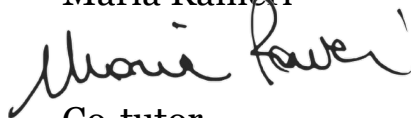


**Corso di Dottorato in
Learning Sciences and Digital Technologies
XXXVIII ciclo**

Artificial Intelligence Literacy.

Dalla costruzione di un framework allo sviluppo di proposte formative nel contesto della scuola e della formazione degli insegnanti

Tutor
Professoressa
Maria Ranieri



Co-tutor
Professoressa
Chiara Panciroli

Dottorando
Gabriele Biagini



Abstract

La rapida diffusione dell'Intelligenza Artificiale (IA) nella società contemporanea ha reso l'alfabetizzazione all'IA (AI literacy) una competenza imprescindibile non solo per gli specialisti, ma per l'esercizio di una cittadinanza consapevole. Questa tesi di dottorato affronta la sfida di definire, misurare e promuovere l'AI Literacy in ambito educativo, con un focus specifico sulla formazione dei futuri insegnanti di scuola dell'infanzia e primaria.

Il lavoro si apre con una revisione sistematica della letteratura che porta alla definizione di un framework teorico originale, articolato in quattro dimensioni fondamentali: **Conoscitiva** (comprensione dei concetti e della storia dell'IA), **Operativa** (capacità di utilizzo e comprensione dei meccanismi sottostanti l'IA), **Critica** (valutazione degli output e riconoscimento dei limiti) ed **Etica** (riflessione sull'impatto dell'AI, responsabilità e valori).

La ricerca si sviluppa attraverso due fasi empiriche principali. In primo luogo, viene descritto il processo di costruzione e validazione del **CAILS (Critical Artificial Intelligence Literacy Scale)**, uno strumento psicometrico atto a misurare le quattro dimensioni del framework. Successivamente, viene presentata la sperimentazione di un curriculum formativo implementato presso l'Università degli Studi di Firenze, che ha coinvolto studenti del corso di Laurea in Scienze della Formazione Primaria. Lo studio adotta un disegno di ricerca misto (mixed-methods), integrando l'analisi quantitativa dei dati pre e post-intervento con un'analisi qualitativa e computazionale (NLP) dei diari metacognitivi dei partecipanti.

I risultati evidenziano un impatto significativo dell'intervento formativo. L'analisi statistica mostra un incremento in tutte le dimensioni dell'AI Literacy, con dimensioni dell'effetto (*effect size*) particolarmente elevate, a dimostrazione dell'efficacia di un approccio didattico che integra teoria, attività pratiche (come il coding unplugged) e metodologie riflessive (come il debate e la Philosophy for Children). L'analisi qualitativa rivela una evoluzione delle concezioni ingenuie dei partecipanti verso un approccio più consapevole. Infatti, da una percezione iniziale dell'IA come entità misteriosa o puramente tecnica, si passa a una visione matura che riconosce la centralità delle dimensioni critiche ed etiche per l'insegnamento.

La tesi si conclude con l'elaborazione di **Linee Guida operative** per l'implementazione dell'AI literacy, proponendo l'alfabetizzazione all'IA non come mero addestramento tecnico, ma come strumento di empowerment democratico e giustizia sociale, essenziale per formare cittadini e educatori capaci di navigare e orientare consapevolmente il futuro digitale.

CAPITOLO 1: Dall'Intelligenza Artificiale all'AI literacy	4
1.1 L'intelligenza artificiale tra inverni e primavere	4
1.1.1 Il concetto di Intelligenza Artificiale: definizioni e prospettive	4
1.1.2 Le origini dell'Intelligenza Artificiale	8
1.1.3 Dai primi sviluppi agli "inverni" dell'IA	11
1.1.4 La rinascita e l'evoluzione contemporanea dell'IA	15
1.1.5 Riflessioni critiche e prospettive future	24
1.2 Revisione sistematica della letteratura sull'AI literacy	29
1.2.1 Contesto e rilevanza dell'AI literacy	29
1.2.2 Metodologia della revisione sistematica	30
1.2.3 Estrazione e codifica dei dati	32
1.2.4 Risultati dell'analisi	34
1.2.5 Definizione e concettualizzazione dell'AI literacy	40
1.2.6 Discussione e implicazioni	44
1.3 Un quadro concettuale per l'AI literacy	49
1.3.1 Dal concetto di literacy all'AI literacy: evoluzione e trasformazione	49
1.3.2 Fondamenti teorici del framework dell'AI literacy	50
1.3.3 Le quattro dimensioni del framework dell'AI literacy	53
1.3.4 Altri framework e modelli	58
1.3.5 Operazionalizzazione e implementazione del framework	60
1.3.6 Sfide e prospettive future	62
CAPITOLO 2: Insegnare l'AI literacy: dai contenuti alle tecniche didattiche	66
2.1 Dal framework ad un curriculum per l'AI literacy: progettare per la comprensione profonda	66
2.1.1 Principi metodologici per la trasposizione del framework in curriculum	67
2.1.2 Progressione delle competenze: verso una padronanza multilivello	68
2.1.3 Selezione e organizzazione dei contenuti per le quattro dimensioni	71
2.1.4 Connessioni interdisciplinari e integrazione curricolare	74
2.2 Sviluppo di strategie didattiche per l'AI literacy: metodologie per un apprendimento attivo e consapevole	76
2.2.1 Approcci pedagogici per l'insegnamento dell'AI literacy	77
2.2.2 Strategie didattiche per la dimensione conoscitiva	78
2.2.3 Strategie didattiche per la dimensione operativa	81
2.2.4 Strategie didattiche per la dimensione critica	83
2.2.5 Strategie didattiche per la dimensione etica	87

CAPITOLO 3: Valutare l'AI literacy: costruzione e validazione di uno strumento di misurazione	93
3.1 Fondamenti teorici e metodologici: cosa dice la ricerca	93
3.1.1 Revisione della letteratura sugli strumenti di misurazione dell'AI literacy.....	93
3.1.2 Analisi comparativa degli strumenti	98
3.1.3 Un modello concettuale per la costruzione di un dispositivo valutativo	101
3.2 Metodologia di sviluppo e validazione dello strumento valutativo	101
3.2.1 Generazione degli item	102
3.2.2 Revisione da parte di esperti e validità di facciata.....	104
3.2.3 Somministrazione pilota e validazione	105
3.2.4 Proprietà psicometriche del questionario	106
3.2.5 Applicazione del questionario sull'AI literacy nella sperimentazione.....	110
3.2.6 Limitazioni e direzioni future.....	112
CAPITOLO 4 La sperimentazione del curriculum per l'AI: il contesto e la metodologia	115
4.1 Il contesto della sperimentazione	115
4.1.1 Obiettivi formativi.....	116
4.1.2 Campione della sperimentazione: futuri insegnanti di fronte all'IA	117
4.1.3 Fasi della sperimentazione.....	120
4.1.4 Struttura e contenuti del curriculum	123
4.2 Le strategie di ricerca	126
4.2.1 Disegno della ricerca	126
4.2.2 Strumenti di raccolta dati.....	128
4.2.3 Procedure di analisi dei dati	130
4.2.4 Considerazioni etiche e riflessioni sulle limitazioni metodologiche.....	134
CAPITOLO 5: La sperimentazione del curriculum per l'AI literacy: risultati e discussione	136
5.1 Analisi dei risultati della sperimentazione: la rilevazione dell'AI literacy.....	136
5.1.1 Caratteristiche dei partecipanti	136
5.1.2 Verifica delle assunzioni statistiche.....	137
5.1.3 Valutazione dell'affidabilità.....	140
5.1.4 Analisi per sottogruppi	142
5.1.5 Risultati dell'analisi statistica	142
5.2 Analisi dei risultati della sperimentazione: i diari metacognitivi.....	153
5.2.1 Analisi del diario 1: dimensione conoscitiva - demistificare l'ignoto.....	161
5.2.2 Analisi del diario 2: dimensione operativa - svelare i meccanismi.....	165
5.2.3 Analisi del diario 3: dimensione critica - riconoscere bias e coltivare giudizio.....	168
5.2.4 Analisi del diario 4: dimensione etica - navigare valori e responsabilità.....	172
5.3 Discussione integrata dei risultati della sperimentazione.....	175

5.3.1 Discussione dei risultati quantitativi	175
5.3.2 Discussione dei risultati qualitativi.....	180
5.3.3 Triangolazione tra risultati quantitativi e qualitativi e implicazioni per l'AI literacy	187
CAPITOLO 6 Elaborazione delle linee guida per l'AI literacy.....	188
6.1 Dalla ricerca alla pratica - verso un'AI literacy per la cittadinanza attiva.....	188
6.1.1 Principi fondamentali per l'implementazione dell'AI literacy nella società	188
6.1.2 Linee guida per la dimensione conoscitiva: costruire le fondamenta concettuali	190
6.1.3 Linee guida per la dimensione operativa: sviluppare competenze pratiche e pensiero computazionale	191
6.1.4 Linee guida per la dimensione critica: sviluppare pensiero analitico e capacità valutative	192
6.1.5 Linee guida per la dimensione etica: promuovere responsabilità e consapevolezza valoriale	194
6.2 Strategie trasversali e approcci metodologici per l'AI literacy.....	195
6.2.1 Implementazione in diversi contesti e per diverse popolazioni	197
6.2.2 Valutazione e monitoraggio dell'AI literacy	199
6.2.3 Sfide e considerazioni etiche nella valutazione.....	200
6.2.4 Prospettive promettenti e innovazioni emergenti	202
6.2.5 Considerazioni: verso un futuro consapevole nell'era dell'Intelligenza Artificiale.....	203

CAPITOLO 1: Dall'Intelligenza Artificiale all'AI literacy

1.1 L'intelligenza artificiale tra inverni e primavera

1.1.1 Il concetto di Intelligenza Artificiale: definizioni e prospettive

Che cos'è l'intelligenza artificiale (IA)? Consultando la letteratura sull'argomento ci accorgiamo subito della difficoltà di proporre una risposta univoca a questa domanda, ma non per difetto, bensì, specie per quanto concerne gli ultimi decenni, per eccesso, se è vero come osserva Quintarelli nel libro da lui curato sull'Intelligenza artificiale: «Tutti siamo capaci di riconoscere la stupidità. Quando parliamo di intelligenza invece ciascuno di noi ha una propria idea e dire con certezza cosa sia è cosa assai ardua, tanto che gli studiosi se ne occupano da secoli. Allo stesso modo, se chiedessimo a 100 informatici di dare una definizione di Intelligenza Artificiale, otterremmo almeno 101 risposte (per via degli errori di arrotondamento)» (Quintarelli et al., 2020, p. 13).

Tentare di definire cosa sia l'Intelligenza Artificiale (IA) significa quindi affrontare un tema complesso che riflette l'incessante carattere poliedrico ed evolutivo di questo campo, come vedremo proponendo un breve spoglio di alcune delle principali definizioni circa la natura e la funzione dell'IA formulate nel corso del tempo.

Il teorico Marco Somalvico, nell'Enciclopedia Italiana (1992), scriveva, ad esempio: «L'i.a. è una moderna disciplina sorta nell'ambito della scienza dei calcolatori e dell'informatica che negli anni recenti, specialmente per merito dell'avvento dei sistemi esperti, ha avuto un'importante influenza nel progresso dell'intera informatica. Inoltre, l'i.a. riveste un indubbio interesse per coloro che siano interessati a un'approfondita valutazione delle potenzialità dell'informatica e degli effetti che il progresso di questa disciplina ha e avrà nei confronti dell'uomo. L'i.a. è innanzitutto la denominazione di una disciplina e non va, quindi, confusa con l'ipotetica connotazione della funzionalità complessa sviluppata in una macchina. Questa precisazione è molto importante perché nella letteratura giornalistica di divulgazione spesso viene indicata con la dizione 'intelligenza artificiale' la funzione avanzata di elaborazione che la macchina è in grado di compiere. [...] L'i.a. è dunque la disciplina recente, appartenente all'informatica, che studia i fondamenti teorici, le metodologie e le tecniche che permettono di progettare sistemi digitali (hardware) e di programmi (software) capaci di fornire all'elaboratore elettronico prestazioni che, a un osservatore comune, sembrerebbero di pertinenza esclusiva dell'intelligenza umana».

La distinzione di Somalvico tra l'IA come disciplina e l'IA come funzionalità (scienza e/o tecnica) è fondamentale perché evidenzia i possibili equivoci nella comprensione del fenomeno. In ogni caso la definizione sottolinea la peculiarità dell'IA di scienza complessa. Una complessità sottolineata anche da Riguzzi (2006) nel richiamare le circostanze della nascita dell'IA:

«Il termine Intelligenza Artificiale [IA] è stato inventato da John McCarthy nel 1956, in occasione di un seminario di due mesi (da lui organizzato al Dartmouth College di Hanover, New Hampshire, USA) che ebbe il merito di far conoscere tra loro 10 studiosi (su teoria degli atomi, reti neurali, e intelligenza) statunitensi e di dare l'imprimatur al termine 'Intelligenza Artificiale' come nome ufficiale del nuovo campo di ricerca. Da allora l'IA si è affermata ed evoluta; oggi è riconosciuta come branca autonoma, sebbene connessa a informatica, matematica, scienze cognitive, neurobiologia e filosofia».

Jerry Kaplan aggiunge una prospettiva più problematica: ovvero la riflessione circa la difficoltà di concepire l'IA in quanto macchina capace di comportarsi come l'intelligenza umana: «'Che cos'è l'intelligenza artificiale?' Ecco una domanda a cui è sia facile che difficile rispondere per due ragioni:

la prima è che non c'è comune accordo su cos'è l'intelligenza artificiale. La seconda è che ci sono poche ragioni, almeno al momento, per ritenere che l'intelligenza delle macchine abbia molto in comune con quella umana. Sono state proposte molte definizioni di intelligenza artificiale (IA), ognuna espressione di un diverso punto di vista, ma tutte più o meno allineate sul concetto di creare programmi informatici o macchine capaci di comportamenti che riterremmo intelligenti se messi in atto da esseri umani. Nel 1955 John Mac Carthy, uno dei padri fondatori della disciplina, descriveva il processo come 'consistente nel far sì che una macchina si comporti in modi che sarebbero definiti intelligenti se fosse un essere umano a comportarsi così'» (2017, p. 15).

È invece la natura artificiale, tecnica, dell'IA, da intendersi in quanto "opera dell'uomo", a caratterizzare la definizione dell'IA adottata da Urwin: «'Artificiale' -L'intelligenza trattata in questo libro è opera dell'uomo, vale a dire che deve essere un prodotto della scienza e della tecnica. Non c'è alcun contributo della magia [...]. Questo è il motivo per il quale l'intelligenza artificiale è sviluppata con mezzi informatici. Tutte le tecnologie trattate in questo libro sono implementate su computer» (Urwin, 2019, pp. 8-9).

Infatti, stesso lo Urwin descrive l'IA «come uno strumento costruito per coadiuvare e sostituire il pensiero umano», precisando che si tratta di un «programma informatico che funziona in modo indipendente all'interno di un centro di elaborazione dei dati o di un personal computer, oppure inserito in un apparecchio come un robot, in grado di dare segni di intelligenza: vale a dire la capacità di acquisire e applicare conoscenze e competenze allo scopo di operare adeguatamente in un dato ambiente» (Urwin, 2019, pp.15-16).

