70. P. Schumacher, Parametricism: A new global style for architecture and urban design, Patrik Schumacher Website, 2008. Disponibile all'indirizzo: <https://patrikschumacher.com/parametricism-a-new-global-style-for-architecture-and-urban-design/> (Consultato il 19-06-2026).
71. DMG MORI, Intelligenza artificiale nella produzione, n.d. Disponibile all'indirizzo: <https://it.dmgmori.com/novita-e-media/blog-e-stories/blog/blg24-17-intelligenza-artificiale-produzione> (Consultato il 19-06-2026).
72. J. Lee, The evolution of AI in design: From CAD to generative models, Number Analytics Blog, 2025. Disponibile all'indirizzo: <https://www.numberanalytics.com/blog/evolution-of-ai-in-design> (Consultato il 19-06-2026).
73. Prisma Tech, Come l'intelligenza artificiale generativa sta cambiando il mondo del design, n.d. Disponibile all'indirizzo: <https://www.prisma-tech.it/news/come-lintelligenza-artificiale-generativa-sta-cambiando-il-mondo-del-design> (Consultato il 19-06-2026).
74. K. Chen, Y. Wang, L. Zhang, How generative AI supports human in conceptual design, in Design Science, 11(e22), 2025. Disponibile all'indirizzo: <https://doi.org/10.1017/dsj.2025.22> (Consultato il 19-06-2026).
75. M. Ozsoy, The impact of generative AI on creative design processes, DergiPark, 2025. Disponibile all'indirizzo: <https://dergipark.org.tr/tr/download/article-file/4387905> (Consultato il 19-06-2026).
76. Domus, The evolution of AI artists: From AARON to Botto, 2025. Disponibile all'indirizzo: <https://www.domusweb.it/en/art/2025/01/24/the-evolution-of-ai-artists.html> (Consultato il 19-06-2026).
77. Officina dei Segni, L'intelligenza artificiale nel design: Collaborazione

- creativa o minaccia?, n.d. Disponibile all'indirizzo: <https://www.officinadisegni.it/intelligenza-artificiale-nel-design-collaborazione-creativa-o-minaccia/> (Consultato il 19-06-2026).
78. The Business Research Company, Artificial Intelligence (AI) In Industrial Design Global Market Report, 2026. Disponibile all'indirizzo: <https://www.thebusinessresearchcompany.com/report/artificial-intelligence-ai-in-industrial-design-global-market-report> (Consultato il 14-04-2026)
 79. Precedence Research, Generative AI in Industrial Design Market, 2025. Disponibile all'indirizzo: <https://www.precedenceresearch.com/generative-ai-in-industrial-design-market> (Consultato il 14-04-2026)
 80. Yahoo Finance, Industrial Design Market Report 2026, 2026. Disponibile all'indirizzo: <https://finance.yahoo.com/news/industrial-design-market-report-2026-095000776.html> (Consultato il 14-04-2026)
 81. Fortune Business Insights, Generative AI in Product Design & Engineering Market, 2026. Disponibile all'indirizzo: <https://www.fortunebusinessinsights.com/generative-ai-in-product-design-engineering-market-115752> (Consultato il 14-04-2026)
 82. Autodesk, Autodesk Fusion 2026 / Neural CAD, Prisma Tech, 2026. Disponibile all'indirizzo: <https://www.prisma-tech.it/news/autodesk-fusion-2026> (Consultato il 14-04-2026)
 83. NVIDIA + Synopsys, Synopsys Showcases NVIDIA Partnership Impact and Ecosystem Innovation at GTC 2026, PR Newswire, 2026. Disponibile all'indirizzo: <https://www.prnewswire.com/news-releases/synopsys-showcases-nvidia-partnership-impact-and-ecosystem-innovation-at-gtc-2026-302715123.html> (Consultato il 14-04-2026)
 84. Ansys, Ansys 2025 R2: Engineering Copilot, 2025. Disponibile all'indirizzo: <https://youtu.be/FV69J1yp0B8> (Consultato il 14-04-2026)
 85. Ansys, Ansys 2025 R2 Accelerates Innovation with AI, 2025. Disponibile all'indirizzo: <https://www.ansys.com/news-center/press-releases/7-29-25-r2-accelerates-innovation-with-ai> (Consultato il 14-04-2026)
 86. Adobe, Adobe Firefly Delivers Groundbreaking AI Audio, Video and Imaging Innovations and New Models in All-In-One Creative AI Studio, 2025. Disponibile all'indirizzo: <https://news.adobe.com/news/2025/10/adobe-max-2025-firefly> (Consultato il 14-04-2026)
 87. Adobe, Adobe and NVIDIA Announce Strategic Partnership to Deliver the Next Generation of Firefly Models and Creative, Marketing and Agentic Workflows, 2026. Disponibile all'indirizzo: <https://news.adobe.com/news/2026/03/adobe-and-nvidia-announce-strategic-partnership> (Consultato il 14-04-2026)
 88. Figma, Your creativity, unblocked with Figma AI, 2025. Disponibile all'indirizzo: <https://www.figma.com/ai/config-2025/> (Consultato il 14-04-2026)
 89. Engineering.com, Leveraging simulation-driven design with PTC Creo, 2025. Disponibile all'indirizzo: <https://www.engineering.com/leveraging-simulation-driven-design-with-ptc-creo/> (Consultato il 14-04-2026)
 90. Tristar Digital Thread Solutions, PTC Creo AI Optimization: How to Make

- Faster Design Decisions, 2026. Disponibile all'indirizzo: <https://www.tristar.com/blog/creo-ai/> (Consultato il 14-04-2026)
91. Dassault Systèmes, IA e innovazione protagoniste al 3DEXPERIENCE World 2026, Tecnelab, 2026. Disponibile all'indirizzo: <https://www.tecnelab.it/news/attualita/ia-e-innovazione-protagoniste-al-3dexperience-world-2026-di-dassault> (Consultato il 14-04-2026)
92. Siemens, Siemens EDA AI DAC 2025, 2025. Disponibile all'indirizzo: <https://news.siemens.com/it-it/siemens-eda-ai-dac-2025/> (Consultato il 14-04-2026)
93. Schaeffler, Schaeffler at CES 2026: Technologies That Move the Future, 2026. Disponibile all'indirizzo: <https://schaeffler-tomorrow.com/en/article/schaeffler-at-ces-2026-technologies-that-move-the-future> (Consultato il 14-04-2026)
94. Baidu, ERNIE 5.0: Unified Multimodal Model, 2026. Disponibile all'indirizzo: <https://ernie.baidu.com/blog/posts/ernie5.0/> (Consultato il 14-04-2026)
95. Tech in Asia, DeepSeek and Tencent researchers launch AI tool for CAD design, 2026. Disponibile all'indirizzo: <https://www.techinasia.com/news/deepseek-tencent-researchers-launch-ai-tool-for-cad-design> (Consultato il 14-04-2026)
96. Enki AI, Samsung AI in Manufacturing 2025: How the NVIDIA Partnership Drives Industrial Efficiency, 2026. Disponibile all'indirizzo: <https://enkiai.com/ai-market-intelligence/samsung-ai-in-manufacturing-2025-nvidia-drives-efficiency> (Consultato il 14-04-2026)
97. Stanford University, Institute for Human-Centered Artificial Intelligence (HAI), 2026. Disponibile all'indirizzo: <https://hai.stanford.edu/> (Consultato il 14-04-2026)
98. Stanford HAI, AI Index Report 2025, Stanford University, 2025. Disponibile all'indirizzo: <https://hai.stanford.edu/ai-index/2025-ai-index-report> (Consultato il 14-04-2026)
99. Stanford University, Stanford AI experts predict what will happen in 2026, 2025. Disponibile all'indirizzo: <https://news.stanford.edu/stories/2025/12/stanford-ai-experts-predict-what-will-happen-in-2026> (Consultato il 14-04-2026)
100. Stanford HAI, Stanford scholars train generative AI to be better creative collaborators, HAI News, 2025. Disponibile all'indirizzo: <https://hai.stanford.edu/news/stanford-scholars-train-generative-ai-to-be-better-creative-collaborators> (Consultato il 14-04-2026)
101. Stanford University, Norms in the Age of Intelligent Machines: Bodies, Knowledge, Governmentality, 2025. Disponibile all'indirizzo: <https://art.stanford.edu/norms-age-intelligent-machines-bodies-knowledge-governmentality> (Consultato il 14-04-2026)
102. Stanford University, How artists can 'interrupt' artificial intelligence, 2025. Disponibile all'indirizzo: <https://art.stanford.edu/news/how-artists-can-interrupt-artificial-intelligence> (Consultato il 14-04-2026)
103. Stanford HAI, From Privacy to 'Glass Box' AI, Stanford Students Are Targeting Real-World Problems, HAI News, 2025. Disponibile all'indirizzo: <https://hai.stanford.edu/news/from-privacy-to-glass-box-ai-stan>