Dal canto suo Ricciardi (2019, p. 5) ha optato per l'accezione scientifica dell'IA pensata proprio in quanto tecnologia in grado di "imitare" l'intelligenza umana: «Una scienza è un insieme di tecniche computazionali che vengono ispirate – pur operando tipicamente in maniera diversa – dal modo in cui gli esseri umani utilizzano il proprio sistema nervoso e il proprio corpo per sentire, imparare, ragionare e agire».

Su questa base di sistema intelligente, per Margaret A. Boden è importante tornare a sottolineare la doppia natura (scientifica e tecnologica) dell'IA: «L'Intelligenza artificiale (IA) cerca di mettere in grado i computer di fare il genere di cose che fanno le menti [...]. L'intelligenza non è a una sola dimensione, ma è piuttosto uno spazio riccamente strutturato di diverse capacità di elaborazione dell'informazione. Di conseguenza, l'IA utilizza varie tecniche, che affrontano molti compiti differenti [...]. L'IA ha due obiettivi principali. Un obiettivo è tecnologico: usare i computer per fare cose utili [...]. L'altro obiettivo è scientifico: usare i concetti e i modelli dell'IA per contribuire a rispondere a interrogativi che riguardano gli esseri umani e gli altri esseri viventi» (2019, pp. 7-8).

Kate Crawford (2021, pp. 15-16) propone invece della natura dell'IA una lettura decisamente critica per il consumo di energia e in quanto «registro di potere»: «[In] questo libro io sostengo che l'IA non è intelligente né artificiale. Piuttosto, l'intelligenza artificiale è sia incarnata che materiale, composta da risorse naturali, combustibili, lavoro umano, infrastrutture, logistica, storie e classificazioni. I sistemi di IA non sono autonomi o razionali, né in grado di discernere alcunché senza una fase di formazione estensiva ma computazionalmente intensiva con grandi set di dati o regole o ricompense predefinite. In effetti l'intelligenza artificiale come la conosciamo dipende interamente da un insieme molto più ampio di strutture politiche e sociali [...]. In questo senso, l'intelligenza artificiale è un registro di potere».

Un po' come Ricciardi, Elliott, introducendo elementi di percezione sociale dell'IA, sottolinea la natura tecnologica dell'IA, dato che, per macchina intelligente si intende una macchina dotata di

capacità di autoapprendimento (Machine Learning): «La relazione tra l'intelligenza artificiale e le sue tecnologie, con particolare riguardo per le esperienze o le opinioni delle persone in tema di IA, è piuttosto complicata. Per le finalità del presente studio, la mia definizione di IA, e dell'apprendimento automatico che ne è diretta conseguenza, includerà qualsiasi sistema computazionale che sia in grado di discernere il contesto per esso rilevante e di reagire in modo intelligente ai dati. Le macchine possono essere definite intelligenti (e ottenere, pertanto, il crisma di IA) quando si registrano determinati livelli di auto-apprendimento, di autocoscienza e di capacità percettive» (2021, p. 27).

Paradossalmente, Melanie Mitchell (2022, p. 4) assume la difficoltà di giungere a una definizione univoca dell'IA come un orizzonte stimolante e aperto per la ricerca: «In un recente report sullo stato corrente dell'IA, un comitato di ricercatori di punta l'ha definita 'una branca della scienza informatica che studia le proprietà dell'intelligenza sintetizzando intelligenza' [...]. Ma lo stesso comitato ha anche ammesso che è difficile definire questo campo, ed è meglio così: La mancanza di una definizione precisa, universale, ha probabilmente contribuito a far crescere, fiorire e progredire questo campo con un ritmo via via crescendo».

Floridi (2022, pp. 34-35) rovesciando l'ottica comune sull'IA (tecnica e scienza) propone una interessante provocazione epistemologica nell'attribuire alla "rivoluzione digitale", e al successo dell'IA che ne è conseguito, "un divorzio tra l'agire e l'intelligenza": «Alcune persone (forse molte) sembrano credere che l'IA riguardi la congiunzione tra agire artificiale e comportamento intelligente in nuovi artefatti. Questo è un malinteso [...], in realtà, è vero il contrario: la rivoluzione digitale ha reso l'IA non solo possibile ma sempre più utile separando la capacità di risolvere un problema o di portare a termine un compito con successo dall'essere intelligenti nel farlo. L'IA ha successo proprio quando è possibile realizzare tale separazione. Le conseguenze derivanti dal comprendere l'IA come un divorzio tra l'agire e l'intelligenza sono profonde, così come lo sono le sfide etiche a cui dà origine tale divorzio».

Infine, Romano (2024, pp. 9-10 e 29) ci propone una definizione che sintetizza i molteplici aspetti che caratterizzano la complessità dell'IA, alcuni dei quali già emersi: «Definire oggi cosa sia l'IA, tracciarne con precisione i confini, e circoscriverne con precisione il dominio, appare sempre più complesso, quasi impossibile, a causa della sua continua evoluzione a ritmi esponenziali e delle molteplici sfaccettature e ottiche con cui è possibile interpretarla. L'IA si configura come un fenomeno che travalica i paradigmi conoscitivi tradizionali, obbligandoci a ripensare il nostro stesso ruolo nel mondo e i principi fondanti delle nostre società. Le sue applicazioni, apparentemente illimitate, spaziando dai campi delle scienze fisiche e naturali a quelle delle scienze sociali, dell'arte, della medicina, fino alla filosofia e all'educazione, offrono strumenti potenti per facilitare e velocizzare le attività umane [...]. Oggi, l'IA è una disciplina profondamente interdisciplinare, che combina informatica, neuro-scienze, scienze cognitive, psicologia, pedagogia, scienze giuridiche e, proprio per questo, è anche sottoposta a continui dibattiti etici ed epistemologici riguardo al suo impatto crescente sulla società e sugli individui, sia alle sfide che pone in termini di responsabilità e trasparenza».

Per molti versi il quadro tracciato dalla studiosa ci riconduce a considerazioni analoghe viste (e che avremmo potuto vedere, palesemente o in filigrana), in altri teorici. Molto convincente ci sembra, ad esempio, la nota di Melanie Mitchell (2022, p. 7) che, riflettendo sulle «differenti dimensioni di intelligenza, emotiva, verbale, logica, artistica, sociale e così via», paragona la stessa parola intelligenza a «una valigia stracolma la cui cerniera è sul punto di rompersi». Il fatto interessante per noi è che, secondo Mitchell, l'IA «eredita questo problema di valigia stipata, sfoggiando significati differenti in contesti differenti». Un simile pensiero ci permette di motivare la pluralità dei pareri

presenti nei vari tentativi di definizione dell'IA, riconducendoli alla complessità intrinseca del fenomeno e alla sua natura profondamente interdisciplinare e in movimento.

Si consideri che solo qualche decennio fa, definizioni così diverse sull'IA non ne avremmo trovate e, a proposito del tema principale relativo alla creazione di macchine intelligenti che imitano (o dovrebbero imitare) in tutto gli esseri umani – magari influenzati dallo scetticismo di un teorico come Searle, circa la possibilità di una versione forte dell'IA, che emula il cervello umano – avremmo invece pensato di dover continuare ad immergerci in un campo di ipotesi fantascientifiche.

Inoltre, dinanzi alla necessità di riflettere su un tema importante come quello dell'Intelligenza Artificiale ormai «arrivata tra noi», non solo ci saremmo trovati dinanzi al problema della scarsità dei riferimenti bibliografici, per tentare di rispondere a quesiti preliminari come "Che cos'è l'Intelligenza Artificiale (IA)?" ma avremmo incontrato la stessa difficoltà a distinguerne principi e modalità (forte/debole/pragmatica; simbolica/subsimbolica; generale/ristretta/generativa) e a delinearne esattamente i risvolti storici.

Attualmente, la situazione di scarsità di riferimenti è molto cambiata: adesso è piuttosto l'abbondanza e la varietà dei dati bibliografici, testi, siti web, voci da compulsare e consultare a creare il rischio di confusione e sovrapposizione nei percorsi storiografici e teorici. In questo contesto ci tornano utili le riflessioni, espresse in altro ambito – quello della rete e dell'accesso alle tecnologie – da Maria Ranieri, quando parla della «necessità di un'educazione all'uso consapevole dell'enorme quantità di dati e informazioni, spesso irrilevanti, futili e inaffidabili che circola in rete», raccomandando un «uso critico, attivo e consapevole dell'informazione» (Ranieri, 2006, p. 13).

Ma il rischio di farne un "uso non critico" non è sempre evitabile: la crescita dell'attenzione speculativa e il mutamento dello scenario degli studi non sono che una conseguenza della necessità di tenere il passo con il cambiamento generalizzato che riflette la misura dell'accelerazione delle conoscenze, del moltiplicarsi dei dibattiti – particolarmente attuali dopo la nota, e recente, proposta di moratoria di sei mesi degli esperimenti avanzati di IA (29 marzo 2023) – sugli effetti dell'enorme sviluppo socioeconomico che ha conosciuto l'applicazione della tecnologia in ogni campo (ricerca scientifica, economia, automazione, informazione, comunicazione, medicina, formazione, lavoro, creazione artistica – film, musica, fotografia ...).

In questa prospettiva non è difficile convenire che, inventato meno di cento anni fa, il computer è ormai parte della nostra quotidianità, ha cambiato le nostre vite e, come ha osservato di recente ancora M. Mitchell, «[i] computer sembrano diventare più intelligenti a un ritmo allarmante» (Mitchell, 2022, p. VII), sebbene – e lo ricorda anche Piero Angela – «i computer [siano] macchine straordinarie, ma da soli non possono funzionare» (Quintarelli, 2020, p. 7).

Il pensiero di Piero Angela è corretto, ma è vero anche che le macchine intelligenti fanno già tutto (o quasi), come nota Nello Cristianini: «Vagliano curricula, concedono mutui, scelgono le notizie che leggiamo: le macchine intelligenti sono entrate nelle nostre vite, ma non sono come ce le aspettavamo. Fanno molte delle cose che volevamo, e anche qualcuna in più, ma non possiamo capirle o ragionare con loro, perché il loro comportamento è in realtà guidato da relazioni statistiche ricavate da quantità sovrumane di dati. Eppure, possono essere in certi casi più potenti di noi: ci osservano continuamente, e prendono decisioni al nostro posto. E allora come incorporarle nella nostra società senza rischi ed effetti collaterali?» (Cristianini, 2023, risvolto di copertina).

La riflessione di Cristianini (e altri molto vicini a questa) sui molteplici usi dell'IA può servire come implicita indicazione a sospendere la domanda ("Cos'è l'IA?") per affiancarla ad altre: "A cosa serve l'IA?" e "Come funziona?" ma, soprattutto, "Quando è nata?" e "Da cosa?" Per quest'ultima domanda

estrapoliamo una risposta da Melanie Mitchell quando accenna ad una «anarchia di metodi». Un'anarchia evidente, secondo la studiosa, sin dal 1956, tra i partecipanti all'incontro di Dartmouth, dove erano presenti grandi pionieri dell'IA: studiosi in matematica (Marvin Minsky), in teoria dell'informazione (Claude Shannon), in ingegneria elettronica (Nathaniel Rochester), in informatica e psicologia cognitiva (Allen Newell), in economia, psicologia e informatica (Herbert Simon, futuro Premio Nobel) e dove fu coniato, da John McCarthy, il termine Intelligenza Artificiale che prevalse soprattutto per l'intenzione di McCarthy di «distinguere questo campo da un'impresa affine, chiamata cibernetica» (Mitchell, 2022, p. 6; Kaplan, 2017, p. 31).

È nato da qui un progetto così complesso che si proponeva sin da allora di simulare l'intelligenza umana e di occuparsi, allora come ora, di problemi di «elaborazione del linguaggio naturale, reti neurali, machine learning, concetti astratti e ragionamento, creatività» e non poteva che richiedere ogni possibile risorsa dell'intelligenza e delle competenze umane. Si trattava soprattutto di direzioni da intraprendere attraverso scelte di metodi e competenze che non potevano essere univoche, trattandosi di conciliare interessi scientifici e settori di ricerca diversi: matematica, informatica, teorie della comunicazione, psicologia, neuroscienze e biologia.

Mancando una conciliazione non restava che decidere di adottare una «famiglia di metodi dell'IA – chiamata collettivamente deep learning o deep neural networks – che, secondo Mitchell, si è elevata al di sopra di questa anarchia diventando il paradigma dominante. Al punto che in buona parte dei media il termine di intelligenza artificiale è identificato con il 'deep learning'. La studiosa chiarisce subito che si è trattato di una «approssimazione infelice». Dal suo punto di vista l'«IA è un campo che include un ampio insieme di approcci, il cui obiettivo è creare macchine intelligenti» e il deep learning sarebbe «soltanto uno degli approcci» (Mitchell, 2022, p. 9).

1.1.2 Le origini dell'Intelligenza Artificiale

Se definire l'IA, i settori e i metodi che ne sono coinvolti, appare arduo, non accade diversamente per i tentativi di proporre una linea di storicizzazione generalizzabile. Abbiamo già accennato al fatto che negli ultimi tempi, intorno all'IA, si è venuto a creare un contesto di riferimenti bibliografici più complesso anche in ambito storiografico quasi a segnare il passaggio ad un approccio storico più ampio e importante dal punto di vista critico e comparato, dove, ad esempio, ad interessare sono anche le ragioni antropologiche, economiche, sociologiche, che spiegano come l'umanità sia arrivata all'approdo dell'Intelligenza Artificiale passando attraverso uno sviluppo scientifico e tecnologico inarrestabile.

Un progresso che ci ricorda come sin dai primordi l'umanità sia stata spinta in avanti dal desiderio di inventare e realizzare tecnologie che sono state spesso frutto di sogni, progetti di scienziati considerati veri e propri padri pionieri. Resta il fatto che, in quanto all'AI, la curva storica «delle tendenze di ricerca più rilevanti e dei maggiori eventi che hanno determinato in seguito sviluppi differenti non è omogenea» (Tonfani e Schneider, 1984, p. 10) ed è segnata da epoche di crescita di interesse, di scoperte e di stasi. Mitchell ha evocato una fase di «inverno dell'IA» situata proprio tra gli anni Settanta e Ottanta (Mitchell, 2022, pp. 21-22).

Tuttavia, anche in questo caso è come se ci si trovasse a dover scegliere tra diverse storie dell'IA, non tutte complete e influenzate dal prevalere di interessi estrinseci (relativi ai campi dell'economia e della finanza) ed intrinseci: gli sviluppi scientifici delle ricerche in vari settori (scienze psicologiche e biologiche, neuroscienze, filosofia del linguaggio, ingegneria elettronica e informatica, teorie dell'informazione). Per fortuna si tratta di una storiografia documentabile, dove i punti fermi si basano su notizie che consentono di stabilire incroci utili tra le varie storicizzazioni.

In questa prospettiva, faremo riferimento ad alcuni capitoli storiografici tra i più recenti, riferibili ad autori che citeremo in corso di lettura – alcuni dei quali già visti, Mitchell, 2019; Quintarelli et al., 2020; Kaplan 2016/2017; Boden, 2019; Urwin, 2019; Ricciardi e Sacco, 2019 – e servono ad evitare di proporre una griglia personale imprecisa e non documentata. La scelta dei riferimenti può sembrare limitata o arbitraria. In realtà, è determinata oltre che dalla fiducia negli autori, anche da uno sguardo interessato a ripercorrere le "origini dell'IA" seguendo punti di vista storico-riflessivi e sistematizzati.

È noto che le radici storiche dell'IA sono state rintracciate ben prima della sua nascita ufficiale come disciplina, se è vero, come rilevano Quintarelli et al., (2020, p. 20), che «le prime idee sull'Intelligenza Artificiale precedono la tecnologia che le ha rese possibili» e possono essere ricondotte al pensiero di filosofi come Hobbes, La Mettrie, Vaucanson e Kant. Insieme all'immaginario artistico e letterario ispirato dalle macchine pensanti, hanno contribuito a gettare le basi concettuali per lo sviluppo dell'IA.

Le prime progettazioni concrete di automi e macchine con schede programmabili vengono ormai considerate "antenate degli elaboratori elettronici". Tra queste, il Mechanical Turk del 1769, un congegno di simulazione robotica usato per vincere al gioco degli scacchi, e il telaio di Joseph Marie Jacquard (1752-1834), basato su dispositivi di schede perforate programmabili, inserite in «una macchina destinata a sostituire il tiratore di lacci nelle stoffe operate», in uso nel passato.

Un passaggio importante nello sviluppo delle idee che avrebbero portato all'IA si ebbe nel 1833, con gli studi della Contessa Ada Lovelace (studiosa di matematica, figlia di Lord Byron – 1815-1852) sulla Macchina analitica di Babbage, una macchina capace di creare calcoli e algoritmi. La Macchina analitica, ideata da Charles Babbage, matematico e filosofo londinese (1791-1871), non fu mai realizzata ma viene tuttora considerata il «primo prototipo di computer meccanico».

Nel 1920, Karel Čapek pubblica R.U.R. Rossum's Universal Robots, un testo teatrale che introdusse il termine "robot" (dal ceco "robota", corvée) per definire degli operai artificiali, andropoidi creati dallo scienziato Rossum. L'opera ha determinato il successo del neologismo (il termine ceco robota, che passa dal femminile al maschile) di cui Karel condivide la paternità con il fratello Josef, pittore. Il termine ha finito per definire l'ibridismo tra umano e non umano e, per slittamento semantico, anche l'universo della meccanizzazione.

Sulla linea di contatto tra speculazione scientifica e immaginario della fantascienza, che ritroveremo spesso, nel film muto Metropolis (1927) Fritz Lang mette sulla scena il Prof. Rotwang, uno scienziato pazzo che ha inventato macchine lavoratrici e un «prototipo androide» dai tratti femminili (Maria), che ha il compito di scatenare una rivolta contro chi detiene il potere. Il film è considerato come il primo tentativo di rappresentazione degli atavici timori degli esseri umani di essere sostituiti da macchine sempre più efficienti e pervasive.

Un contributo importante, seppure indiretto, allo sviluppo tecnologico che avrebbe poi permesso l'emergere dell'IA venne nel 1942, con il brevetto di Hedy Kiesler Markey (diva del cinema nota come Hedy Lamarr, 1914-2000) e George Antheil (musicista e inventore). Si trattava di un «nuovo sistema di modulazione (oggi noto come Frequency-hopping spread spectrum) in uso per la codifica di informazioni da trasmettere su frequenze radio, importante per comandare a distanza siluri e mezzi navali». Il sistema si rivelò provvidenziale per le sorti degli alleati nell'ultima guerra mondiale.

Sempre nel 1942, Isaac Asimov, scienziato e scrittore, nel suo racconto Runaround (Circolo Vizioso) formula le Tre Leggi della Robotica, evidenziando le dimensioni culturali ed etiche dell'impatto dell'IA sulla società. Queste le leggi: «1. Un robot non può recar danno a un essere umano, né permettere che, a causa della propria negligenza, un essere umano patisca danno. 2. Un robot deve

sempre obbedire agli ordini degli esseri umani, a meno che contrastino con la prima legge. 3. Un robot deve proteggere la propria esistenza, purché questo non contrasti con la Prima e la Seconda legge». In seguito, Asimov aggiunse la legge zero: «Nessun robot può causare danni all'umanità, né permettere, con la sua inazione, che l'umanità subisca danni».

Nel 1943, Warren McCulloch (neurofisiologo) e Walter Pitts (matematico e neuroscienziato) compiono ricerche fondamentali nell'ambito delle reti neurali artificiali e, in un saggio (*A Logical Calculus of Ideas Immanent in Nervous Activity*), propongono quello che viene considerato il primo modello matematico di neurone artificiale.

Nel 1949, Grace Murray Hopper (studiosa di matematica, informatica e militare statunitense, 1906-1992) entra a far parte della «Eckert-Mauchly Computer Corporation», la società che aveva sviluppato l'ENIAC, uno dei primi computer digitali in circolazione. Hopper, lavorando sull'idea di «compilatore» (programma informatico che traduce istruzioni) ebbe un ruolo primario nella progettazione e nello sviluppo del COBOL (linguaggio di programmazione in uso in molti software commerciali). È nota anche per aver diffuso l'uso del termine "bug".

In modo preliminare possiamo dire che uno dei punti di incontro più frequenti tra i vari percorsi della storia dell'IA, da noi esaminati, consiste nell'indicazione di una fase di cambiamento intorno agli anni Sessanta. È la fase che segna la cesura fondamentale tra un prima e un dopo Turing (in pratica prima e dopo gli anni centrali del Novecento). Non a caso un simile discrimine è stato indicato dallo studioso Mario Ricciardi nel suo commento alla silloge dei testi dei maggiori teorici (Turing, Wiener, Busch, Engelbart), dove vengono caratterizzate due diverse fasi di studi e ricerche: la prima definita come fase di ricerca del «primato della mente» (Turing e la sua idea di "costruire un cervello" e per Vannevar Bush interrogarsi su come potremmo pensare – "As We May Think" –) che situa i protagonisti in un contesto di «totalitarismi, due guerre mondiali, l'Olocausto e il cervello atomico», consapevoli però di aver contribuito alla «nascita di una nuova scienza [...] che comporta sviluppi tecnici con grandi possibilità per il bene e per il male. La seconda fase ha inizio negli anni Sessanta» (Ricciardi, in Ricciardi e Sacco, 2019, pp. 6-7).

Con tutta evidenza la prima fase abbraccia i testi dei "grandi padri", descritti e analizzati nel volume citato. Un volume che puntualizza anche i pochissimi riconoscimenti e gli antagonismi persecutori che accompagnarono la vicenda esistenziale dei due maggiori autori della «svolta comunicativa del XX». «La mela avvelenata» del titolo del libro non è solo quella reale che uccide Turing, ma anche quella metaforica – politica – diretta contro Wiener, inserito nella lista nera della Commissione Mc Carthy.

Una "svolta" che si sostanzia nelle ricerche di Alan M. Turing (nel saggio *On Computable Numbers, with an Application to the Entscheidungsproblem* del 1936, in cui si descrive un modello concettuale del computer che fa riferimento alla cosiddetta «macchina di Turing») sull'avanzamento degli studi sulla comunicazione. Norbert Wiener, secondo Ricciardi «trasforma, attraverso la cibernetica, il computer in macchina per comunicare e quindi in medium universale» (Ricciardi, in Ricciardi e Sacco, 2019, p. 33); e nel progetto concepito da Vannevar Bush, ingegnere informatico, che nel 1945 pubblica l'articolo *As We May Think* [Come potremmo pensare] dove prefigura una macchina enciclopedica futuribile (Memex: acronimo di memory expansion, corredata da una vera e propria rete associativa di dati, in pratica una rete di ipertesti).

Negli anni centrali del Novecento (il secondo dopoguerra), con lo sviluppo dei «sistemi di trasmissione», prende veramente corpo «l'idea che le macchine possano pensare» (Quintarelli et al., 2020, p. 12).

Il 1950 è considerato un anno capitale nella storia dell'IA. Alan M. Turing (1912-1954), già famoso come filosofo, matematico, crittografo (si deve al suo genio la decodifica di Enigma, i messaggi cifrati trasmessi dai nazisti durante la guerra), pioniere assoluto dell'informatica (la macchina di Turing è considerata l'antesignana teorica diretta del computer) e dell'Intelligenza Artificiale, pubblicò un saggio (dal famoso incipit: «Possono pensare le macchine?») in cui propose una sorta di gioco simulato (Computing machinery and intelligence) noto come il Test di Turing. Nel test, Turing simula un gioco che consiste nell'interrogare, attraverso un terminale, due interlocutori (un uomo e un computer) distanti fra loro, al fine di determinare se le loro risposte sono distinguibili. Siccome non lo sono, Turing deduce che la macchina abbia dato prova di «comportamento intelligente» (Quintarelli et al., 2020, p. 21).

Come abbiamo già visto, il termine "Intelligenza Artificiale" fu ufficialmente coniato nel 1956 durante la Dartmouth Summer Research project on Artificial Intelligence. Durante otto settimane, in un seminario deciso l'anno prima, ne discussero John McCarthy – che suggerì il termine – Claude Shannon, Marvin Minsky, Nathaniel Rochester e altri sei partecipanti: Ray Solomonoff, Oliver Selfridge, Trenchard More, Arthur Lee Samuel (responsabile del concetto Machine Learning – ML o Apprendimento automatico), Allen Newell e Herbert Simon. In quella stessa occasione vennero indicati come campi di ricerca specifici dell'IA gli studi sulle «reti neurali, la teoria della computabilità, la creatività e l'elaborazione del linguaggio naturale».

Per tutti gli storici dell'IA si tratta dell'evento fondativo di cui si serve anche Kaplan per segnare la nascita ufficiale dell'IA in quanto disciplina autonoma e l'inizio di quella che è spesso indicata come una precoce "primavera dell'IA", caratterizzata da ottimismo e volontà di sperimentazione grazie a significativi finanziamenti per la ricerca nel campo. Si veda: «L'uso originario dell'espressione 'intelligenza artificiale' può essere attribuito a un individuo specifico, John McCarthy, che nel 1956 era assistente universitario di matematica al Dartmouth College di Hanover, New Hampshire. Insieme ad altri tre ricercatori con maggiore esperienza (Marvin Minsky di Harvard, Nathan Rochester della IBM e Claude Shannon dei Bell Telephone Laboratories), McCarthy organizzò a Dartmouth un convegno estivo sull'argomento. Parteciparono numerosi ricercatori molti dei quali avrebbero poi offerto contributi fondamentali alla materia. [...] La proposta di Dartmouth coprì una gamma sorprendentemente ampia di argomenti, tra cui le reti neurali – un antenato delle più potenti tecniche di IA di oggi –, e l'elaborazione da parte del computer del linguaggio umano» (Kaplan. 2017, pp. 31, 33).

Nel descrivere la seconda fase che «ha inizio nel 1962» – quindi, cronologicamente più vicina alle teorie di Engelbart (informatico e ingegnere elettronico, fondatore dell' Augmentation Research Center insieme ad altri ingegneri informatici e considerato l'inventore del primo mouse) e a quelle di Ted H. Nelson (padre del concetto di ipertesto) – Ricciardi adotta una classificazione cronologica che evidenzia meglio le scoperte e gli sviluppi tecnologici dell'IA nonché la progressione della crescita degli interessi e degli studi sull'IA, delineando già un abbozzo di prime ricerche caratterizzate da fasi di entusiasmo seguite da raffreddamenti.

Nel 1951, Asimov pubblica il racconto fantascientifico, dove in pratica due bambini del 2157, Tommy e Margie, scoprono «un libro vero», «antichissimo» (Asimov, 1991, p. 35) e il mondo della scuola del XX secolo fatto di aule, classi e un maestro umano. Da qui il titolo venato di nostalgia – Chissà come si divertivano! – per quel mondo trascorso tutt'altro che solipsistico.

1.1.3 Dai primi sviluppi agli "inverni" dell'IA

Nel 1963, la DARPA (Defense Advanced Research Project Agency) finanziò al MIT (Massachusetts Institute of Technology) programmi di ricerca relativi all'Intelligenza Artificiale, segnando un importante riconoscimento istituzionale della rilevanza strategica dell'IA.

Tra il 1964 e il 1966, Joseph Weizenbaum presenta ELIZA, il primo Chatbot, un software che interagiva con un utente in una conversazione simulando una seduta di psicoterapia. Si tratta di un passo significativo verso lo sviluppo di sistemi in grado di interagire con gli esseri umani attraverso il linguaggio naturale.

Nel 1965, Gordon Moore, cofondatore di Intel Corporation, formula quella che sarebbe diventata nota come Legge di Moore, stabilendo che il numero di componenti sul chip di un computer raddoppia all'incirca ogni due anni e che la velocità e la memoria dei computer crescono in modo esponenziale. Questa legge ha avuto un impatto profondo sullo sviluppo dell'IA, determinando l'evoluzione delle capacità computazionali necessarie per i sistemi di intelligenza artificiale. Lo stesso anno Ted Nelson nel corso di una conferenza nazionale usa il termine ipertesto da lui coniato nel 1963.

Un importante riconoscimento culturale dell'idea di intelligenza artificiale si ebbe tra il 1966 e il 1969 con la prima diffusione della celebre serie televisiva Star Trek, ideata da G. Roddenberry, che raccontava le vicende degli uomini del futuro e le loro avventure nell'esplorazione dell'universo a bordo dell'astronave Enterprise. L'universo fantascientifico di Star Trek, con i suoi personaggi e le sue tecnologie avanzate, ha contribuito a diffondere nell'immaginario collettivo storie di macchine intelligenti capaci di interagire con gli esseri umani. Ormai i tempi erano maturi per quella che Kaplan definisce storia intellettuale dell'IA (2017, p. 31). Una storia fatta di libri e film.

Sulla stessa linea di libri e film si colloca il film 2001: Odissea nello spazio (1968) dove Stanley Kubrick immaginava un supercomputer HAL 9000 (acronimo di Heuristically Programmed Algorithmic Computer), quale guida e controllo dell'astronave Discovery. Hal 9000 non ha forma umanoide, è annunciato da un occhio rosso ed è dotato di linguaggio e di intelligenza artificiale. Possiede dell'AI tutte le funzioni: il riconoscimento facciale e vocale, la capacità di dialogare con gli umani e di percepire emozioni (alla fine verrà depotenziato per un disguido di programmazione).

Nel 1970, Kenneth Colby (psichiatra) presenta PARRY, un Chatbot (crasi tra Chat+robot) che venne sottoposto ad una prova analoga al Test di Turing in una conversazione tra pazienti (psichiatrici) umani e virtuali. Anche in questo esperimento la simulazione non fu riconosciuta, dimostrando la crescente sofisticazione dei sistemi di elaborazione del linguaggio naturale.

Tuttavia, nonostante questi progressi iniziali, le limitazioni nella potenza di calcolo e le aspettative eccessivamente ambiziose portarono ad una prima fase di stasi negli anni '70, quando sia il governo statunitense che quello britannico ridussero i finanziamenti per la ricerca sull'IA dopo una serie di risultati deludenti.

Melanie Mitchell (2022, pp., 21-22), come si è già ricordato, fa riferimento a questo periodo definendolo "Inverno dell'IA", per i segni di sfiducia nelle potenzialità della disciplina, attribuiti, in parte, alla discrepanza tra le ambiziose promesse formulate dai pionieri dell'IA e i risultati concreti ottenuti. In realtà non ci fu un solo inverno. Mitchell situa il primo a metà degli anni Settanta in coincidenza con un certo successo dei sistemi esperti ma con mancate risposte ai progetti di più ampio respiro e di innovazione informatica che provocarono riduzioni dei finanziamenti.