- ford-students-are-targeting-real-world-problems (Consultato il 14-04-2026)
104. MIT News, Robot, make me a chair: sistema di progettazione AI-driven, 2025. Disponibile all'indirizzo: <https://news.mit.edu/2025/robot-makes-chair-1216> (Consultato il 14-04-2026)
 105. MIT Media Lab, DESAI25: Workshop on Generative AI for Design, 2025. Disponibile all'indirizzo: <https://www.media.mit.edu/events/generative-ai-for-design-workshop-2025-desai25/> (Consultato il 14-04-2026)
 106. University of California Berkeley, As chatbots get smarter, humans' unique language abilities are becoming less special, 2025. Disponibile all'indirizzo: <https://news.berkeley.edu/2025/07/14/as-chatbots-get-smarter-humans-unique-language-abilities-are-becoming-less-special/> (Consultato il 14-04-2026)
 107. University of California Berkeley, 11 things UC Berkeley AI experts are watching for in 2026, 2026. Disponibile all'indirizzo: <https://news.berkeley.edu/2026/01/13/what-uc-berkeley-ai-experts-are-watching-for-in-2026/> (Consultato il 14-04-2026)
 108. University of California San Diego, Human-centered eXtended Intelligence Research Lab (HXI), 2026. Disponibile all'indirizzo: <https://hxi.ucsd.edu/> (Consultato il 14-04-2026)
 109. Tsinghua University, Tsinghua University releases comprehensive guiding principles for AI use in education, 2025. Disponibile all'indirizzo: <https://www.tsinghua.edu.cn/en/info/1245/14598.htm> (Consultato il 14-04-2026)
 110. Carnegie Endowment for International Peace, Making Sense of China's AISI Equivalent: Institutional Design and Key Actors, 2025. Disponibile all'indirizzo: <https://carnegieendowment.org/research/2025/06/how-some-of-chinas-top-ai-thinkers-built-their-own-ai-safety-institute> (Consultato il 14-04-2026)
 111. Tongji University, Design AI Lab, 2026. Disponibile all'indirizzo: <https://www.oreateai.com/blog/where-design-meets-the-digital-frontier-unpacking-the-design-ai-lab/88-da68bcc3497a4c413b808c2638-dc7b> (Consultato il 14-04-2026)
 112. University of Oxford, Oxford becomes first UK university to offer ChatGPT Edu to all staff and students, 2025. Disponibile all'indirizzo: <https://www.ox.ac.uk/news/2025-09-19-oxford-becomes-first-uk-university-offer-chatgpt-edu-all-staff-and-students> (Consultato il 14-04-2026)
 113. University of Cambridge, Using Generative AI: Guidance for Students, Cambridge Blended Learning, 2025. Disponibile all'indirizzo: <https://blendedlearning.cam.ac.uk/artificial-intelligence-and-education/using-generative-ai> (Consultato il 14-04-2026)
 114. University of Cambridge, Ricerca sull'integrazione dell'AI nell'istruzione human-centrica: AI literacy e curriculum, 2025. Disponibile all'indirizzo: <https://blendedlearning.cam.ac.uk/artificial-intelligence-and-education/using-generative-ai> (Consultato il 14-04-2026)
 115. University of Cambridge, Can AI be a good creative partner?, 2025. Disponibile all'indirizzo: <https://www.cam.ac.uk/research/news/can-ai-be-a-good-creative-partner> (Consultato il 14-04-2026)

116. Royal College of Art, Leading with a Design Mindset, 2026. Disponibile all'indirizzo: <https://www.rca.ac.uk/study/programme-finder/leading-with-a-design-mindset/> (Consultato il 14-04-2026)
117. Royal College of Art, Strategic Plan 2026–2030, 2026. Disponibile all'indirizzo: <https://www.rca.ac.uk/more/organisation/corporate-publications/strategic-plan-2026-30/> (Consultato il 14-04-2026)
118. TU Delft, Design at Scale Lab — Human-AI Collaboration in Design for Social Good, TU Delft AI Labs programme, 2026. Disponibile all'indirizzo: <https://www.tudelft.nl/en/ai/design-at-scale-lab> (Consultato il 14-04-2026)
119. TU Delft, AiDAPT — AI for a sustainable and resilient built environment, TU Delft AI Labs programme, 2026. Disponibile all'indirizzo: <https://www.tudelft.nl/en/ai/aidapt> (Consultato il 14-04-2026)
120. ETH Zurich, Design++ — Center for Augmented Computational Design in Architecture, Engineering and Construction, 2026. Disponibile all'indirizzo: <https://designplusplus.ethz.ch/> (Consultato il 14-04-2026)
121. Technical University of Munich, TUM GNI International Symposium 2026: Artificial Intelligence for the Built World, Georg Nemetschek Institute, 2026. Disponibile all'indirizzo: <https://events.gni.tum.de/ai-symposium-2026/> (Consultato il 14-04-2026)
122. Aalto University, Aalto University in 2025: strong results in research, education and innovation activities, 2026. Disponibile all'indirizzo: <https://www.aalto.fi/en/news/aalto-university-in-2025-strong-results-in-research-education-and-innovation-activities> (Consultato il 14-04-2026)
123. Aalto University / Aalto Executive Education, Generative AI in Creating Visual Content, corso di formazione continua, 2025. Disponibile all'indirizzo: <https://aalto.fi/en/lifewide-learning-courses-and-programmes/generative-ai-in-creating-visual-content> (Consultato il 14-04-2026)
124. The Investor / Herald Corporation, South Korea's KAIST to launch stand-alone AI college in 2026 amid global talent race, 2025. Disponibile all'indirizzo: <https://www.theinvestor.co.kr/article/10635648> (Consultato il 14-04-2026)
125. KAIST News Center, Korea's AI Research Hub: Global AI Frontier Symposium 2025, 2025. Disponibile all'indirizzo: https://www.kaist.ac.kr/news/en/html/news/?mng_no=54570&mode=V (Consultato il 14-04-2026)
126. K-Moonshot, Seoul National University (SNU): AI Research & Strategy, 2026. Disponibile all'indirizzo: <https://kmoonshot.com/ecosystem/research/snu/> (Consultato il 14-04-2026)
127. StartSe / BNamericas, USP lança o maior cluster de IA da América Latina e acelera a pesquisa no Brasil, 2026. Disponibile all'indirizzo: <https://www.startse.com/artigos/usp-lanca-o-maior-cluster-de-ia-da-america-latina-e-acelera-a-pesquisa-no-brasil/> (Consultato il 14-04-2026)
128. Universidad de Buenos Aires — FADU / Instituto de Arte Americano, Seminario abierto IAA — MAHCA-DU: El Uso de las IA en los Procesos de Investigación Urbano-Arquitectó-