La studiosa non indica date precise ma quando scrive «tutto questo è avvenuto, con diverse intensità, nel corso di cicli compresi tra i cinque e i dieci anni» aggiunge una nota personale che ne offre una piuttosto indicativa: il 1990, data del conseguimento del proprio dottorato, indicandola come il

momento in cui «questo campo stava vivendo uno dei propri inverni e si era guadagnato un'immagine così negativa che addirittura mi consigliarono di eliminare la formula 'intelligenza artificiale' dalle mie domande di lavoro» (Mitchell, 2022, pp. 22-23). Le date corrispondono a quelle più diffuse (1974-1980; 1987-1993).

Tuttavia, si può notare che, mentre l'inizio del primo inverno coincide con la prima crisi del petrolio, la fine dello stesso – la primavera –, generalmente attribuita all'invenzione dei «sistemi aperti» ("inventati" da Edward Feigenbaum) e l'inizio del secondo inverno hanno una causa affine: i sistemi aperti che si rivelarono nel giro di pochi anni troppo costosi. Da qui la breve durata dei successi. La fine del secondo inverno è invece coincidente con l'uso degli algoritmi dell'apprendimento automatico (Machine Learning) in grado di pescare nei Big Data.

Intorno agli stessi anni (1977), in Italia, un momento significativo per lo sviluppo dell'IA corrisponde alla nascita del GLIA, primo gruppo di lavoro sull'IA di Sandro Incerti. Comprende professori e ricercatori del Politecnico di Milano (tra cui Marco Somalvico), del Politecnico di Torino e delle Università di Genova, Roma, Firenze, Pisa e Salerno. Nel 1982 ci furono adesioni di altre sedi universitarie (L'Aquila, Bari, Cosenza...) e si attivarono diversi centri di ricerca (Ispra, Enea, Fiat, Csel). Era l'anno in cui L'Italsiel entrò al MIT di Boston.

Nel 1980, Edward Feigenbaum (Expert System, 1936) illustra il concetto di "sistemi esperti". È considerato il padre di questi sistemi basati sulla conoscenza e capaci di risolvere problemi in un ambito limitato. Si tratta di software che riproducono le prestazioni di una o più persone esperte in un campo specifico (medicina, economia) e funzionali nelle decisioni da prendere (es. configuratore di prodotto, consulente medico).

Nello stesso anno, John R. Searle (1932), filosofo del linguaggio e scienza della mente, pubblica un primo saggio, *Minds, Brains and Programs* (*Behavioral and Brain Sciences*, n. 3, 1980, pp. 417–424), che confutava le premesse dell'IA (forte) e dove emulava l'«esperimento mentale» realizzato da Turing nel suo Test. Il saggio è conosciuto per l'argomento della "stanza cinese" e prende lo spunto da un lavoro di Roger Schank et alia (Schank e Abelson, 1977).

In pratica si tratta di una contro-argomentazione dell'idea che la mente sia un computer e che la macchina possa pensare. Per la dimostrazione, Searle immagina sé stesso solo e «chiuso in una stanza» dove gli vengono presentati dei fogli scritti in cinese (di cui dovrebbe servirsi per realizzare un programma), una lingua che sa di non conoscere; dopo il primo invio immagina però di ricevere anche altri fogli con istruzioni in lingua inglese, la lingua che sa maneggiare e che gli permetterà di elaborare un programma di numeri e simboli formali in cinese pur continuando a non capire la lingua. Tuttavia, la sua capacità di elaborazione del programma in caratteri cinesi induce a credere che nella stanza ci sia un individuo di madrelingua.

La simulazione è riuscita, come nel Test di Turing, ma questa volta serve a dimostrare che la macchina non pensa e non capisce: riesce solo a "far finta" di capire, perché è capace di manipolare simboli e seguire regole sintattiche di un linguaggio di cui ignora il significato ovvero le regole della semantica; inoltre non è dotata di intenzionalità. Il risultato dell'esperimento mentale, secondo Searle, evidenzia che la mente umana non è una macchina computazionale perché è capace di interpretare il senso dei simboli e di manifestare intenzionalità (dato che «[g]li stati e gli eventi mentali sono letteralmente un prodotto del cervello, mentre il programma non è nello stesso modo un prodotto del computer») (Tonfoni, 1984, p. 67). Il computer può solo essere programmato per simulare processi biologici della mente ma «le manipolazioni di simboli formali di per sé non hanno alcuna intenzionalità, esse sono del tutto prive di significato» (Tonfoni, 1984, p. 66).

Ancora nel 1980, Seymour Papert (1928-2016), matematico, informatico e pedagogista, pubblica *Mindstorms*, introducendo il concetto di "costruzionismo" nelle teorie dell'apprendimento. Dopo aver collaborato (1958-1963) in Svizzera con Jean Piaget, nel 1964 lo scienziato si trasferì al MIT per insegnare e collaborare con il gruppo che si stava occupando di Intelligenza Artificiale, e in particolare, con Marvin Minsky. Papert si mostrò entusiasta delle possibilità didattiche nell'uso del computer (che definiva «oggetto per pensare» e «ambiente per l'apprendimento» – Papert, 1984, p. 30) proprio per sviluppare la naturale tendenza dei bambini a costruire la propria mente secondo una loro spontanea logica e creatività durante le prime esperienze scolastiche.

A lui è attribuita una prima definizione di «pensiero computazionale» (Papert, 1980) – poi perfezionata da Jeanette Wing – e l'invenzione di LOGO, un linguaggio di programmazione digitale destinato ai bambini e realizzato graficamente come un «robot di programmazione visuale». Il progetto, ideato sin dal 1967 (intorno alle date di "nascita" dell'IA) fu messo in opera insieme a Wallace Feurteig e Daniel Bobrow (in seguito al trio si unì la ricercatrice Cynthia Salomon), anche loro scienziati ed esperti di informatica. La realizzazione è nota anche come "tartaruga di Papert", ovvero: un'interfaccia grafica capace di concretizzare sullo schermo lo spostamento di un oggetto, applicando un insieme di procedimenti informatici (realizzati attraverso quattro movimenti: avanti, indietro, sinistra, destra).

Nel 1982, venne presentato a Tokyo dal Japan Information Processing Development Center il Japan Fifth Generation Computer System (FGCS), un supercomputer avente 1000 gigabyte di memoria. Analoghi progetti di ricerca furono iniziati in diversi altri Paesi.

Nello stesso anno, uscì nelle sale *Blade Runner*, un film di Ridley Scott ispirato a un'opera di Philip K. Dick (*Do androids Dream of Electric Sheep?*, tradotto in italiano come *Il cacciatore di androidi*). La vicenda, ambientata nel 2019 in una Los Angeles distopica, racconta di un detective privato, Rick Deckard, incaricato di dare la caccia a dei replicanti, robot dall'apparenza umana. Il film contribuì a diffondere nell'immaginario collettivo il tema del confine sempre più sfumato tra umano e artificiale.

Nel 1984, iniziò la serie di film *Terminator*, diretta da James Cameron, che aveva come antagonista l'androide cibernetico Terminator, inviato dal futuro per uccidere Sarah Connor prima che abbia un figlio (John Connor) predestinato a sconfiggere l'intelligenza artificiale, salvando l'umanità dalla ribellione delle macchine. Si deve a questa saga cinematografica il timore di una possibile rivolta delle macchine intelligenti contro l'umanità.

Nel 1988, fu costituita l'AIxIA (Associazione italiana per IA). Tra i fondatori e prima Presidente, Luigia Carlucci Aiello, professore emerito dell'Università "La Sapienza" di Roma. Considerata la «madre dell'IA in Italia», si è laureata in matematica presso l'Università di Pisa, ma lei stessa in un'intervista dichiara: «Pisa mi ha fornito le basi, Stanford grandi stimoli intellettuali: lì ho imparato a fare ricerca». A Stanford, presso l'Artificial Intelligence Laboratory, collabora con John McCarthy, «un visionario» che «voleva creare macchine in grado di pensare come gli esseri umani». Attualmente continua ad essere responsabile di progetti e di gruppi di lavoro in IA presso l'Università di Roma.

Nel 1990, John R. Searle in un altro suo scritto critico si chiedeva: *Is the Brain a Digital Computer?* (Il cervello è un computer digitale?). Il testo, edito dalla rivista "Proceedings of the American Philosophical Association", rappresenta l'ultima versione dell'intervento dedicato all'argomento della "stanza cinese". Nel suo saggio, lo studioso torna qui a muovere la sua critica filosofica all'IA forte, riassumendo i passaggi del noto esperimento. Tuttavia, nel confermare le sue premesse circa l'impossibilità di identificare la mente con un programma di computer, va oltre, ponendo di fatto il quesito scientifico su come conosce il cervello e chiedendosi se sia possibile considerare le complesse

attività cerebrali alla stregua di sofisticate operazioni computazionali. Ancora una volta la sua idea è nettamente negativa e nel suo pensiero conclusivo torna a coinvolgere l'elemento intenzionale scrivendo: «Non è possibile 'scoprire' che il cervello o qualsiasi altra cosa sia intrinsecamente un computer digitale, sebbene sia possibile attribuirgli un'interpretazione computazionale» (Searle, 1990, p. 457).

1.1.4 La rinascita e l'evoluzione contemporanea dell'IA

La riflessione di Searle sembra chiudere il periodo di crisi dell'IA. Dopo un periodo di stallo, proprio intorno agli anni '90 l'interesse per l'IA mostra già segni di ripresa. Nel giro di pochi anni troviamo nuove scoperte. Nel 1989 «Tim Berners-Lee lancia World Wide Web (www)» e produce il primo server per il World Wide Web. In seguito (1990) «scrive la prima versione del linguaggio HTML» (Hypertext Markup Language) e nel 1991 «Berners Lee pubblica il primo web al mondo, presso il CERN di Ginevra» (Ricciardi, in Ricciardi, Sacco p. 8). In pratica diventa il padre di Internet. Sebbene le origini di Internet si trovino in ARPANET, una rete di computer costituita nel settembre del 1969 negli USA da ARPA [Advanced Research Projects Agency].

Tra il 1993 e il 1997 viene fondata la RoboCup (tra i fondatori Manuela M. Veloso, Professore in scienze informatiche e specialista in robotica, presso l'Università di Pittsburgh) e avviata nel 1997. Si tratta di una iniziativa «ideata nel 1993 con l'obiettivo di realizzare, entro il 2050, una squadra di robot umanoidi autonomi in grado di sfidare e, possibilmente, battere la squadra di calcio campione del mondo» (robocup.org).

Un momento simbolico fondamentale per la rinascita dell'IA è la vittoria, nel 1997, di Deep Blue, il potente computer concepito appositamente dall'IBM per il gioco degli scacchi, contro il campione del mondo in carica Gary Kasparov. Questo evento segna la prima volta in cui una macchina è riuscita a battere il campione mondiale di scacchi, dimostrando le crescenti capacità dell'IA in domini specifici.

Lo stesso anno è data importante anche per l'altra storia (quella intellettuale e artistica). Esce Nirvana, un film di anticipazione prodotto da Gabriele Salvatores, ambientato nel mondo dei videogiochi. Nel film, un virus infetta il programma di un videogioco chiamato Nirvana, e il protagonista virtuale del gioco, Solo, scoprendo di esistere solo virtualmente, chiede aiuto all'inventore Jimi Dini di essere liberato e di cancellare il videogioco per evitare di essere replicato all'infinito.

Nella medesima scia, i fratelli Wachowski realizzarono Matrix (1999) un film di fantascienza (seguito da Matrix Revolutions nel 2003 e Matrix Resurrections nel 2021) ambientato in una realtà virtuale creata da un'intelligenza artificiale. Il protagonista, Thomas Anderson (Neo), scopre che la realtà in cui vive è frutto di una creazione fittizia ed è l'esito di una devastazione compiuta dalle intelligenze artificiali in lotta contro gli esseri umani. Il film è una sorta di favola nera della tecnologia e dà corpo a tutti i fantasmi e le paure che accompagnano da secoli i progressi del mondo della scienza. Queste paure sono state successivamente giudicate irrealistiche dal «Gruppo di esperti di alto livello sull'IA della Commissione europea» (Quintarelli, 2020, p. 20).

È del 1999 anche L'uomo bicentenario, un film di Chris Columbus tratto da un racconto di Asimov (The Bicentennial Man, 1976) e da un romanzo di Asimov scritto con Robert Silverberg (The Positronic Man, 1992). Il film narra la storia di un robot domestico, Andrew, dalle prerogative insolite, che percepisce emozioni, ha il senso dell'umorismo e sviluppa capacità non programmate, inaspettate doti artistiche e intellettuali. Andrew, in pratica, auto apprende e intende realizzare il proprio desiderio di diventare simile ad un uomo, un desiderio che impiegherà duecento anni a realizzare fino a quando, riconosciuto "membro a tutti gli effetti del genere umano", dopo aver conosciuto amore e successo, vorrà condividere con gli esseri umani anche l'ultima prova: la morte.

Nel 2001, Steven Spielberg realizza A.I. - Artificial Intelligence, basato su un progetto di Stanley Kubrick e ispirato a un racconto del 1969 di Brian Aldiss (Supertoys che durano tutta l'estate). Ambientato in un contesto post apocalittico di disastri ecologici e leggi sulla limitazione delle nascite, il film racconta la storia di David, una sorta di Pinocchio futuribile: in realtà un robot bambino in grado di provare sentimenti.

Il 21 gennaio dello stesso anno nasce Wikipedia, una enciclopedia multilingue collaborativa, Open Data, ideata da Jimmy Wales e Larry Sanger inizialmente nell'edizione in lingua inglese, e successivamente in tutte le lingue. Sebbene non direttamente legata all'IA, Wikipedia rappresenta un esempio significativo di come la tecnologia digitale possa facilitare la collaborazione e la condivisione della conoscenza su scala globale.

Nel 2004, viene presentato Io robot, un film realizzato da Alex Proyas e ispirato all'opera di Asimov. Ambientato nel 2035, un'epoca in cui i robot sono pienamente integrati nella vita quotidiana degli esseri umani e godono della loro fiducia, il film segue le indagini del detective Del Spooner, che sospetta un robot di aver contravvenuto alle leggi della robotica. Il film esplora temi etici relativi all'IA e alla possibilità che le macchine possano interpretare le leggi della robotica in modi inaspettati.

A cavallo tra il 2005 e il 2006, 50 anni dopo Dartmouth, è tempo di bilanci per valutare speranze e obiettivi centrati dai promotori (M. Minsky, C. Shannon, J. McCarthy, N. Rochester) di quell'evento che aveva dato corpo all'idea di "progettare una macchina (dunque artificiale) intelligente quanto l'uomo", realizzando il progetto di Turing. Da allora, tra "primavere e inverni dell'IA" (1974-1980), i successi in ogni campo non erano mancati, ma neppure i ripensamenti (per lo più legati al timore di un futuro popolato da automi e di uomini senza più lavoro). Le sfide di sviluppo che restavano ancora aperte continuavano a ruotare intorno alla necessità di bilanciamento della doppia natura dell'IA di presentarsi come disciplina tecnologica («finalizzata alla costruzione di macchine che permettano all'uomo una vita migliore») e come disciplina psicologica («rivolta, cioè, alla costruzione di agenti che riproducano, o che quantomeno tentino di riprodurre, alcune caratteristiche cognitive umane allo scopo di poterci rivelare qualcosa in più sui misteri del funzionamento della nostra mente»).