- nicos, 2026. Disponibile all'indirizzo: <https://www.iaa.fadu.uba.ar/tag/intelligenza-artificial/> (Consultato il 14-04-2026)
129. Rivoltella, Pier Cesare, Uno dei temi del dibattito sull'IA in scuola è legato al rischio della delega cognitiva, Adnkronos, 2026. Disponibile all'indirizzo: https://www.adnkronos.com/economia/rivoltella-alma-mater-universita-di-bologna-uno-dei-temi-del-dibattito-sullintelligenza-artificiale-in-scuola-e-legato-al-rischio-della-delega-cognitiva_7Ltkoi7YVtzMwiX0dxf4rs (Consultato il 14-04-2026)
 130. Perin, Lorenzo, L'intelligenza artificiale come progetto umano. Intervista a Luciano Floridi, Scienza in rete, 2026. Disponibile all'indirizzo: <https://www.scienzainrete.it/articolo/lintelligenza-artificiale-come-progetto-umano-intervista-luciano-floridi/lorenzo-perin> (Consultato il 14-04-2026)
 131. ISIA Firenze, Dottorato in Cognitive Design, 2025. Disponibile all'indirizzo: <https://www.isiadesign.fi.it/orientamento/ammissione-dottorato/> (Consultato il 14-04-2026)
 132. Università IUAV di Venezia, Ethics and Aesthetics of Artificial Images (EA-AI 2025), 2025. Disponibile all'indirizzo: <https://www.iuav.it/it/notizie/venezia-la-prima-grande-conferenza-internazionale-sulletica-e-lestetica-dellintelligenza> (Consultato il 14-04-2026)
 133. Sapienza Università di Roma, Digital Education Hub, La Sfida della Didattica Universitaria nell'era dell'Intelligenza Artificiale, 2026. Disponibile all'indirizzo: <https://tlc-s.web.uniroma1.it/it/4-e-5-marzo-2026-la-sfida-della-didattica-universitaria-nellera-della-intelligenza-artificiale> (Consultato il 14-04-2026)
 134. Politecnico di Milano, Design through Generative AI, GSOM, 2026. Disponibile all'indirizzo: <https://www.gsom.polimi.it/course/programma-executive-creativita-design-management/design-through-generative-ai/> (Consultato il 14-04-2026)
 135. Politecnico di Milano, Samsung Innovation Campus 2026 — Passion in Action: Design for Trust, 2026. Disponibile all'indirizzo: <https://www.polimi.it/formazione/passion-in-action/dettaglio/samsung-innovation-campus-2026> (Consultato il 14-04-2026)
 136. la Repubblica Torino, Il Politecnico di Torino permette l'utilizzo dell'AI per le tesi, 2026. Disponibile all'indirizzo: https://torino.repubblica.it/cronaca/2026/01/19/news/ai_tesi_politecnico_torino-425103740/ (Consultato il 14-04-2026)
 137. la Repubblica Tecnologia, Il Politecnico di Torino tra le eccellenze italiane selezionate da Google per la ricerca sull'IA, 2026. Disponibile all'indirizzo: https://www.repubblica.it/tecnologia/2026/02/26/news/il_politecnico_di_torino_tra_le_eccellenze_italiane_selezionate_da_google_per_la_ricerca_sull_ia-425186092/ (Consultato il 14-04-2026)
 138. Politecnico di Torino, Hub AI@-PoliTo, 2026. Disponibile all'indirizzo: <https://ai-h.polito.it/> (Consultato il 14-04-2026)
 139. Politecnico di Torino, Polito Design Workshop 2026 — 26 anni di innovazione, 2026. Disponibile all'indirizzo: <https://www.polito.it/ateneo/comunicazione-e-ufficio-stampa/poliflash/polito-design-workshop-2026-26-anni-di-innovazione-e> (Consultato il 14-04-2026)

to il 14-04-2026)

140. Lawfare Media, Understanding Global AI Governance Through a Three-Layer Framework, 2025. Disponibile all'indirizzo: <https://www.lawfaremedia.org/article/understanding-global-ai-governance-through-a-three-layer-framework> (Consultato il 08-05-2026)
141. Medium, The Mirage of a Global Framework for AI Governance, 2024. Disponibile all'indirizzo: <https://medium.com/carre4/the-mirage-of-a-global-framework-for-ai-governance-35b88a36615c> (Consultato il 08-05-2026)
142. Mirage News, Global Push to Establish Urgent AI Governance, 2025. Disponibile all'indirizzo: <https://www.miragenews.com/global-push-to-establish-urgent-ai-governance-1540499/> (Consultato il 08-05-2026)
143. Egyptian Government / AI Authority, Global AI Governance Frameworks: A Comparative Study, 2024. Disponibile all'indirizzo: <https://ai.gov.eg/SynchedFiles/en/Resources/Global%20AI%20Governance%20Frameworks%20A%20Comparative%20Study.pdf> (Consultato il 08-05-2026)
144. Lazard, The Geopolitics of Artificial Intelligence, 2025. Disponibile all'indirizzo: <https://www.lazard.com/research-insights/the-geopolitics-of-artificial-intelligence/> (Consultato il 08-05-2026)
145. Mind Foundry, AI Regulations Around the World, 2025. Disponibile all'indirizzo: <https://www.mindfoundry.ai/blog/ai-regulations-around-the-world> (Consultato il 08-05-2026)
146. OneTrust, Where AI Regulation is Heading in 2026: A Global Outlook, 2026. Disponibile all'indirizzo: <https://www.onetrust.com/blog/where-ai-regulation-is-heading-in-2026-a-global-outlook/> (Consultato il 08-05-2026)
147. Parlamento Europeo e Consiglio dell'Unione Europea, Regolamento (UE) 2024/1689 – Legge sull'Intelligenza Artificiale, 2024. Disponibile all'indirizzo: <https://eur-lex.europa.eu/eli/reg/2024/1689/oj> (Consultato il 08-05-2026)
148. Council of Europe, Convention on Artificial Intelligence, Human Rights, Democracy and the Rule of Law (CETS No. 225), 2024. Disponibile all'indirizzo: <https://www.coe.int/en/web/conventions/full-list?module=treaty-detail&treatynum=225> (Consultato il 08-05-2026)
149. Commissione europea, Orientamenti per i fornitori di modelli di IA per finalità generali, 2025. Disponibile all'indirizzo: <https://digital-strategy.ec.europa.eu/it/policies/guidelines-gpai-providers> (Consultato il 08-05-2026).
150. Avv. A. Ingoglia, AI Act: Guida Completa al Regolamento Europeo sull'Intelligenza Artificiale, 2026. Disponibile all'indirizzo: <https://privacy-legale.it/risorse/ai-act-guida-completa> (Consultato il 08-05-2026).
151. Business People, Fermare l'AI Act: l'appello delle startup italiane ed europee, 2025. Disponibile all'indirizzo: <https://www.businesspeople.it/business/economia/fermare-lai-act-lappello-delle-startup-italiane-ed-europee> (Consultato il 08-05-2026).
152. Commissione europea, Orientamenti per i fornitori di modelli di IA per finalità generali, 2025. Disponibile all'indirizzo: <https://digital-strategy.ec.europa.eu/it/policies/guideli>

nes-gpai-providers (Consultato il 08-05-2026).

153. Euronews, Regole sull'AI rimandate: l'UE divisa tra tutele e pressione delle grandi aziende, 2025. Disponibile all'indirizzo: <https://it.euronews.com/my-europe/2025/11/19/regole-sullai-rimandate-lue-divisa-tra-tutele-e-pressione-delle-grandi-aziende> (Consultato il 08-05-2026).
154. Deirdre Ahern, Operationalising AI Regulatory Sandboxes under the EU AI Act: The Triple Challenge of Capacity, Coordination and Attractiveness to Providers, Cambridge University Press, 2025. Disponibile all'indirizzo: <https://arxiv.org/pdf/2509.05985> (Consultato il 08-05-2026).
155. European Parliament, Recommendation on the Conclusion of the Council of Europe Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law (A10-0007/2026), 2026. Disponibile all'indirizzo: https://www.europarl.europa.eu/doceo/document/A-10-2026-0007_EN.html (Consultato il 12-05-2026).
156. Council of Europe, Responsible Artificial Intelligence – A Council of Europe Priority, 2026. Disponibile all'indirizzo: <https://www.coe.int/en/web/portal/artificial-intelligence> (Consultato il 12-05-2026).
157. Stoyanova, V., Council of Europe Framework Convention on Artificial Intelligence: Context, Regulatory Approach and Scope of Obligations, European Journal of Risk Regulation, Cambridge University Press, 2025. Disponibile all'indirizzo: <https://www.cambridge.org/core/journals/european-journal-of-risk-regulation/article/council-of-europe-framework-convention-on-artificial-intelligence-context-regulatory-approach-and-scope-of-obligations/0E34E-FE5AA7C9628F7053720F88F48C> (Consultato il 12-05-2026).
158. Linkiesta, Alla fine gli Stati Uniti hanno bisogno dell'Europa sui chip, 2026. Disponibile all'indirizzo: <https://www.linkiesta.it/2026/04/alla-fine-gli-stati-uniti-hanno-bisogno-dell-europa-sui-chip/> (Consultato il 08-05-2026).
159. The White House, Executive Order on the Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence (EO 14110), 2023. Disponibile all'indirizzo: <https://www.federalregister.gov/documents/2023/11/01/2023-24283/safe-secure-and-trustworthy-development-and-use-of-artificial-intelligence> (Consultato il 08-05-2026).
160. National Institute of Standards and Technology (NIST), Artificial Intelligence Risk Management Framework (AI RMF 1.0), NIST AI 100-1, 2023. Disponibile all'indirizzo: <https://www.nist.gov/artificial-intelligence/ai-risk-management-framework> (Consultato il 08-05-2026).
161. The White House, Removing Barriers to American Leadership in Artificial Intelligence, 2025. Disponibile all'indirizzo: <https://www.whitehouse.gov/presidential-actions/2025/01/removing-barriers-to-american-leadership-in-artificial-intelligence/> (Consultato il 08-05-2026).
162. Parliament of Canada, Bill C-27: An Act to enact the Consumer Privacy Protection Act, the Personal Information and Data Protection Tribunal Act and the Artificial Intelligence and Data Act, 2023. Disponibile all'indirizzo: <https://www.parl.ca/legisinfo/en/bill/44-1/c-27> (Consultato il 08-05-2026).

05-2026)

163. Government of Canada – Innovation, Science and Economic Development Canada, The Artificial Intelligence and Data Act (AIDA) – Companion Document, 2023. Disponibile all'indirizzo: <https://ised-isde.canada.ca/site/innovation-better-canada/en/artificial-intelligence-and-data-act-aida-companion-document> (Consultato il 08-05-2026)
164. Senado Federal do Brasil, Projeto de Lei nº 2338, de 2023 (PL 2338/2023), 2023. Disponibile all'indirizzo: <https://www.senado.leg.br/noticias/materias/2024/05/29/senado-aprova-projeto-de-lei-de-inteligencia-artificial> (Consultato il 08-05-2026)
165. Mattos Filho, Regulatory Framework for Artificial Intelligence Passes in Brazilian Senate, 2024. Disponibile all'indirizzo: <https://www.mattosfilho.com.br/en/unico/regulatory-framework-for-artificial-intelligence/> (Consultato il 08-05-2026)
166. Gobierno de Chile – Ministerio de Ciencia, Tecnología, Conocimiento e Innovación, Estrategia Nacional de Inteligencia Artificial, 2021. Disponibile all'indirizzo: <https://www.minciencia.gob.cl/politicalIA> (Consultato il 08-05-2026)
167. Ministerio de Ciencia, Tecnología e Innovación de Argentina, Lineamientos Éticos para el Desarrollo y Uso de la Inteligencia Artificial, 2021. disponibile all'indirizzo: <https://www.argentina.gob.ar/ciencia/lineamientos-eticos-ia> (Consultato il 08-05-2026)
168. ScienceDirect, Agile and Iterative Governance: China's Regulatory Response to AI, 2025. Disponibile all'indirizzo: <https://www.sciencedirect.com/science/article/abs/pii/S2212473X25000562> (Consultato il 08-05-2026).
169. Hogan Lovells, China Finalizes Generative AI Regulation, agosto 2023. Disponibile all'indirizzo: <https://www.hoganlovells.com/en/publications/china-finalizes-generative-ai-regulation> (Consultato il 08-05-2026).
170. Future of Privacy Forum, China's Interim Measures for the Management of Generative AI Services: A Comparison Between the Final and Draft Versions of the Text, 2024. Disponibile all'indirizzo: <https://fpf.org/blog/chinas-interim-measures-for-the-management-of-generative-ai-services-a-comparison-between-the-final-and-draft-versions-of-the-text/> (Consultato il 08-05-2026)
171. Geopolitechs, China's AI Law: Recent Developments and Legislative Signals, 2025. Disponibile all'indirizzo: <https://www.geopolitechs.org/p/chinas-ai-law-recent-developments> (Consultato il 08-05-2026)
172. Fondo Monetario Internazionale, Kohei Asao et al., The Impact of Aging and AI on Japan's Labor Market: Challenges and Opportunities, IMF Working Papers 2025/184, 2025. Disponibile all'indirizzo: <https://doi.org/10.5089/9798229024747.001> (Consultato il 08-05-2026).
173. International Bar Association, Japan's Emerging Framework for Responsible AI: Legislation, Guidelines and Guidance, 2025–2026. Disponibile all'indirizzo: <https://www.ibanet.org/japan-emerging-framework-ai-legislation-guidelines> (Consultato il 08-05-2026).

174. Japan AI Safety Institute (AISI), Fact Sheet 2024, 2024. Disponibile all'indirizzo: <https://aisi.go.jp/en> (Consultato il 08-05-2026).
175. Korea Legislation Research Institute, Framework Act on the Development of Artificial Intelligence and Establishment of Trust (AI Basic Act), 2025. Disponibile su: <https://e-law.klri.re.kr> (Consultato il 08-05-2026)
176. Future of Privacy Forum, South Korea's New AI Framework Act: A Balancing Act Between Innovation and Regulation, 2025. Disponibile all'indirizzo: <https://fpf.org/blog/south-koreas-new-ai-framework-act-a-balancing-act-between-innovation-and-regulation/> (Consultato il 08-05-2026)
177. ITIF, One Law Sets South Korea's AI Policy — and One Weak Link Could Break It, 2025. Disponibile all'indirizzo: <https://itif.org/publications/2025/09/29/one-law-sets-south-koreas-ai-policy-one-weak-link-could-break-it/> (Consultato il 08-05-2026)
178. DD News, India AI Governance Guidelines: Enabling Safe and Trusted AI Innovation, 2026. Disponibile all'indirizzo: <https://ddnews.gov.in/en/india-ai-governance-guidelines-enabling-safe-and-trusted-ai-innovation/> (Consultato il 08-05-2026).
179. Economic Laws Practice (ELP), Navigating AI Regulation in India: Unpacking the MeitY Advisory on AI in a Global Context, 2024. Disponibile all'indirizzo: <https://elplaw.in/leadership/navigating-ai-regulation-in-india-unpacking-the-meity-advisory-on-ai-in-a-global-context/> (Consultato il 08-05-2026).
180. Brookings Institution, Cameron F. Kerry, Sovereignty, Safety, and Scale: Takeaways from the India AI Impact Summit, 2026. Disponibile all'indirizzo: <https://www.brookings.edu/articles/takeaways-from-the-india-ai-impact-summit/> (Consultato il 08-05-2026).
181. African Union, Continental Artificial Intelligence Strategy, 2024. Disponibile su: <https://au.int/en/documents/artificial-intelligence-strategy> (Consultato il 08-05-2026)
182. Carnegie Endowment for International Peace, Understanding Africa's AI Governance Landscape: Insights from Policy, Practice, and Dialogue, 2025. Disponibile all'indirizzo: <https://carnegieendowment.org/research/2025/09/africa-ai-governance> (Consultato il 08-05-2026)
183. Government of South Africa — Department of Communications and Digital Technologies, National AI Policy Framework, 2024. Disponibile all'indirizzo: <https://www.unesco.org/ethics-ai/en/southafrica> (Consultato il 08-05-2026)
184. National Centre for Artificial Intelligence and Robotics (NCAIR) Nigeria — NITDA, National Artificial Intelligence Strategy, 2024. Disponibile all'indirizzo: <https://nitda.gov.ng/national-artificial-intelligence-strategy/> (Consultato il 08-05-2026)
185. UAE AI Office, UAE Charter for the Development and Use of AI, 2024. Disponibile all'indirizzo: <https://ai.gov.ae/uae-charter-for-the-development-and-use-of-ai/> (Consultato il 08-05-2026)
186. Saudi Data and AI Authority (SDAIA), AI Adoption Framework, 2024. Disponibile all'indirizzo: <https://sdaia.gov.sa/en/MediaCenter/Do>

cLib/AI-Adoption-Framework.pdf
(Consultato il 08-05-2026)