Nel 2006, Jeanette Wing, scienziata informatica (allora Direttrice del Dipartimento di informatica della Carnegie Mellon University), pubblica il suo «viewpoint» (Computational Thinking, wing@cs.cmu.edu), destinato agli insegnanti (L'associazione Computing At School – CAS), dove definisce il «pensiero computazionale quale competenza originale e tipica dell'informatica, utile a tutti», precisando che si tratta di «processi mentali coinvolti nel formulare problemi e le loro soluzioni presentate in una forma tale da poter essere eseguite da un agente di elaborazione dell'informazione umano o da una combinazione di uomini e macchine» (Chioccarello, pensiero computazionale.idt.cnr.it, 2015, p.7).

Nel 2011 Muore padre Roberto Busa (1913-2011), gesuita, linguista e informatico, considerato il pioniere della linguistica computazionale è stato professore presso diverse sedi universitarie, tra cui la Pontificia Università Gregoriana, l'Università Cattolica e, dal 1995-2000, presso il Politecnico di Milano, nell'ambito dei corsi di IA e robotica del Professor Marco Somalvico. Autore di molti saggi teorici, realizzò l'Index Thomisticus ovvero la "lemmatizzazione" dell'opera omnia di San Tommaso d'Aquino inaugurando il cosiddetto metodo "concordanziale", un metodo d'analisi fondamentale nella pratica delle "Digital Humanities" (umanistica informatica). L'opera, ideata nel 1946, fu pubblicata tra il 1974 e il 1980 in 56 volumi (oggi è disponibile anche su CD-ROM).

Lo stesso anno, il supercomputer Watson di IBM – del medesimo tipo di quello utilizzato per schedare le parole dell'Index Thomisticus, realizzato dallo studioso Roberto Busa – vince un quiz in TV.

Realizzato da IBM in un programma di ricerca Deep Q&A per sviluppare un software in grado di comprendere le domande e dare risposte corrette. Il computer, guidato da David Ferrucci, superò in velocità e precisione due concorrenti rivali umani, campioni del gioco in TV. Pietra miliare nel riconoscimento vocale, le capacità di Watson trovarono applicazione inizialmente nelle strutture ospedaliere per estendersi successivamente a svariati altri ambiti.

Sempre nel 2011, Sherry Turkle, sociologa, tecnologa ed esperta in psicoanalisi (Massachusetts Institute of Technology), pubblica *Alone Together* (Insieme ma soli. Perché ci aspettiamo sempre più dalla tecnologia e sempre meno dagli altri, 2012) un libro importante che riflette sull'impatto psicologico, cognitivo e sociale dell'interazione tra gli esseri umani e le tecnologie, aprendo un'area di ricerca per una nuova ricerca sul sé. Ne è derivata una descrizione degli effetti percettivi ed emotivi dovuti a uno status di connettività ininterrotta che risuona di centinaia di messaggi postati sui social, creando una illusione di vita reale. Ma in realtà si tratta di un vissuto alimentato virtualmente da miliardi di utenti che mostrano una sorta di dipendenza seduttiva nei confronti delle operazioni di connessione, al punto che il "loro secondo io" (Turkle, 1985) preferisce intrattenere rapporti solipsistici con lo schermo di una macchina piuttosto che coltivare rapporti interpersonali.

Maria Ranieri pubblica il testo *Le insidie dell'ovvio. Tecnologie educative e critica della retorica tecnocentrica* (2011), un libro critico, sì, ma nei confronti dei luoghi comuni in cui si riflette sull'ampia diffusione delle tecnologie dell'informazione e della comunicazione (ITC) focalizzando la sua attenzione sui risvolti che derivano dalle iniziative di introdurre in un ambito particolarmente sensibile ed importante come quello della formazione, dove quelle iniziative si preparano ad entrare accompagnate da dibattiti con "voci contrarie" o invece pareri ottimisti circa le «potenzialità delle tecnologie per l'educazione» (Ranieri, 2001, p. 9).

Nel 2013, Viktor Mayer-Schönberger (1966, professore presso la Oxford University) e Kenneth Cukier (1968, specialista in economia) pubblicano il volume *Big Data* (trad. Roberto Merlini, Milano, Garzanti, 2013/2017). Un testo che darà ampia diffusione alla locuzione inglese e a un concetto che indica «genericamente una raccolta di dati informatici così estesa in termini di volume, velocità e varietà da richiedere tecnologie e metodi analitici specifici per l'estrazione di valore e conoscenza» di capitale importante in statistica ed informatica. Non a caso gli studiosi ne parlano come del «nuovo petrolio» (Clive Humby, 2006) di una nuova era.

Nello stesso anno si diffonde il concetto di Open Data [dati aperti], ovvero dati resi utilizzabili da chiunque e che hanno assunto un ruolo centrale (Chignard, 2013). La locuzione è stata registrata nel Lessico del XXI secolo della Treccani nel 2013. L'*open data handbook* ne dà la seguente definizione: «I dati aperti sono dati che possono essere liberamente utilizzati e ridistribuiti da chiunque, soggetti eventualmente alla necessità di citarne la fonte e di condividerli con lo stesso tipo di licenza con cui sono stati originariamente rilasciati». Fa parte di un insieme di concetti correlati come open access (libero accesso), open science (libera scienza), open content (contenuto aperto: Wikipedia, ad es.), open source (libera fonte).

Nel 2014, Wally Pfister realizza *Transcendence*, un film sulla Singolarità e le sue possibili conseguenze negative e pericolose. Il film narra la storia del Dottor Will Caster, esperto di intelligenza artificiale, che costruisce una macchina che possiede tutta la conoscenza accumulata dall'umanità unita alla capacità di simulare la coscienza umana. Assassinato da un gruppo di terroristi tecnofobi, continuerà a vivere grazie alla moglie Evelyn che inserirà i dati delle sue ricerche dentro un supercomputer.

Il computer come macroscopio. Big data e approccio computazionale per comprendere i cambiamenti sociali e culturali (2015), è il libro del sociologo Davide Bennato che in questo suo lavoro propone

una suggestiva "metafora" di definizione del computer come ulteriore approfondimento di quelle già proposte da lui stesso in un suo lavoro precedente. La scelta questa volta è una "metafora globale" usata come formula per condurre una approfondita ricognizione degli effetti prodotti nelle nostre società dai Big data e dai nuovi mezzi di informazione (i "social media") largamente diffusi. In quest'ambito l'autore ha aperto un inedito campo di studi che coniuga scienze sociali e informatica per comprendere le dinamiche di società complesse come le nostre governate da algoritmi che sfuggono alla comprensione del grande pubblico.

Nello stesso anno, Luciano Floridi, filosofo italiano (di nazionalità britannica), professore ordinario di Filosofia ed etica dell'informazione presso l'Oxford Internet Institute della Oxford University, pubblica *The Fourth Revolution. How the infosphere is reshaping human reality* (tradotto in italiano nel 2017 *La quarta rivoluzione. Come l'infosfera sta trasformando il mondo*). Si deve a lui un'ampia revisione di teorie e metodi dei paradigmi della disciplina. Appartengono ai suoi studi concetti (e neologismi) come Iperstoria (è il termine con il quale Floridi designa le società che, per il loro funzionamento e sviluppo, come la nostra, dipendono sempre più dalle tecnologie dell'informazione e della comunicazione – ICT – Information and Communication Technologies), Infosfera (riferito alla pervasività delle tecnologie digitali) e Onlife – on+life (connesso) opposto a Offline (non connesso).

Da tutt'altro campo, Nick Bostrom, scienziato e filosofo danese, figura centrale (insieme a voci autorevoli come quelle di Kate Crawford, Manfred Spitzer, Sherry Turkle o Helga Nowotny) del dibattito allarmistico che vede nella tecnologia una entità in grado di superare e sostituire gli esseri umani, pubblica *Superintelligence: Paths, Dangers Strategies* (2014) (*Superintelligenza: tendenze, pericoli, strategie*, 2018), un libro estremamente circostanziato che alimenta l'idea che l'umanità, attualmente, si illude di addomesticare l'AI, mentre – come accade in una celebre fiaba che racconta di improvvidi passeri e di un vorace gufo (in pratica noi umani e l'AI) – non si accorge del probabile (e futuro) avvento di una superintelligenza in grado di occupare ogni ambito decisionale, rimuovendo tutte le leve del controllo umano sull'AI. Non a caso l'autore dedica il suo voluminoso resoconto all'unico passero «cieco» e inascoltato che ha percepito quel rischio sfuggito ai passeri.

A preoccupare sono i rapidi avanzamenti tecnologici. Ricordiamo che la svolta decisiva nell'evoluzione dell'IA avviene con la riscoperta delle reti neurali artificiali e l'avvento del deep learning negli anni 2000 (Boden, 2019; Mitchell, 2022).

Nel 2016, il programma Alpha GO sconfigge il campione mondiale Lee Sedoll nel gioco da tavolo considerato più difficile del mondo, il Go, un antico gioco cinese che prevede un numero pressoché infinito di mosse possibili, enormemente più grande rispetto agli scacchi. Questa vittoria è particolarmente significativa perché il gioco richiede di fare affidamento sull'intuito e Alpha Go è stato progettato per imitare la capacità di intuizione caratteristica dell'uomo, prendendo a modello le reti neurali del cervello umano.

Ancora nel 2016, appare una delle prime definizioni del concetto di AI literacy da parte di Burgsteiner et al. e Kandlhofer et al. come la "capacità di comprendere le tecniche e i concetti alla base dell'intelligenza artificiale in diversi ambiti". La concettualizzazione verrà precisata da Long & Magerko (2020) nel loro contributo intitolato *What Is AI literacy? Competencies and Design Considerations*.

Il 2017 vede la pubblicazione di *Life 3.0* di Max Tegmark, fisico teorico e fondatore del Future of Life Institute, che definì l'epoca contemporanea come fase dell'"evoluzione tecnologica", successiva all'"evoluzione biologica" – in cui l'uomo non era in grado di progettare il proprio hardware – e all'"evoluzione culturale" – che ha reso l'individuo capace di progettare gran parte del proprio

software. Nell'attuale "evoluzione tecnologica", secondo Tegmark, è possibile all'individuo «riprogettare drasticamente non solo il proprio software ma anche il proprio hardware senza tener conto dei ritmi normali di cambiamento» (Tegmark, 2018, pp. 44-45).

In questi anni, la linea degli intellettuali "preoccupati" o "concilianti" in tema di IA diventa netta. Ad esempio, Yuval Noah Harari pubblica *21 Lessons for the 21st Century* (21 lezioni per il XXI secolo) (2018), una delle voci più critiche al intorno al tema dell'eccessivo flusso di informazioni e dell'esigenza di lucidità mentale che ne consegue per non cadere nelle "acque torbide" della disinformazione pilotata. Si tratta di una riflessione critica che ritroveremo in opere successive, ma in questa Harari si domanda: «Cosa sono state e rischiano di diventare le tappe della progressione antropologico-culturale che ha reso possibile la nostra storia?» e «Cosa fa rischiare persino alla mente umana – in materia di possibilità di lavoro, rischi di discriminazione, scelte imposte, consigli subliminali, profilazioni – tutto il progresso tecnologico delle IA specie se saranno gli algoritmi a decidere per noi chi siamo e che cosa dovremmo sapere di noi stessi?» Tuttavia, Harari non si limita a saldare le sue lezioni con una visione del tutto pessimistica, ma avverte: «Ancora per pochi anni o decenni, avremo facoltà di scegliere. Se ci impegniamo, potremmo ancora indagare chi siamo davvero. Ma per cogliere questa opportunità dobbiamo farlo subito» (Harari, 2018, p. 464).

La pubblicazione del libro di Harari coincide con quella dell'Età della frammentazione. Cultura del libro e scuola digitale (2018) in cui Gino Roncaglia ripercorre tappe della cosiddetta "rivoluzione digitale" prendendo in esame gli indubbi vantaggi che derivano dalla quantità di risorse e contenuti offerti da Internet, specie nell'ambito dell'apprendimento scolastico e della formazione. Tuttavia, l'autore non manca di rilevare che il carattere frammentato e granulare dei contenuti offerti dalla rete comporta il rischio di un macroscopico appiattimento della complessità verticale "tipica della cultura del libro". Da qui la riflessione sulle strategie da usare per evitare un simile rischio, introducendo di nuovo l'attitudine e il desiderio di affrontare la complessità propria della cultura del libro anche nel mondo della nuova scuola, immaginando percorsi e iniziative di promozione di "lettura digitale" praticata con piccoli gruppi ed esperienze di lettura aumentata discutendone la natura, i limiti e la possibile efficacia dell'esperimento.

2018 esce *Teoria e pratica delle New Media Literacy*, Ranieri, M. (a cura di), con 8 contributi (Bruni L., Carioli S., Fabbro F., Nardi A., Peria L., Raffaghelli J. E. Ranieri M., Rosa A.). Si tratta di un volume teorico-pratico (basato su una ricerca sul campo della didattica in collaborazione con Istituti di scuola Primaria e secondaria coinvolgendo alunni dai 3 ai 18 anni, docenti e genitori, nonché altre associazioni) in cui Maria Ranieri fa il punto intorno ad un quesito centrale circa le «prospettive della media education» e agli "studi sulle nuove literacy (modalità di alfabetizzazione) ed implicazioni educative" con una finalità: la «costruzione di una cittadinanza digitale sicura» (Ranieri, 2018, p.10) attraverso la sperimentazione di tracciati didattici ispirati a concetto ampio di literacy e di testualità.

Nel 2019, Neil Selwyn, studioso di educazione digitale e professore presso l'Università di Melbourne, pubblica *Should Robots Replace Teachers? AI and the future of education* (I robot dovrebbero sostituire gli insegnanti? AI e il futuro dell'educazione), interrogandosi sulle problematiche e le sfide dei cambiamenti che aspettano gli educatori nella loro professione, trovandosi dinanzi alla pressante richiesta di una profonda innovazione tecnologica (con robot in classe, tutor intelligenti, e valutatori dell'apprendimento automatizzati) da controbilanciare con un'altra esigenza altrettanto forte: quella di non perdere di vista la ricchezza e la qualità emotiva del rapporto interpersonale degli insegnanti umani con i discenti. La riflessione analizza e argomenta esiti e rischi ma si chiude su una sorta di saggio patteggiamento basato sull'auspicio che l'integrazione di supporti tecnologici nell'educazione sia prospettata sempre nel rispetto di una libera scelta degli educatori.

Nel 2022 esce la traduzione italiana del libro *Le macchine di Dio*. Gli algoritmi predittivi e l'illusione del controllo [In *AI We Trust: Power, Illusion and Control of Predictive Algorithms*, 2021), ancora una voce piuttosto critica, quella di Helga Nowotny che, però, ha scelto un'angolazione peculiare in aggiunta alle riflessioni di quanti si preoccupano dei rischi e delle sfide di un'umanità immersa in un mondo digitale, sempre più coinvolto in una «traiettoria coevolutiva di esseri umani e macchina» (Prencipe in Nowotny, 2022, p. 7), centrando i suoi obiettivi sui cambiamenti della concezione di tempo, spazio e contesto sociale dovuti alle tecnologie.