187. Saudi Data and AI Authority (SDAIA), AI Ethics Principles, 2023. Disponibile all'indirizzo: <https://sdaia.gov.sa/en/SDAIA/about/Documents/AI-Ethics-En-Final.pdf> (Consultato il 08-05-2026)
188. United Nations, Global Digital Compact, 2024. Disponibile all'indirizzo: <https://www.un.org/global-digital-compact> (Consultato il 08-05-2026)
189. OECD, OECD AI Principles, 2024. Disponibile all'indirizzo: <https://oecd.ai/en/ai-principle> (Consultato il 08-05-2026)
190. G7, Hiroshima AI Process International Code of Conduct for Organizations Developing Advanced AI Systems, 2023. Disponibile all'indirizzo: <https://oecd.ai/en/g7> (Consultato il 08-05-2026)
191. Goldman Sachs Global Institute, The generative world order: AI, geopolitics, and power, 2023. Disponibile all'indirizzo: <https://www.goldman-sachs.com/insights/articles/the-generative-world-order-ai-geopolitics-and-power> (Consultato il 08-05-2026)
192. Claudio Novelli et al., Getting Regulatory Sandboxes Right: Design and Governance Under the AI Act, in European Journal of Risk Regulation, 2026. Disponibile all'indirizzo: <https://www.cambridge.org/core/journals/european-journal-of-risk-regulation/article/getting-regulatory-sandboxes-right-design-and-governance-under-the-ai-act/E85D318E227C721BF4823F7930DDA937> (Consultato il 08-05-2026)
193. Bunge, Mario, Verso una tecnocritica, *The Monist*, 60(1), pp. 96–107, 1977. Disponibile all'indirizzo: <https://www.encyclopedia.com/science/encyclopedias-almanacs-transcripts-and-maps/technoethics> (Consultato il 22-05-2026).
194. Vallor, Shannon, *Technology and the Virtues: A Philosophical Guide to a Future Worth Wanting*, Oxford University Press, 2016. Disponibile all'indirizzo: <https://global.oup.com/academic/product/technology-and-the-virtues-9780190498511> (Consultato il 22-05-2026).
195. Vallor, Shannon, *The AI Mirror: How to Reclaim Our Humanity in an Age of Machine Thinking*, Oxford University Press, 2024. Disponibile all'indirizzo: <https://global.oup.com/academic/product/the-ai-mirror-9780197759066> (Consultato il 22-05-2026).
196. Notre Dame Ethics, Dr. Shannon Vallor Explains Virtue Ethics and Technomoral Futures on Lab Podcast, 2024. Disponibile all'indirizzo: <https://ethics.nd.edu/news-and-events/news/dr-shannon-vallor-explains-virtue-ethics-and-technomoral-futures-on-lab-podcast/> (Consultato il 22-05-2026).
197. Il Riformista, L'intelligenza non è un algoritmo: il limite nascosto dell'intelligenza artificiale, 2026. Disponibile all'indirizzo: <https://www.il-riformista.it/lintelligenza-non-e-un-algoritmo-il-limite-nascosto-dellintelligenza-artificiale-506104/> (Consultato il 22-05-2026).
198. Agenda Digitale, Intelligenza artificiale o stupidità reale: le sferzate di Noam Chomsky all'IA, 2024. Disponibile all'indirizzo: <https://www.agendadigitale.eu/cultura-digitale/intelligenza-artificiale-o-stupidita-reale-le-sferzate-di-noam-chomsky->

allia/ (Consultato il 22-05-2026).

199. Anqi Shao, Beyond Misinformation: A Conceptual Framework for Studying AI Hallucinations in (Science) Communication, arXiv, 2025. Disponibile all'indirizzo: <https://arxiv.org/html/2504.13777v1> (Consultato il 22-05-2026).
200. Wired Italia, AI, diritto d'autore e copyright nell'arte, 2025. Disponibile all'indirizzo: <https://www.wired.it/article/ai-diritto-autore-copyright-arte/> (Consultato il 22-05-2026).
201. Charis Avlonitou, Eirini Papadaki, AI: An Active and Innovative Tool for Artistic Creation, MDPI, vol. 14, n. 3, 52, 2025. Disponibile all'indirizzo: <https://www.mdpi.com/2076-0752/14/3/52> (Consultato il 22-05-2026).
202. Il Foglio, Quei rompicoglioni degli intellettuali che non vedono nel progresso solo progresso, 2026. Disponibile all'indirizzo: <https://www.il-foglio.it/cultura/2026/03/28/news/quei-rompicoglioni-degli-intellettuali-che-non-vedono-nel-progresso-solo-progresso--268335> (Consultato il 22-05-2026).
203. Jayesh Rane, Suraj Kumar Mallick, Ömer Kaya, Nitin Rane, Enhancing Black-Box Models: Advances in Explainable Artificial Intelligence for Ethical Decision-Making, ResearchGate, 2024. Disponibile all'indirizzo: https://www.researchgate.net/publication/385131009_Enhancing_black-box_models_Advances_in_explainable_artificial_intelligence_for_ethical_decision-making (Consultato il 22-05-2026).
204. Luciano Floridi, Josh Cows, A Unified Framework of Five Principles for AI in Society, Harvard Data Science Review, 2019. Disponibile all'indirizzo: <https://hdsr.mitpress.mit.edu/pub/10jsh9d1/release/8> (Consultato il 22-05-2026).
205. Commissione Europea, Regulatory Framework for AI, Digital Strategy EC, 2024. Disponibile all'indirizzo: <https://digital-strategy.ec.europa.eu/en/policies/regulatory-framework-ai> (Consultato il 22-05-2026).
206. Reda Hassan, Nhien Nguyen, Stine Rasdal Finserås, Lars Adde, Inga Strümke, Ragnhild Støen, Technological Forecasting and Social Change, Elsevier, 2025. Disponibile all'indirizzo: <https://www.sciencedirect.com/science/article/pii/S0040162525002963> (Consultato il 22-05-2026).
207. University of Cambridge, Cambridge Dictionary Names 'Hallucinate' Word of the Year 2023, 2023. Disponibile all'indirizzo: <https://www.cam.ac.uk/research/news/cambridge-dictionary-names-hallucinate-word-of-the-year-2023> (Consultato il 22-05-2026).
208. Satyadhar Joshi, Comprehensive Review of AI Hallucinations: Impacts and Mitigation Strategies for Financial and Business Applications, Preprints.org, 2025. Disponibile all'indirizzo: <https://www.preprints.org/manuscript/202505.1405/v1> (Consultato il 22-05-2026).
209. la Repubblica, Dall'errore alla disobbedienza: se l'intelligenza artificiale inizia a ignorare i comandi, 2026. Disponibile all'indirizzo: https://www.repubblica.it/tecnologia/2026/03/31/news/dall_errore_alla_disobbedienza_se_l_intelligenza_artificiale_inizia_a_ignorare_i_comandi-425255295/ (Consultato il 22-05-2026).

210. Page Laubheimer, AI Hallucinations, Nielsen Norman Group, 2025. Disponibile all'indirizzo: <https://www.nngroup.com/articles/ai-hallucinations/> (Consultato il 22-05-2026).
211. Crawford, Kate, Atlas of AI: Power, Politics, and the Planetary Costs of Artificial Intelligence, Yale University Press, New Haven, 2021. Disponibile all'indirizzo: <https://kate-crawford.net/atlas> (Consultato il 22-05-2026).
212. Lorenzo Belenguer, AI bias: exploring discriminatory algorithmic decision-making models and the application of possible machine-centric solutions adapted from the pharmaceutical industry, PMC / NCBI, 2022. Disponibile all'indirizzo: <https://pmc.ncbi.nlm.nih.gov/articles/PMC8830968/> (Consultato il 22-05-2026).
213. Nouar AlDahoul, Talal Rahwan & Yasir Zaki, AI-generated faces influence gender stereotypes and racial homogenization, Nature Scientific Reports, 2025. Disponibile all'indirizzo: <https://www.nature.com/articles/s41598-025-99623-3> (Consultato il 22-05-2026).
214. Luhang Sun, Mian Wei, Yibing Sun, Yoo Ji Suh, Liwei Shen, Sijia Yang, Smiling women pitching down: auditing representational and presentational gender biases in image-generative AI, Journal of Computer-Mediated Communication, vol. 29, n. 1, 2024. Disponibile all'indirizzo: <https://academic.oup.com/jcmc/article/29/1/zmad045/7596749> (Consultato il 22-05-2026).
215. MediaLaws, Impact of Meta's Policy Changes on Gender-Based Discrimination in Its Advertising Model, 2025. Disponibile all'indirizzo: <https://www.medialaws.eu/impact-of-metas-policy-changes-on-gender-based-discrimination-in-its-advertising-model-focusing-on-opinion-2025-17-by-dutch-college-for-human-rights> (Consultato il 22-05-2026).
216. G. Cecere, C. Jean, F. Le Guel, M. Manant, Technological Forecasting and Social Change, Elsevier, 2024. Disponibile all'indirizzo: <https://www.sciencedirect.com/science/article/abs/pii/S0040162523008892> (Consultato il 22-05-2026).
217. Jeremy Baum, John Villasenor, Rendering Misrepresentation: Diversity Failures in AI Image Generation, Brookings Institution, 2024. Disponibile all'indirizzo: <https://www.brookings.edu/articles/rendering-misrepresentation-diversity-failures-in-ai-image-generation/> (Consultato il 22-05-2026).
218. Greg Demirchyan, Algorithmic fairness: challenges to building an effective regulatory regime, Frontiers in Artificial Intelligence, 2025. Disponibile all'indirizzo: <https://www.frontiersin.org/journals/artificial-intelligence/articles/10.3389/frai.2025.1637134/full> (Consultato il 22-05-2026).
219. Thilo Hagendorff, The Ethics of AI Ethics: An Evaluation of Guidelines, Minds and Machines, Springer, 2020. Disponibile all'indirizzo: <https://link.springer.com/article/10.1007/s11023-020-09517-8> (Consultato il 22-05-2026).
220. Emmie Malone, Saleh Afroogh, Jason D'Cruz, Kush R. Varshney, When Trust Is Zero-Sum: Automation's Threat to Epistemic Agency, ResearchGate, 2025. Disponibile all'indirizzo: https://www.researchgate.net/publication/392475550_When_trust_is_zero_sum_automatization

- tion's_threat_to_epistemic_agency (Consultato il 22-05-2026).
221. Xia Cheu, How to Outsource Your Frontal Lobes, UMass Media, 2024. Disponibile all'indirizzo: <https://umassmedia.com/40534/opinions/how-to-outsource-your-frontal-lobes/> (Consultato il 22-05-2026).
222. Floridi Luciano, Capitale semantico: la vera posta in gioco della rivoluzione digitale, Wired Italia, 2025. Disponibile all'indirizzo: <https://www.wired.it/article/capitale-semantico-la-vera-posta-in-gioco-della-rivoluzione-digitale-luciano-floridi/> (Consultato il 22-05-2026).
223. Rainer Mühlhoff, Subjectivity and Power in the Ethics of AI, The Ethics of AI: Power, Critique, Responsibility, Bristol University Press, JSTOR Open Access Monograph, 2025. Disponibile all'indirizzo: http://jstor.org/content/oa_chapter_monograph/jj.20433222.10 (Consultato il 22-05-2026).
224. Floridi Luciano, Dalla sabbia alla perla: il capitale semantico nell'era dell'intelligenza artificiale, AI News, 2025. Disponibile all'indirizzo: <https://ainews.it/luciano-floridi-dalla-sabbia-alla-perla-il-capitale-semantico-nellera-dellintelligenza-artificiale/> (Consultato il 22-05-2026).
225. Digital Education Council, AI and the Future of Work: Who Benefits and Who Falls Behind, 2024. Disponibile all'indirizzo: <https://www.digitaleducationcouncil.com/post/ai-and-the-future-of-work-who-benefits-and-who-falls-behind> (Consultato il 22-05-2026).
226. World Economic Forum, Generative AI and Creative Jobs, 2023. Disponibile all'indirizzo: <https://www.weforum.org/stories/2023/05/generative-ai-creative-jobs/> (Consultato il 22-05-2026).
227. Xu Cheng, Yanqi Sun, Haibo Zhou, Artificial Aesthetics and Ethical Ambiguity: Exploring Business Ethics in the Context of AI-driven Creativity, Journal of Business Ethics, Springer, 2024. Disponibile all'indirizzo: <https://link.springer.com/article/10.1007/s10551-024-05837-2> (Consultato il 22-05-2026).
228. Domus Web, Intelligenza artificiale: libri fondamentali per capirla, 2025. Disponibile all'indirizzo: <https://www.domusweb.it/it/notizie/2025/06/27/intelligenza-artificiale-libri-fondamentali-per-capirla.html> (Consultato il 22-05-2026).
229. Cervantes, Juan-Antonio et al., Ascribing Responsibility for the Actions of Learning Automata, Philosophy & Technology, Springer, 2021. Disponibile all'indirizzo: <https://link.springer.com/article/10.1007/s13347-021-00450-x> (Consultato il 22-05-2026).
230. Ellen Hohma, Auxane Boch, Rainer Trauth, Christoph Lütge, Investigating accountability for Artificial Intelligence through risk governance: A workshop-based exploratory study, PMC / NCBI, 2023. Disponibile all'indirizzo: <https://pmc.ncbi.nlm.nih.gov/articles/PMC9905430/> (Consultato il 22-05-2026).
231. DiCulther, Educare al futuro: l'incontro tra intelligenza artificiale e patrimonio culturale, 2025. Disponibile all'indirizzo: <https://www.diculther.it/rivista/educare-al-futuro-lincontro-tra-intelligenza-artificiale-e-patrimonio-culturale-coltivando-la-human-centricity-olivettiana/> (Consultato il 22-05-2026).
232. Thierry Poibeau, Disinformation,

- Misinformation and the Crisis of Trust in AI-Generated Content, ResearchGate, 2025. Disponibile all'indirizzo: https://www.researchgate.net/publication/398762304_Disinformation_misinformation_and_the_crisis_of_trust_in_AI-generated_content (Consultato il 22-05-2026).
233. Tijs Portegies, Oumaima Hajri, Lotte Willemsen, Maaïke Harbers, Decoding AI ethics: What do ethical tools tell us about ethics?, Ethics and Information Technology, Springer, 2025. Disponibile all'indirizzo: <https://link.springer.com/article/10.1007/s10676-025-09860-3> (Consultato il 22-05-2026).
234. Chiara De Vecchis, Paolo Trianiello, La proprietà del pensiero. Il diritto d'autore dal Settecento a oggi, Carocci Editore, 2012. Disponibile all'indirizzo: <https://www.bibliotecheoggi.it/media/download/get/48edc693-6d98-4f3b-8976-21e6b9843da5/original> (Consultato il 22-05-2026).
235. Lala Hajibayova, Authorship in the Age of Artificial Intelligence, ResearchGate, 2025. Disponibile all'indirizzo: https://www.researchgate.net/publication/394504703_Authorship_in_the_Age_of_Artificial_Intelligence (Consultato il 22-05-2026).
236. Wired Italia, Intelligenza artificiale e diritto d'autore: i limiti, 2025. Disponibile all'indirizzo: <https://www.wired.it/article/intelligenza-artificiale-diritto-dautore-limiti/> (Consultato il 22-05-2026).
237. VareseNews, L'intelligenza artificiale secondo Maurizio Ferraris a Gallarate, 2025. Disponibile all'indirizzo: <https://www.varesenews.it/2025/02/lintelligenza-artificiale-secondo-maurizio-ferraris-a-gallarate/2165529/> (Consultato il 22-05-2026).
238. Ferraris, Maurizio, Non c'è niente di più umano dell'intelligenza artificiale, Il Foglio, 2025. Disponibile all'indirizzo: <https://www.ilfoglio.it/tecnologia/2025/04/04/news/non-ce-niente-di-piu-umano-dellintelligenza-artificiale-parla-il-filosofo-maurizio-ferraris--114466> (Consultato il 22-05-2026).
239. Maria Danielsen, Sabine Roeßer, Why emotions and works of art are pertinent to the assessment of the ethical risks of AI, Ethics and Information Technology, Springer, 2025. Disponibile all'indirizzo: <https://link.springer.com/article/10.1007/s10676-025-09849-y> (Consultato il 22-05-2026).
240. Leone XIV, Lettera Enciclica Magnifica Humanitas sulla custodia della persona umana nel tempo dell'intelligenza artificiale, 2026, n. 110. Disponibile all'indirizzo: <https://www.vatican.va/content/leo-xiv/it/encyclicals/documents/20260515-magnifica-humanitas.html> (Consultato il 23-06-2026).
241. Fritjof Capra, Il punto di svolta. Scienza, società e cultura emergente, Feltrinelli, Milano, 1984 (ed. orig. The Turning Point: Science, Society, and the Rising Culture, 1982). Disponibile all'indirizzo: <https://www.ibs.it/punto-di-svolta-scienza-societa-libro-fritjof-capra/e/9788807882319> (Consultato il 23-06-2026).
242. Agnese Azzola et al., Generative AI in the Design Process, FrancoAngeli, 2025. Disponibile all'indirizzo: <https://series.francoangeli.it/index.php/oa/catalog/book/1422> (Consultato il 10-06-2026).