Dopo averci fatto conoscere, in altri lavori, il concetto di "presente digitale esteso" – comprensivo non di solo di passato ma anche di futuro immediato (Nowotny, 1989) – la studiosa riprende le fila della sua ricerca muovendo da un quesito – «In che modo l'AI cambia la nostra concezione del futuro e la nostra esperienza del tempo?» – e si estende fino a coinvolgere le conseguenze di cambiamento socio-economico-antropologico derivate dalle «nuove relazioni temporali con coloro i quali sono fisicamente distanti ma digitalmente vicini» per cui l'umanità non sta «vivendo solo in un'epoca digitale», ma anche in «una macchina del tempo digitale» (Nowotny, 2022, p. 22). Una macchina di cui occorre conoscere gli scenari e i meccanismi, per poter realizzare quel necessario processo di reciproca evoluzione ("coevoluzione"), che permetterà all'umanità di «vivere assieme in maniera migliore» (Nowotny 2022, p. 38).

2022 Nuovo agire didattico (a cura di Rivoltella, Rossi et Alti). Il "nuovo agire" del titolo si richiama a iniziative didattiche innovative già sperimentate con ben due manuali (*Agire didattico*, 2012 e 2017) ma il nuovo testo non è una mera edizione aggiornata e il rinvio alle esperienze precedenti serve principalmente per segnare gli obiettivi mirati dai curatori di questa nuova pubblicazione. Si tratta di obiettivi basati su criteri che privilegiano in egual misura pratica e teoria mirati a una proposta d'uso innovativa: la possibilità di una lettura per "frammenti". Una lettura resa possibile dalla divisione in due parti del manuale: una prima parte di carattere applicativo (Attività di insegnamento e apprendimento) che comporta quattro sezioni e ventiquattro lezioni complessive con fulcro metodologico nel concetto di EAS (Episodio di apprendimento situato).

Il mese di novembre (2022) vede il lancio di ChatGPT di Open AI (acronimo di Chat Generative Pre-trained Transformer prodotto da Open AI). È basata su un modello di apprendimento automatico (ML) Si tratta di un software capace di conversare come e con un essere umano, in continuo aggiornamento (è già arrivata la quinta versione) ed è capace di generare testi (tradurre) grazie ad enormi quantità di dati. Con l'intelligenza generativa, secondo Paolo Legrenzi, siamo entrati nella terza fase storica dell'AI (dopo la creazione del computer – meno di un secolo fa –; la nascita della rete – nel 1989 –; e la fase storica che vede «nel corso del 2023 l'affermarsi dell'intelligenza generativa di cui ChatGPT, grazie a Nvidia H100 Tensor Core, è lo strumento principe accompagnato oggi da Gemini di Google» (Legrenzi, 2024, p.14). Ma l'autore ricorda che ChatGPT non è intelligente, può fornire risposte sbagliate o superate ed è possibile farne un uso non etico (plagio) L'opinione generale, come al solito, è divisa. A pesare è proprio il rischio di abusi proprio nell'ambito della formazione (plagio, compiti e lavori non personali). Da qui le iniziative "no plagio" e le proposte di limiti nell'uso.

Nel 2022 Ancora una voce critica (che l'autore assume in quarta di copertina come dichiarazione identitaria: «Non siamo inforg; Non siamo robot; Non siamo algoritmi; Non siamo macchine; Siamo essere umani»). È quella di Francesco Varanini che pone il suo libro (*Le cinque leggi bronzee dell'era digitale*. E perché conviene trasgredirle) sotto il segno del pensiero di Leopardi (citato anche nella chiusa "Vaghe stelle"), Goethe, Lasalle e Marx, proponendo una specie di manuale di contro alfabetizzazione digitale presentato come una sequenza di "capi di legge" da derogare per sfuggire ai rischi del macchinismo spinto dell'era digitale che promette «ciò che era inimmaginabile: ogni lavoro

svolto fino ad oggi da un essere umano può essere ora svolto da una macchina. La minaccia sviscerata gli umani. Schiacciando verso il basso le loro remunerazioni» (Varanini, 2020, p. 20).

Nel 2022, Rocco Tanica pubblica *Non siamo mai stati sulla Terra*, un romanzo scritto con la collaborazione di Out0mat-B13, ovvero un modello di deep learning sviluppato da OpenAI. Il software è «in grado di riconoscere il testo prodotto dall'autore, apprendere le caratteristiche e generare nuovi contenuti a integrazione di quelli esistenti» (Quarta di copertina del libro) e opera secondo un procedimento statistico.

Quell'anno, Federico Faggin – a lui si deve la realizzazione del microprocessore (1971-1973) – pubblica *Irriducibile*. La coscienza, la vita, i computer e la nostra natura. Un libro che entra nel vivo del dibattito su intelligenza artificiale vs intelligenza umana e sulle preoccupazioni di quanti pensano (e temono) che le "macchine ci sostituiranno". Faggin, coniugando scienza e spiritualità, per altro tematizzati in altre sue opere (ad esempio, l'ultima del 2024) avanza e discute con passione e rigorosa argomentazione scientifica (ispirata alle teorie della fisica quantistica) un'unica tesi: quella del carattere irriducibile della natura umana basata come è su coscienza, libero arbitrio e intenzionalità, per cui nessuna macchina potrà mai sostituire l'essere umano.

Nel 2023, si pubblicano diversi testi significativi sul tema dell'IA. Maurizio Ferraris pubblica l'articolo *A chi fa paura l'intelligenza artificiale*. Gli improbabili timori circa la presa di potere delle macchine, dove invita ad esercitare l'"intelligenza naturale" da parte dell'uomo nei confronti di «improbabili timori circa la presa del potere delle macchine» e a ricordare che la macchina intelligente, a differenza dell'uomo, ignora di essere macchina (e quindi è priva di autocoscienza e di intenzionalità), che può rispondere o chiedere solo sulla base di dati e costituisce una «strepitosa fonte di reddito in informatica».

Chiara Panciroli e Pier Cesare Rivoltella pubblicano *Pedagogia algoritmica*. Per una riflessione educativa sull'Intelligenza artificiale, un libro che tratta di un argomento attuale, con l'aggiunta di un orientamento molto specifico perché parla di "Pedagogia algoritmica" e intende "accompagnare, guidare" il lettore nel campo dell'IA secondo una prospettiva educativa, non guardando però solo alla dimensione della didattica digitale, bensì focalizzandosi sui meccanismi matematici che ne determinano il funzionamento. Il libro è ispirato a criteri pedagogici di cui l'IA, in ambito extraitaliano, si occupa già sotto la formula AIED (acronimo di Artificial Intelligence In Education) e sviluppa percorsi educativi in cui è previsto «Educare l'Intelligenza Artificiale» (cap. 5), oltre che «Educare con l'Intelligenza Artificiale» (cap. 3) e «Educare all'Intelligenza Artificiale» (cap. 4)» (Panciroli-Rivoltella, p. 9).

Nel marzo, Noam Chomsky pubblica sul New York Times un contributo critico sull'IA, mettendo in guardia dal rischio di «sviluppare l'etica incorporando nella tecnologia una concezione fondamentalmente errata del linguaggio e della conoscenza». Lo studioso spiega che il funzionamento della mente umana «non è, come ChatGPT e i suoi simili, un motore statistico ingombrante per la corrispondenza dei modelli, che si ingozza di centinaia di terabyte di dati ed estrapola la risposta più probabile a una conversazione o la risposta più probabile a una domanda scientifica». È bensì, «un sistema sorprendentemente efficiente e persino elegante che opera con piccole quantità di informazioni».

È del 2023 anche il libro intitolato *Tecnosofia*. Tecnologia e umanesimo per una scienza nuova di Maurizio Ferraris e Guido Saracco. Un filosofo e un tecnologo, dunque un umanista e uno scienziato sono gli autori di un libro dal titolo che «suona come un minotauro concettuale» e prende forma in un grattacielo, un concreto palazzo con visuale sulle officine Fiat e abitato per «i primi trentatré anni della [sua] vita» da Guido Saracco. Un palazzo di nove piani, tutti indicati nell'immagine del

grattacielo che apre il libro, una figura che, però, assume la dimensione metaforica del cammino che i due autori intendono percorrere per attuare il loro progetto di "scienza nuova" – la Tecnosofia – superando le sfide e i pregiudizi che ostacolano l'unione a lungo negata tra (scienza) tecnologia e umanesimo.

A settembre del 2023 L'UNESCO pubblica la prima Guida per l'Intelligenza Artificiale generativa nell'educazione e nella ricerca (Generative AI and the future of education), un testo guida all'uso dell'intelligenza Artificiale Generativa (GenIA o GenAI) in settori specifici. Il documento è accompagnato da un Preambolo a cura della Vicedirettrice generale dell'UNESCO Stefania Giannini che descrive le potenzialità della GenAI, delle sue capacità di creazione automatica di contenuti in testi e immagini e ne esamina gli effetti specie nel campo dell'insegnamento e dell'apprendimento invitando a "ripensare perché, cosa e come apprendiamo" e toccando questioni di sicurezza, privacy, da cui derivano le ragioni e l'urgenza di regolamentazione ufficiale.

«Su ChatGpt prestiamo attenzione, ma senza allarmismi. Parla Cristianini. Intervista di Simona Sotgiu a Nello Cristianini, professore di IA preso l'Università di Bath (di recente ha pubblicato *La scorciatoia. Come le macchine sono diventate intelligenti senza pensare in modo umano*, 2023; *Machina Sapiens. L'algoritmo che ci ha rubato il segreto della conoscenza e Sovrumano. Oltre i limiti della nostra intelligenza*, 2025). Il contributo dell'intervista va nello stesso senso dell'opinione non allarmistica espressa dal filosofo Maurizio Ferraris, ma prende una piega personale richiamando alcune tematiche che sono a fondamento del suo primo libro e ruotano intorno al presupposto di un concetto esteso di intelligenza. Cristianini pensa che ci siano tanti modi di essere intelligenti e quindi ci siano molte intelligenze, "intelligenze aliene": naturali, non solo umane ma vegetali, animali, compresa quella artificiale che, nella versione più recente, facendo uso di metodi statistici (grazie a una massiccia raccolta di dati), ha adottato un linguaggio proprio della macchina, abbandonando l'idea di emulare il comportamento e il modo di pensare umani.

Nel dicembre 2023, nel suo Messaggio di fine anno, il Presidente della Repubblica Italiana Sergio Mattarella dichiarava: «La tecnologia ha sempre cambiato gli assetti economici e sociali. Adesso, con l'intelligenza artificiale che si autoalimenta, sta generando un progresso inarrestabile, destinato a modificare profondamente le nostre abitudini professionali, sociali, relazionali. Ci troviamo nel mezzo di quello che verrà ricordato come il grande balzo storico dell'inizio del terzo millennio. Dobbiamo fare in modo che la rivoluzione che stiamo vivendo resti umana. Cioè, iscritta dentro quella tradizione di civiltà che vede, nella persona, e nella sua dignità, il pilastro irrinunciabile».

Sempre nel 2023, Gino Roncaglia pubblica *L'architetto e l'oracolo. Forme digitali del sapere da Wikipedia a ChatGPT*, in cui si concentra «su due modelli [...] di organizzazione delle conoscenze: quello architettonico, basato sull'idea di un sapere fortemente strutturato [...] Il primo modello è sostanzialmente quello della tradizione enciclopedica [...] Wikipedia ne è la manifestazione più nota [...] la biblioteca. Il secondo modello ha invece il suo esempio paradigmatico nelle reti neurali che hanno portato allo sviluppo delle AI generative come ChatGPT [...] spesso paragonate a oracoli, basati su associazioni statistico/probabilistiche e questa contrapposizione fra modello architettonico deterministico e modello oracolare-probabilistico è alla radice anche del libro» (Roncaglia, 2024, p. XV).

Nel 2024, Paolo Benanti, teologo, presbitero terziario dell'ordine di San Francesco, professore presso la pontificia Università Gregoriana e presso l'Università di Seattle, già consigliere di Papa Francesco sui temi dell'IA, dell'etica e della tecnologia, viene nominato Presidente della Commissione sull'IA per l'informazione e l'editoria della Presidenza del Consiglio dei ministri (detta "Commissione algoritmi"). È sua la formula "Algoetica", che coinvolge l'inedita relazione uomo-macchina del

nostro universo digitale e che Benanti evoca come necessità di conciliare «valori numerici con valori etici».

Nel giugno 2024, Papa Francesco partecipa alla sessione del G7 sull'intelligenza artificiale a Borgo Egnazia (Puglia) con un discorso intitolato *Uno strumento affascinante e tremendo*. Nel suo intervento, il Pontefice riflette sugli effetti dell'intelligenza artificiale circa il futuro dell'umanità, sottolineando potenzialità positive e rischi connessi a questa tecnologia.

Sempre nel 2024, Nello Cristianini pubblica *Machina Sapiens*. L'algoritmo che ci ha rubato il segreto della conoscenza, un libro che invita i lettori a percorrere uno straordinario viaggio nel complicato mondo dell'intelligenza artificiale, partendo da ChatGPT, l'attuale «abile conversatore» (Cristianini, 2024, p.7) che ci è diventato familiare, procedendo a ritroso, e ponendo al centro la domanda «Le macchine possono pensare»? Era la domanda che apriva il noto saggio (*Computing machinery and intelligence*, 1950) dello scienziato Alan Turing, primo ideatore dell'IA, con il quale l'autore intreccerà una sorta di lungo e ragionato dialogo. È quella domanda a fare da guida a una vera e propria storia dell'intelligenza artificiale, che ripercorre le prime scoperte computazionali fino agli esiti attuali. L'autore parla di storia e ne nomina i tre principali protagonisti: gli scienziati (coloro che hanno inventato l'IA), gli utenti (che sono a contatto con essa) e le macchine (gli agenti intelligenti, potenzialmente sempre più creativi e pericolosamente più autonomi). La definizione di protagonisti è del tutto giustificata, infatti come in un racconto è intorno ad essi che Cristianini ricostruisce le tappe della storia dell'AI, dalle scoperte di Turing, fino all'approdo con un "nuovo mondo" intelligente, aperto da ChatGPT e dalla potenzialità dei suoi diversi modelli di linguaggio (i LLM attuali che permettono alle macchine di conversare con noi: ChatGPT, Bard, Eliza, Gemini). L'autore conclude con una riflessione provocatoria: «Quando anche le macchine possono pensare, allora noi chi siamo? Siamo pronti ad incontrare un Oracolo che combina tutte le conoscenze di tutti i libri mai scritti? Questo sta diventando rapidamente possibile, e presto dovremo decidere se vogliamo continuare su questa strada o fermarci», a cui l'autore stesso dà l'unica risposta possibile: «Alla fine, credo di sapere cosa faremo, perché so che siamo della specie di Pandora e Prometeo. Dove saremmo se non avessimo giocato con il fuoco» (Cristianini, 2024, p. 146).