243. Esa Tammisto, Usage of Artificial intelligence in industrial design processes, Lappeenranta–Lahti University of Technology LUT, 2025. Disponibile all'indirizzo: https://lutpub.lut.fi/bitstream/handle/10024/169011/Diplomityo_Tammisto_Esa.pdf?sequence=3&isAllowed=y (Consultato il 10-06-2026).
244. Joana Cerejo, Human-Centered AI Design Process – Part 1, Medium, 2021. Disponibile all'indirizzo: <https://medium.com/design-bootcamp/human-centered-ai-design-process-part-1-8cf7e3ce00> (Consultato il 10-06-2026).
245. Joana Cerejo, Human-Centered AI Design Process – Part 2: Empathize & Hypothesis, Medium, 2021. Disponibile all'indirizzo: <https://medium.com/design-bootcamp/human-centered-ai-design-process-part-2-empathize-hypothesis-6065db967716> (Consultato il 10-06-2026).
246. Yingchao Feng, Artificial Intelligence-Assisted Product Design: From Concept to Prototype, Pioneer Academic Publishing Limited, 2024. Disponibile all'indirizzo: <https://www.pioneerpublisher.com/jpeps/article/view/1130> (Consultato il 10-06-2026).
247. Xilin Tang et al., Human-AI Collaboration Within Industrial Design, AHFE, 2025. Disponibile all'indirizzo: <https://doi.org/10.54941/ahfe1006425> (Consultato il 10-06-2026).
248. Yangluxi Li et al., A Review of Artificial Intelligence in Enhancing Architectural Design Efficiency, Applied Sciences MDPI, 2025. Disponibile all'indirizzo: <https://doi.org/10.3390/app15031476> (Consultato il 10-06-2026).
249. Liu, Y., Fu, Z., Li, T., How Can Artificial Intelligence Transform the Future Design Paradigm and Its Innovative Competency Requisition: Opportunities and Challenges, HCI International 2023 – Late Breaking Papers, Lecture Notes in Computer Science, vol. 14059, Springer, Cham, 2023, pp. 131–148. Disponibile all'indirizzo: https://link.springer.com/chapter/10.1007/978-3-031-48057-7_9 (Consultato il 11-06-2026).
250. O'Regan, S., The Design Difficulties of Human–AI Interaction, 2024. Disponibile all'indirizzo: <https://www.simonoregan.com/short-thoughts/the-design-difficulties-of-human-ai-interaction> (Consultato il 10-06-2026).
251. Sedef Süner-Pla-Cerda et al., Designer experiences and perspectives on the role of generative AI in industrial design, Springer Nature, 2025. Disponibile all'indirizzo: <https://doi.org/10.1007/s00146-025-02504-6> (Consultato il 10-06-2026).
252. Yu-Min Fang, The Role of Generative AI in Industrial Design: Enhancing the Design Process and Education, International Conference on Innovation, Communication and Engineering, 2023. Disponibile all'indirizzo: <https://doi.org/10.1049/icp.2024.0303> (Consultato il 10-06-2026).
253. Jeremiah Folorunso et al., Product Design: The Evolving Role of Generative AI in Creative Workflows, International Journal of Scientific and Management Research, 2025. Disponibile all'indirizzo: <http://doi.org/10.37502/IJSMR.2025.8401> (Consultato il 10-06-2026).
254. Henriikka Vartiainen et al., Emerging human-technology rela-

- tionships in a co-design process with generative AI, Elsevier, 2024. Disponibile all'indirizzo: <https://doi.org/10.1016/j.tsc.2024.101742> (Consultato il 10-06-2026).
255. Netgroup, The New Playbook for Innovation: Stingray Model Design, 2024. Disponibile all'indirizzo: <https://netgroup.com/blog/stingray-model/> (Consultato il 10-06-2026).
 256. José Carlos López Cervantes & Cintya Eva Sánchez Morales, Design in the Age of Predictive Architecture: From Digital Models to Parametric Code to Latent Space, MDPI, 2026. Disponibile all'indirizzo: <https://doi.org/10.3390/architecture6010025> (Consultato il 10-06-2026).
 257. Denny Royal, Designing for an AI World, Tietoevry, 2025. Disponibile all'indirizzo: <https://www.tietoevry.com/en/blog/2025/10/designing-for-ai-world/> (Consultato il 10-06-2026).
 258. Pablo Fernández Vallejo, Redefine Your Design Skills to Prepare for AI, Nielsen Norman Group, 2024. Disponibile all'indirizzo: <https://www.nngroup.com/articles/prepare-for-ai/> (Consultato il 10-06-2026).
 259. Cerejo, J., Designing for Automation vs Augmentation, Medium, 2021. Disponibile all'indirizzo: <https://medium.com/swlh/designing-for-automation-vs-augmentation-d2367a3ede42> (Consultato il 11-06-2026).
 260. Dhanwani, R., Designing the Future with AI: Trends & Insights for 2026, Parallel, 2026. Disponibile all'indirizzo: <https://www.parallelhq.com/blog/future-of-design-with-ai> (Consultato il 11-06-2026).
 261. Sharma, S., AI tools + AI fluency + human advantage = AI-native designer, UX Collective, 2025. Disponibile all'indirizzo: <https://uxdesign.cc/ai-tools-ai-fluency-human-advantage-ai-native-designer-bade3494e7aa> (Consultato il 11-06-2026).
 262. World Economic Forum, The Future of Jobs Report 2025, Ginevra, World Economic Forum, 2025. Disponibile all'indirizzo: <https://www.weforum.org/publications/the-future-of-jobs-report-2025/> (Consultato il 11-06-2026).
 263. The Human Side of Generative AI: Creating a Path to Productivity, McKinsey & Company, 2024. Disponibile all'indirizzo: <https://www.mckinsey.com/capabilities/people-and-organizational-performance/our-insights/the-human-side-of-generative-ai-creating-a-path-to-productivity> (Consultato il 11-06-2026).
 264. The Intersection of Design Thinking and AI: Enhancing Innovation, IDEO U, 2025. Disponibile all'indirizzo: <https://www.ideo.com/blogs/inspiration/ai-and-design-thinking> (Consultato il 11-06-2026).
 265. Rodriguez, H., Think about what's happening here. Hundreds, if not thousands of micro-decisions, LinkedIn, AlxCreative, 2024. Disponibile all'indirizzo: https://www.linkedin.com/posts/hrodriguez1_aifor-design-productdesign-industrialdesign-ugcPost-7167174809406099456-YTQ3/ (Consultato il 11-06-2026).
 266. World Economic Forum, Why AI literacy is now a core competency in education, 2025. Disponibile all'indirizzo: <https://www.weforum.org/stories/2025/05/why-ai-literacy-is-now-a-core-competency-in-education/> (Consultato il 17-06-2026).

267. Digital Education Council, AI adoption is nearly universal among students, but confidence is not, DEC Insights, 2025. Disponibile all'indirizzo: <https://www.digitaleducation-council.com/post/ai-adoption-is-nearly-universal-among-students-but-confidence-is-not> (Consultato il 17-06-2026)
268. L'Espresso, Ministero dell'Istruzione, intelligenza artificiale, storia e geografia nei licei, 2026. Disponibile all'indirizzo: <https://lespresso.it/c/giovani/2026/4/23/ministero-istruzione-intelligenza-artificiale-storia-geografia-licei/61532> (Consultato il 17-06-2026)
269. Guido Saracco, Alleati digitali: la nostra AI personale, Rizzoli, 2026.
270. Chun-Hsin Wu et al., The influence of subjective knowledge, technophobia and perceived enjoyment on design students' intention to use artificial intelligence design tools, International Journal of Technology and Design Education, Springer Link, 2024. Disponibile all'indirizzo: <https://link.springer.com/article/10.1007/s10798-024-09897-3> (Consultato il 17-06-2026)
271. The Conversation, How AI can help in the creative design process, The Conversation, 2024. Disponibile all'indirizzo: <https://theconversation.com/how-ai-can-help-in-the-creative-design-process-244718> (Consultato il 17-06-2026)
272. World Economic Forum, AI skills at scale: How to keep people in the loop, 2026. Disponibile all'indirizzo: <https://www.weforum.org/stories/2026/06/ai-skills-scale-people-in-the-loop/> (Consultato il 17-06-2026)
273. Younjung Hwang et al., The influence of generative artificial intelligence on creative cognition of design students: a chain mediation model of self-efficacy and anxiety, Frontiers in Psychology, 2024. Disponibile all'indirizzo: <https://www.frontiersin.org/journals/psychology/articles/10.3389/fpsyg.2024.1455015/full> (Consultato il 17-06-2026)
274. Stanford HAI, AI challenges core assumptions in education, Stanford Institute for Human-Centered Artificial Intelligence, 2026. Disponibile all'indirizzo: <https://hai.stanford.edu/news/ai-challenges-core-assumptions-in-education> (Consultato il 17-06-2026)
275. Stanford Teaching Commons, Understanding AI Literacy, Stanford University Teaching Guides, 2026. Disponibile all'indirizzo: <https://teachingcommons.stanford.edu/teaching-guides/artificial-intelligence-teaching-guide/understanding-ai-literacy> (Consultato il 17-06-2026)
276. OECD.AI, A socio-technical approach to AI literacy, OECD.AI Wonk, 2026. Disponibile all'indirizzo: <https://oecd.ai/en/wonk/socio-technical-approach-ai-literacy> (Consultato il 17-06-2026)
277. Nick Shockey, Kathryn Zickuhr et al., Red Teaming Generative AI in Classrooms and Beyond, Tech Policy Press, 2024. Disponibile all'indirizzo: <https://www.techpolicy.press/red-teaming-generative-ai-in-classrooms-and-beyond/> (Consultato il 17-06-2026)
278. Wired Italia, Capitale semantico: la vera posta in gioco della rivoluzione digitale secondo Luciano Floridi, 2026. Disponibile all'indirizzo: <https://www.wired.it/article/capitale-semantico-la-vera-posta-in-gioco-della-rivoluzione-digitale-luciano-floridi/> (Consultato il 17-06-2026)

279. Binny Jose et al., The cognitive paradox of AI in education: between enhancement and erosion, *Frontiers in Psychology*, 2025. Disponibile all'indirizzo: <https://www.frontiersin.org/journals/psychology/articles/10.3389/fpsyg.2025.1550621/full> (Consultato il 17-06-2026)
280. Jingwan Li et al., Exploring Challenges and Opportunities to Support Designers in Learning to Co-create with AI-based Manufacturing Design Tools, *Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems*, ACM Digital Library, 2023. Disponibile all'indirizzo: <https://dl.acm.org/doi/10.1145/3544548.3580999> (Consultato il 17-06-2026)
281. Chun-Hsiang Chen et al., Enhancing human computer collaboration in design process: How AI-generated content elevates students' creativity, *International Journal of Educational Technology in Higher Education / ScienceDirect*, 2025. Disponibile all'indirizzo: <https://www.sciencedirect.com/science/article/abs/pii/S1871187125002160?via%3Dihub> (Consultato il 17-06-2026)
282. John T. E. et al., Cultivating independent thinkers: The triad of artificial intelligence, Bloom's taxonomy and critical thinking in assessment pedagogy, *Education and Information Technologies*, Springer Link, 2025. Disponibile all'indirizzo: <https://link.springer.com/article/10.1007/s10639-025-13476-x> (Consultato il 17-06-2026)
283. World Economic Forum, Beyond data culture: human judgement, institutions and AI, *World Economic Forum*, 2026. Disponibile all'indirizzo: [https://www.weforum.org/stories/2026/06/beyond-data-culture-](https://www.weforum.org/stories/2026/06/beyond-data-culture-human-judgement-institutions-ai/)
[human-judgement-institutions-ai/](https://www.weforum.org/stories/2026/06/beyond-data-culture-human-judgement-institutions-ai/) (Consultato il 17-06-2026)
284. Domus, Danny Boyle: per salvare la vita, scegliete la cultura, *Domusweb*, 2020. Disponibile all'indirizzo: <https://www.domusweb.it/it/arte/2020/02/28/danny-boyle-per-salvare-la-vita-scegliete-la-cultura.html> (Consultato il 17-06-2026)
285. Orizzonte Scuola, Intelligenza artificiale a scuola, ecco le 2.100 istituzioni scolastiche finanziate: pubblicato il Decreto, 2026. Disponibile all'indirizzo: <https://www.orizzonte-scuola.it/intelligenza-artificiale-a-scuola-ecco-le-2-100-istituzioni-scolastiche-finanziate-pubblicato-il-decreto/> (Consultato il 17-06-2026)
286. Libero Tecnologia, Intelligenza artificiale nei licei italiani, *Libero*, 2026. Disponibile all'indirizzo: <https://www.libero.it/tecnologia/intelligenza-artificiale-licei-italiani-115903> (Consultato il 17-06-2026)
287. UNESCO, AI Competency Framework for Teachers, *UNESCO Digital Library*, 2024. Disponibile all'indirizzo: <https://www.unesco.org/en/articles/ai-competency-framework-teachers> (Consultato il 17-06-2026)
288. EduNews24, AI Index Report Stanford 2026: l'intelligenza artificiale conquista il mondo, ma l'Italia resta indietro sull'istruzione, 2026. Disponibile all'indirizzo: <https://edunews24.it/mondo/ai-index-report-stanford-2026-lintelligenza-artificiale-conquista-il-mondo-ma-litalia-resta-indietro-sullistruzione> (Consultato il 17-06-2026)
289. Gwo-Jen Hwang et al. (Eds.), *Artificial Intelligence in Education: 26th International Conference, AIED 2025, Palermo, Italy, July 22–26,*

2025, Proceedings, Part III, Springer Link, 2025. Disponibile all'indirizzo: <https://link.springer.com/book/10.1007/978-3-031-98420-4> (Consultato il 17-06-2026)

tificiale fa sentire in colpa, Domusweb, 2025. Disponibile all'indirizzo: <https://www.domusweb.it/it/notizie/2025/09/12/perche-l-intelligenza-artificiale-fa-sentire-in-colpa.html> (Consultato il 17-06-2026)

290. Xinyun Zhang et al., From Consumption to Co-Creation: A Systematic Review of Six Levels of AI-Enhanced Creative Engagement in Education, *Multimodal Technologies and Interaction (MTI)*, MDPI, 2025. Disponibile all'indirizzo: <https://www.mdpi.com/2414-4088/9/10/110> (Consultato il 17-06-2026)

291. Juho Kahila, Henriikka Vartiainen et al., Creative partnerships with generative AI: Possibilities for education and beyond, *International Journal of Educational Research / ScienceDirect*, 2024. Disponibile all'indirizzo: <https://www.sciencedirect.com/science/article/pii/S1871187124002682?via%3Dihub> (Consultato il 17-06-2026)

292. Imran Khan et al., Artificial intelligence in design education: evaluating ChatGPT as a virtual colleague for post-graduate course development, *Design Science*, Cambridge University Press, 2024. Disponibile all'indirizzo: <https://www.cambridge.org/core/journals/design-science/article/artificial-intelligence-in-design-education-evaluating-chatgpt-as-a-virtual-colleague-for-postgraduate-course-development/49EF-FE11646E5BCB09129E5995AD2246> (Consultato il 17-06-2026)

293. Miguel Ángel Escotet, The Future of Artificial Intelligence in Higher Education, *Escotet.org*, 2025. Disponibile all'indirizzo: <https://escotet.org/2025/12/the-future-of-artificial-intelligence-in-higher-education/> (Consultato il 17-06-2026)

294. Domus, Perché l'intelligenza ar-