Andrea Prencipe (economista, ex-Rettore dell'Università Luiss) pubblica il testo *Università generativa* (2024) un titolo inconsueto, una formula innovativa e un libro con vasti orizzonti teorico-pratici. È diviso in due parti: la prima è preceduta da una riflessione che si concentra su una cronistoria dei modelli educativi per finire abbracciando i percorsi di approdo all'Università generativa, descrivendo le tecnologie che la supportano e soffermandosi a trattare sfide e opportunità che derivano dall'ingresso delle AI nelle Università. La Seconda parte descrive modalità e obiettivi dell'Università generativa, soffermandosi sulle questioni di internazionalizzazione, interdisciplinarietà e innovazione (presentata quest'ultima come fattore «trainante del progresso»). Il libro si chiude con una provocazione (*Imparare a disimparare* ovvero un invito a non fermarsi dinanzi alla necessità di continuare a confrontarsi con il sapere). Nella presentazione del volume si parla di una missione che «si concretizza nel conservare le qualità 'umane' del pensiero, sapientemente integrate con le opportunità offerte dalla tecnologia». Tra gli obiettivi mirati figura l'auspicio che le Università recuperino una centralità sociale orientando la loro attività educativa verso «il 'saper essere' e il 'saper diventare' prospettando una «crescita personale e professionale continua» se è vero che «la formazione universitaria [...] costituisce una pratica del 'saper essere al mondo', oltre che un processo di creazione di strumenti per conoscere e costruire futuri, e quindi 'saper diventare'» (Prencipe, 2024 p. 28).

Il 13 marzo 2024, il testo del documento – denominato EU AI Act – che determina il progetto di regolamentazione europea dell'IA, ideato sin dal 2021, è stato adottato dal Parlamento europeo. L'AI Act è in vigore dal 2 agosto del 2024 e sarà pienamente applicabile dal 2026. L'applicazione di alcune parti del Regolamento verrà però anticipata al 2025. Si tratta di un documento molto dettagliato di 419 pagine e 13 allegati la cui peculiarità consiste nell'essere proposto come "Regolamento applicativo dei diritti e dei doveri in materia di IA". I temi ricorrenti non contemplano proposte o direttive relative educazione/formazione, bensì ribadiscono generali esigenze e preoccupazioni etiche, di trasparenza, sostenibilità, affidabilità, tutela dei diritti e illustrano regole armonizzate sull'intelligenza artificiale e piani di modifica dei precedenti regolamenti dell'Unione Europea che fanno tutti riferimento alla necessità di una supervisione umana.

Il documento segue un approccio basato su quattro livelli di rischio (inaccettabile, alto, limitato, minimo o inesistente) associati all'uso. Il rischio inaccettabile comporta il divieto dell'impiego per diversi scopi. Ad esempio: manipolare le persone, sfruttarne la vulnerabilità, raccogliere dati sulle persone al fine di valutare (prevedere) la possibilità che commettano un reato; creare o ampliare banche dati di riconoscimento facciale; classificare le persone sulla base dei loro dati biometrici per trarre deduzioni o inferenze in merito a razza, opinioni politiche, convinzioni religiose o filosofiche, vita sessuale o orientamento sessuale; in alcuni casi è proibita l'identificazione biometrica remota in tempo reale (è il sistema che permette di identificare persone fisiche a distanza mettendo a confronto i dati biometrici di una persona con i dati biometrici contenuti in una banca dati di riferimento. Il divieto, ovviamente, non vale in caso di ricerca di persone scomparse, atti criminali).

Ecco un estratto del Preambolo dell'EU AI Act: (1) «Lo scopo del presente regolamento è migliorare il funzionamento del mercato interno istituendo un quadro giuridico uniforme in particolare per quanto riguarda lo sviluppo, l'immissione sul mercato, la messa in servizio e l'uso di sistemi di intelligenza artificiale (sistemi di IA) nell'Unione, in conformità dei valori dell'Unione, promuovere la diffusione di un'intelligenza artificiale (IA) antropocentrica e affidabile, garantendo nel contempo un livello elevato di protezione della salute, della sicurezza e dei diritti fondamentali sanciti dalla carta dei diritti fondamentali dell'Unione Europea ("Carta") compresi la democrazia, lo Stato di diritto e la protezione dell'ambiente, proteggere contro gli effetti nocivi dei sistemi di IA nell'Unione, nonché promuovere l'innovazione. Il presente regolamento garantisce la libera circolazione transfrontaliera di beni e servizi basati sull'IA, impedendo così agli Stati membri di imporre restrizioni allo sviluppo, alla commercializzazione e all'uso di sistemi di IA, salvo espressa autorizzazione del presente regolamento. [...] (2) Il presente regolamento dovrebbe essere applicato conformemente ai valori dell'Unione sanciti dalla Carta agevolando la protezione delle persone fisiche, delle imprese, della democrazia e dello Stato di diritto e la protezione dell'ambiente, promuovendo allo stesso tempo l'innovazione e l'occupazione e rendendo l'Unione un leader nell'adozione di un'IA affidabile.» (1PE-24-INIT_it.pdf, p. 1).

1.1.5 Riflessioni critiche e prospettive future

L'evoluzione dell'IA, caratterizzata, come si è spesso ricordato, da epoche di euforia ("primavere") e di disillusione ("inverni"), ha generato non solo aspettative per le sue potenzialità, ma anche preoccupazioni per i suoi prevedibili riflessi sociali, etici ed economici. In ogni caso, è vero che «L'IA è arrivata ed è qui tra noi» (Ricciardi, 2019, p. 5); ed è anche vero, come osserva Cristianini (2023), che le macchine intelligenti fanno già tutto (o quasi) in maniera "sovrumana".

Le preoccupazioni coinvolgono diversi ambiti, dalla privacy alla sorveglianza, dai bias algoritmici alla disoccupazione tecnologica. Nel 2021, il commissario ONU per i diritti umani, Michelle De Bachelet, ha richiesto «un'azione dei paesi affiliati per sospendere lo sviluppo degli strumenti di

Intelligenza Artificiale, almeno sino a quando non siano ben chiari i limiti e la sua controllabilità e non sia possibile tutelare i diritti personali».

Il dibattito contemporaneo sull'IA rimane comunque piuttosto vivace e fa emergere diverse posizioni. Alcuni studiosi, come Kate Crawford, Manfred Spitzer, Yuval N. Harari, Sherry Turkle, Helga Nowotny e altri, vedono nella tecnologia un'entità potenzialmente in grado di superare e sostituire gli esseri umani. Nowotny (2022, p. 38) indicava, invece, la via conciliante di una "coevoluzione" tra macchina e umanità nell'idea di «vivere assieme in maniera migliore».

Ma è ancora l'ottica critica basata sulla necessità di riscatto dell'essere umano a prevalere nelle riflessioni che Mauro Crippa e Giuseppe Girgenti esprimono in *Umano poco umano. Esercizi spirituali contro l'intelligenza artificiale* (2024). Nel loro libro, Crippa e Girgenti propongono una sorta di prontuario di contravveleni per esorcizzare la pervasività che l'intelligenza artificiale ha acquisito in ogni campo, soprattutto paventando il pericolo che l'umanità si stia avviando verso lo scenario di un universo «postumano digitalmente modificato» (Crippa e Girgenti, 2024, p. 20). I due autori invitano a «fare di necessità virtù [...]» e a «capire chi siamo per difendere la nostra identità». La tesi centrale del libro consiste nell'idea che si non si debba temere che la nostra intelligenza sia presto superata da quella artificiale, bensì che l'uso dell'AI indebolisca la mente e l'anima degli esseri umani. I nove esercizi spirituali indicati nel titolo seguono un percorso di "ritorni guidati" alla radice di noi stessi in compagnia di figure esemplari (Socrate, Platone, Epicuro, Galeno, Sant'Agostino, Omero, Eraclito, Aristotele, Sant'Ignazio di Loyola).

Altri studiosi, come Maurizio Ferraris, adottano una posizione più mediana. Ferraris, nell'articolo ricordato sopra, interviene sul timore diffuso che «le macchine [abbiano] la tendenza a portare via i lavori agli umani», spiegando: «Niente di male, per carità, se un lavoro può essere automatizzato in genere è indegno di un umano. Ma sappiamo che finora il salario che retribuiva il lavoro è stata la più diffusa maniera per ridistribuire i beni all'interno di una società. Dobbiamo trovarne un'altra: concentriamoci su questo, e lasciamo perdere tutte le vane fantasie sul Golem che prenderà il potere».

Mentre Cristianini, come si è già detto, tendendo a sottolineare il fatto che esistono tanti modi di essere intelligenti e quindi che ci sono molte intelligenze, compresa quella artificiale, fa notare come, in modo innovativo, nella versione più recente, grazie all'uso di metodi statistici (e una massiccia raccolta di dati), l'IA abbia adottato il linguaggio proprio della macchina, al punto da giustificare l'interrogativo: «È possibile che le macchine comprendano cose che noi non potremo mai capire?» (Cristianini, 2024, p. 8).

Tutto questo delinea un contesto di prospettive future fatto di trasformazioni e di cambiamenti, dove sarà necessario continuare a confrontarsi con concetti come Algocrazia, Algoritica, Singolarità - come fanno Floridi e Cabitza (2021) - proponendo soluzioni ridisegnate ontologicamente da una «Intelligenza ibrida» (basata sull'idea di una collaborazione responsabile tra intelligenza umana e intelligenza artificiale, sempre più diffusa e automatizzata).

Le notizie degli anni più recenti (dal 2008 in poi) appaiono sempre più intrecciate con il successo delle invenzioni tecnologiche (come lo smartphone – responsabile del moltiplicarsi delle connessioni in rete), lo sviluppo degli eventi finanziari di Silicon Valley e il moltiplicarsi di società e di veri e propri giganti digitali (EBay, Twitter, Skype, Netflix, LinkedIn, Uber e Airbnb...). Ricciardi nel suo libro evocava quanti parlavano dell'esistenza di un vero e proprio «capitalismo delle piattaforme», facendo notare come «[i] nomi da cui non si può prescindere sono pochissimi. Comprare? Amazon. Connettersi? Apple. Cercare? Google. Comunicare? Facebook (o Whatsapp). Il risultato è scritto nei

numeri. Nel 2005, tutte assieme, le quattro aziende fatturavano 28,7 miliardi di dollari. L'anno scorso (2018) sono arrivate a 350 miliardi» (Ricciardi, in Ricciardi e Sacco, 2019, p. 12).

Nel momento di fare un bilancio dello stato di cose determinate dall'"arrivo dell'IA tra noi" lo studioso ci presentava un quadro sociale culturalmente cambiato, governato da una economia ricca, praticata da pochi, dove il valore dipende esclusivamente dai mezzi di connessione ed è reso possibile attingendo alle informazioni fornite da ogni singolo utente, che, di fatto, diviene «consumatore e collaboratore a titolo gratuito» (Ricciardi, in Ricciardi e Sacco, 2019, p. 13). Il quadro tuttora valido e più fitto di iniziative social o altro. In definitiva ora come allora siamo «noi stessi il prodotto che le aziende digitali 'vendono' sul mercato», sulla base di algoritmi di «profilazione personale» che ci identificano e permettono di essere raggiunti da pubblicità e offerte di acquisto.

Ricciardi sottolineava anche la piega negativa di una possibile deriva di controllo da parte di un «sistema centralizzato» che pretende un «patto di credenza» e l'abdicazione della propria intelligenza in favore di quella artificiale. Da qui la riflessione finale di quel suo scorcio di visione storica dell'IA che, invece di toccare le questioni di privacy, o di minaccia per il lavoro umano, si soffermava a valutare la distanza dall'idea di Turing circa la capacità di inventare macchine in grado di «distribuire intelligenza in forma collaborativa a una platea di umani la più vasta e più consapevole possibile» (Ricciardi, in Ricciardi e Sacco, 2019, p. 15). Un'idea utopica, secondo Ricciardi, quanto, si potrebbe aggiungere, quella di «intelligenza collettiva» capace di «valorizzazione tecnica, economica, giuridica e umana», immaginata da Pierre Lévy.

Il concetto di "intelligenza collettiva", in realtà, rimane un concetto fertile, considerato fondante in molti ambiti, compreso quello della formazione. Ad esempio, è indicato quale presupposto ampiamente commentato da Maria Ranieri e Ubaldo Fadini che lo ritengono essenziale per «promuovere la costruzione collaborativa di conoscenza» (Formazione e cyberspazio. Divari e opportunità nel mondo della rete, 2006, p. 6), grazie anche all'uso critico della rete.

Tuttavia, parallelamente all'evolversi di queste riflessioni, emerge con sempre maggiore forza la necessità di un'educazione all'IA, o "AI literacy", come aspetto critico dell'educazione contemporanea. L'AI literacy comprende un ampio insieme di competenze, includendo non solo la conoscenza tecnica dei meccanismi dell'IA, ma anche la capacità di valutare criticamente i vari riflessi nella società e nella cultura, considerazioni etiche e applicazioni nell'apprendimento permanente (Long & Magerko, 2020; Miao et al., 2021).

Il rapporto Future of Jobs del World Economic Forum riferisce che l'IA rimodellerà significativamente il mercato del lavoro, spiazzando molti ruoli ma creandone anche di nuovi (World Economic Forum, 2020). Una prospettiva possibile in ambito scolastico è descritta da Pacchioni (2021, p. 169), quando immagina una didattica futura svolta in aule universitarie dove «ci saranno delle meravigliose proiezioni olografiche di un celeberrimo professore di Harvard, il quale presenterà un corso, disponibile in ben 139 idiomi e dialetti, tratto da uno sconfinato catalogo di offerta culturale in formato elettronico, il tutto arricchito di stupendi contenuti multimediali e di avvincente realtà virtuale. Qualche società offrirà questo servizio, ovviamente in modo gratuito. Si chiamerà Google Knowledge o Facebook University o roba del genere». Con qualche adattamento, lo scenario è un po' simile al contesto scolastico immaginato da Asimov per Margie e Tommy (1951). Tutto questo espone i ricercatori dinanzi alla necessità di un continuo adeguamento all'inedito modo dell'IA di fare innovazione.

Indubbiamente questo nostro abbozzo storico-bibliografico presenta lacune in materia di date, dati, informazione ed esaustività. Ma su un aspetto la nostra ricerca rimane un'esperienza che ci ha

insegnato molto perché ci ha permesso di renderci conto – specie agli inizi di questa ricerca –, che è la stessa AI una fonte che si autoalimenta e si storicizza tramite i propri motori di ricerca (Google) e i vari siti. Tuttavia, tali fonti di informazione sono in continuo cambiamento, per cui dati e notizie sembrano avere una validità limitata nel tempo e devono essere continuamente rivisti. I quadri storici di Melanie Mitchell, Mario Ricciardi e Sara Sacco, Stefano Quintarelli et al., Margaret Boden, Richard Urwin e altri, che abbiamo scelto per seguire percorsi già tracciati – aggiungendo elementi di integrazione –, ci sono stati utili proprio per non incorrere nel rischio di irrilevanza e provvisorietà di dati a cui si espone ogni tentativo di storicizzazione dell'IA (delle AI), se è vero, come è vero, lo stato di complessità a cui fa riferimento il pensiero dello studioso Paolo Benanti:

«Forse la verità straordinaria dell'età moderna è che certi tipi di tecnologia avanzano non in modo lineare, ma su curve esponenziali [...] Ogni anno una parte sempre più ampia del mondo della tecnica viene risucchiato in queste curve esponenziali. A grandi linee, ciò significa che ogni anno vede più innovazione rispetto a tutti gli anni prima messi insieme. Ciò implica che i prossimi vent'anni presenteranno cambiamenti tecnologici così profondi da rendere quasi irrilevante tutto ciò che è avvenuto prima. Questa velocità ci interroga e lo scenario diviene così complesso che sfida la nostra capacità di comprendere la tecnologia e i suoi prodigi» (Benanti, 2018, p. 9).

E l'irrilevanza di "tutto ciò che è avvenuto prima" non riguarda solo il perfezionarsi delle tecnologie e l'avanzamento del loro uso presso gli individui e le società, bensì anche il banco della riprova storica. Tuttavia, non ci sembra sbagliato pensare che l'unica alternativa possibile fosse proprio quella di tentare comunque di intraprendere una storia fatta di continui ricominciamenti in modo da tenere il passo con le idee, gli eventi, e i libri. Questo ci ha incoraggiati a proporre una visione più generale, allargata alla "storia intellettuale" e ai dibattiti che hanno accompagnato e accompagnano tuttora l'evolversi della rivoluzione AI.

Ad esempio, oltre alle note fasi di "AI Winters" da una certa data in poi (intorno agli anni 2000 e, in crescendo, fino al "lancio" dell'AI generativa, che ne rappresenta il punto di arrivo più recente) notiamo che nella scansione cronologica abbiamo dovuto prendere in considerazione oltre alle date-evento, relative ai progressi delle innovazioni tecniche, dei processi di sviluppo dell'importanza scientifica delle ICT (Information and Communications Technology) nonché dei successi applicativi dell'AI in ogni campo, anche le date del diffondersi di un fenomeno che A. Elliott ha definito "Cultura dell'intelligenza artificiale" (Elliott, 2019). Una cultura caratterizzata da una vera e propria fioritura di metafore teoriche ("Quarta rivoluzione", "Infosfera", "Onlife", "Iperstoria": Floridi; "Età della frammentazione", "Architetto", "Oracolo": Roncaglia; "Algoretica": Benanti; "Pedagogia algoritmica": Panciroli e Rivoltella; "Tecnosofia": Ferraris e Saracco; "Intelligenza inorganica aliena": Harari; "Sesta svolta": Bennet; "Sovrumana", "Intelligenza aliena": Cristianini; "Intelligenza collettiva": Lévy –; "Intelligenza condivisa": Mollick; "Intelligenza ibrida": Cabitza).

Da qui l'esigenza di fare riferimento anche a opere che hanno dato spessore a quella cultura, l'hanno definita, spiegata, introdotta in ambiti inediti (come quello della scuola e dell'educazione per le future generazioni) e resa per quanto possibile familiare, "discussa" dai pareri dei teorici (Chomsky, Selwyn, Faggin, Harari, Bostrom, Turkle, Spitzer, Nowotny, Varanini, Grippa, Girenti e altri studiosi...), intenti a riflettere criticamente sulla possibilità di operare scelte non etiche nell'uso dell'AI, accentuandone i rischi e i risvolti negativi in ambito filosofico/cognitivo/concettuale e socioeconomico, specie per ciò che concerne il rapporto uomo-macchina. Un rapporto visto sempre più dipendente e sbilanciato a favore della capacità e rapidità di calcolo dei sistemi tecnologici che comporta promesse – impensabili fino a qualche decennio fa – ma anche molte sfide etiche per l'intera umanità. Una umanità chiamata a scegliere di confrontarsi responsabilmente con macchine che

sembrano diventare sempre più intelligenti e dove i dilemmi, circa la minaccia di perdere ogni controllo umano, non sono poi molto diversi da quelli che poneva il Cardinale Martini, quando osservava:

«Sono così minacciose tutte le tecnologie del virtuale? L'intero cammino verso l'intelligenza artificiale finirà per svalutare il valore della persona, riducendola a pura meccanica? O, invece, saranno i valori dell'uomo a indurre la scienza ad aprire nuovi fronti grazie alle conquiste tecnologiche? [scenario questo] molto incoraggiante, purché l'intelligenza umana rimanga padrona dei processi» (Cardinale Martini citato da Cabitza, in Floridi e Cabitza, 2021, p. 24).

Crediamo che nei prossimi decenni, simili dubbi e alternative in cerca di risposta continueranno ad avere larghissima eco e alimenteranno durevolmente il confronto e i dibattiti tra gli specialisti della storia delle AI che resta ancora da scrivere.

La storia dell'IA, con le sue primavere e i suoi inverni, si è determinata lungo un percorso complesso e discontinuo, fatto comunque di progressi scientifici e riflessioni etiche che ci ha messo davanti un quadro in rapida trasformazione tecnologica e di crescente rilevanza dell'IA. Si tratta di un contesto che ci vede spesso dubbiosi e impreparati, dove attualmente risulta anche fondamentale chiedersi a che punto sono le disposizioni e le regole utili per rispondere alla necessità di realizzare quella "intelligenza ibrida" evocata sopra, che renderà possibile la collaborazione tra macchina e umani. Una collaborazione da acquisire mediante un serio processo di preparazione ovvero una vera e propria AI literacy, una materia relativamente recente, che intendiamo approfondire e di cui ci occuperemo nelle sezioni successive di questo capitolo, esaminando la letteratura esistente intorno a questo tema estremamente nuovo, descrivendone lo status, interrogandoci sulla natura del concetto e le prospettive del campo delle ricerche che lo riguardano.

1.2 Revisione sistematica della letteratura sull'AI literacy

1.2.1 Contesto e rilevanza dell'AI literacy

Come abbiamo visto negli ultimi vent'anni l'Intelligenza Artificiale (IA) è passata da ambito di ricerca specialistico a tecnologia-infrastruttura che informa il tessuto socio-economico globale. Algoritmi di apprendimento automatico, visione computazionale e elaborazione del linguaggio naturale alimentano servizi, oggetti e processi che usiamo quotidianamente – dalla selezione di notizie online all'assistenza sanitaria preventiva – e, di conseguenza, plasmano opportunità, rischi e competenze richieste a cittadini, lavoratori e decisori politici.

Oggi, l'IA non è solo uno strumento ma parte integrante di un ecosistema digitale in continua evoluzione. Le moderne tecnologie di IA, come l'apprendimento automatico, la percezione sensoriale, l'elaborazione del linguaggio naturale e la visione artificiale, si basano su enormi quantità di dati e potenza di calcolo, rese possibili dall'ascesa del cloud computing, dell'analisi dei big data e di architetture hardware avanzate (Russell & Norvig, 2021; Wang, 2019; Zawacki-Richter et al., 2019). Il passaggio da sistemi basati su regole a sistemi adattivi guidati dai dati ha trasformato il modo in cui l'IA interagisce con il mondo, permettendo alle macchine di apprendere dai dati, riconoscere schemi e prendere decisioni con un intervento umano minimo. Questa evoluzione ha posizionato l'IA al centro di industrie come la sanità, l'istruzione, il marketing e i sistemi autonomi, dove le sue applicazioni continuano a espandersi (Chen et al., 2020; Hwang et al., 2020).

La rapida integrazione dell'IA in questi campi ha generato l'emergere dell'AI literacy come aspetto critico dell'educazione contemporanea. Gli studiosi sostengono che l'educazione all'IA dovrebbe essere considerata con lo stesso livello di importanza delle competenze di alfabetizzazione tradizionali come la lettura e la scrittura, a causa della sua ampia natura interdisciplinare e della crescente influenza (Kandlhofer et al., 2016). L'AI literacy comprende un ampio insieme di competenze, includendo non solo la conoscenza tecnica dei meccanismi dell'IA, ma anche la capacità di valutare criticamente i suoi impatti sociali, considerazioni etiche e applicazioni nell'apprendimento permanente (Long & Magerko, 2020; Miao et al., 2021). Queste competenze sono essenziali per interagire efficacemente con i sistemi di IA e per prendere decisioni informate sulla loro governance e applicazione all'interno del più ampio panorama digitale.

Con l'evoluzione continua dell'IA, il rapporto Future of Jobs del World Economic Forum suggerisce che entro il 2025, l'IA rimodellerà significativamente il mercato del lavoro, spiazzando molti ruoli ma creandone anche di nuovi (World Economic Forum, 2020). Tuttavia, la maggior parte degli individui interagisce con tecnologie guidate dall'IA senza comprendere appieno i principi sottostanti o le sfide etiche, come le preoccupazioni per la privacy, i bias e la sorveglianza, che accompagnano queste innovazioni (Burgsteiner et al., 2016; Ghallab, 2019). Data la profondità di questi impatti, c'è una crescente enfasi sull'essere "AI literate" – una competenza che permette agli individui di impegnarsi criticamente con i sistemi di IA e facilita la collaborazione e la comunicazione efficace con l'IA (Long & Magerko, 2020; Ng et al., 2021).

Come moderno contraltare alle alfabetizzazioni fondamentali come la digital literacy e la media literacy (Calvani et al., 2012; Gilster, 1997; Livingstone, 2004), l'AI literacy sottolinea l'importanza di comprendere e interagire con l'IA nel contesto del più ampio ambiente digitale in rapida evoluzione. Specialmente per i non esperti, aumentare l'AI literacy è cruciale quanto formare esperti di IA, poiché questo gruppo probabilmente utilizzerà, collaborerà e coesisterà insieme alle tecnologie di IA nella vita quotidiana (Ng et al., 2021).

In questo contesto di rapida trasformazione tecnologica e crescente rilevanza dell'IA, risulta fondamentale esaminare la letteratura esistente sull'AI literacy per comprendere come questo concetto è stato definito, come si è evoluto nel tempo e quali sono le principali prospettive di ricerca in questo campo. La presente revisione sistematica della letteratura si propone di rispondere a queste domande, offrendo una panoramica critica e approfondita dello stato dell'arte della ricerca sull'AI literacy.

Alla luce di queste dinamiche, diventa essenziale comprendere come la comunità scientifica definisca, studi e promuova l'AI literacy. La presente revisione sistematica si propone di colmare questa lacuna ponendosi tre domande di ricerca esplicite:

RQ1 – Sviluppo e diffusione: Quali sono le tendenze temporali e la distribuzione geografica delle pubblicazioni sull'AI literacy e quali gruppi target (studenti, docenti, cittadini, professionisti) vengono maggiormente indagati?

RQ2 – Definizioni: In che modo la letteratura accademica definisce il construct di AI literacy e quali dimensioni/competenze vengono comunemente incluse?

RQ3 – Temi emergenti: Quali macro-temi, metodologie di ricerca e prospettive future emergono dagli studi esistenti sul tema?

1.2.2 Metodologia della revisione sistematica

La revisione sistematica della letteratura sull'AI literacy è stata condotta seguendo il framework metodologico stabilito da Cooper e colleghi (2019), in conformità con le linee guida PRISMA (Moher et al., 2009). L'approccio metodologico adottato è stato sistematico e rigoroso, mirato a ridurre bias ed errori, producendo così risultati affidabili.

Il processo di revisione sistematica è stato articolato in cinque fasi distinte: a) definizione del problema e delle domande di ricerca, b) identificazione della letteratura rilevante, c) valutazione dei risultati della ricerca, d) categorizzazione, codifica e sintesi dei risultati, e) disseminazione dei risultati della revisione. La ricerca ha coinvolto diverse banche dati accademiche, tra cui Web of Science (Clarivate), ERIC (Institute of Education Sciences), IEEE Xplore (IEEE) e SCOPUS (Elsevier), ed è stata condotta nel marzo 2023, con un successivo aggiornamento nel settembre 2023.

La strategia di ricerca è stata elaborata con cura per catturare il più ampio spettro possibile di pubblicazioni rilevanti sull'AI literacy. Ciò ha portato alla seguente stringa di ricerca completa: "TITLE-ABS-KEY (((("artificial intelligence literac*" OR "AI literac*" OR "AI awareness" OR "AI readiness"))) AND NOT ("machine learning") AND NOT ("computer-based learning") AND NOT ("intelligent tutoring system"))". La ricerca è stata limitata agli ultimi 10 anni e si è concentrata su pubblicazioni in lingua inglese.

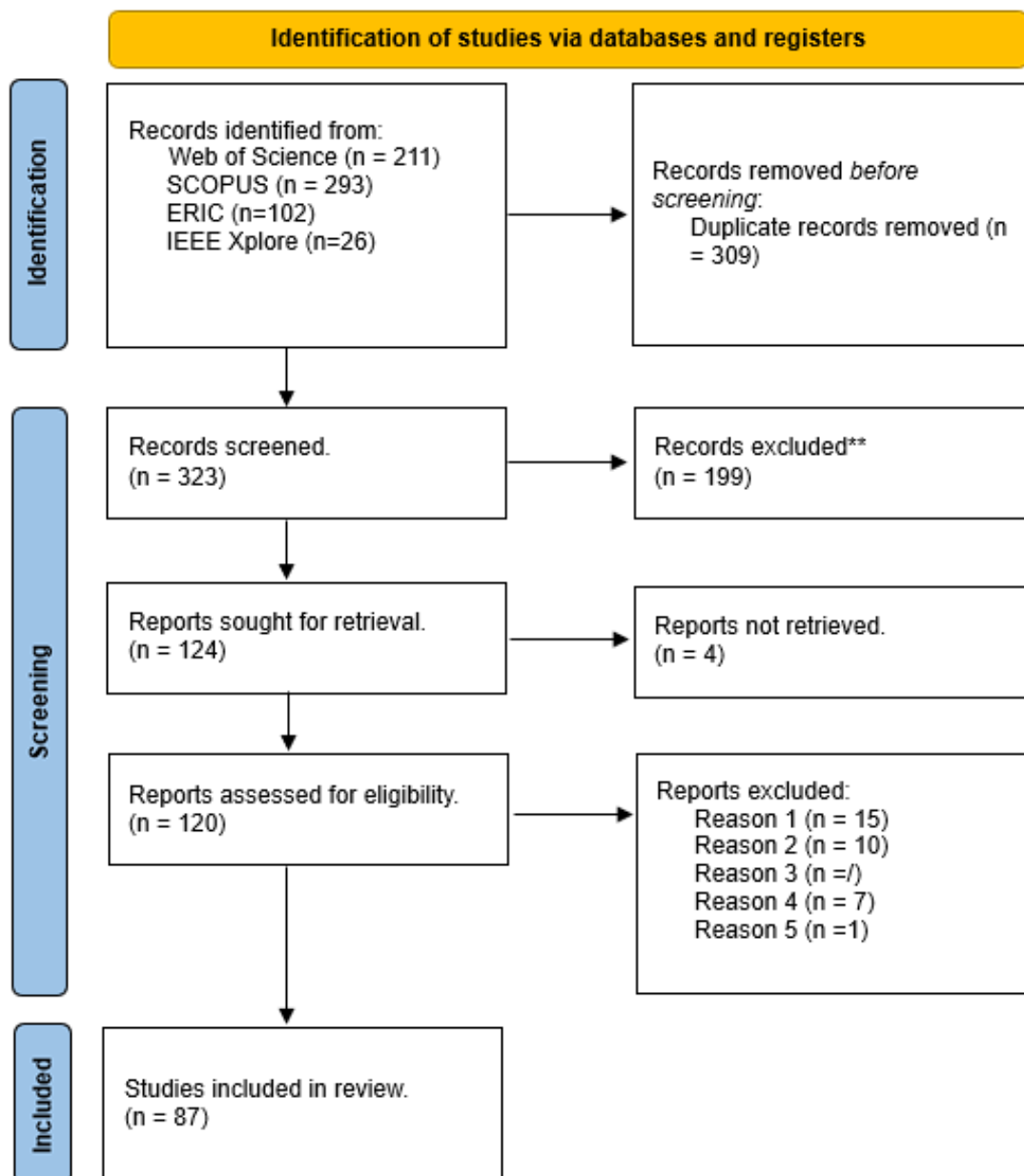
Per garantire la pertinenza e la qualità degli studi inclusi nella revisione, sono stati formulati e applicati cinque criteri di inclusione-esclusione durante le fasi di revisione degli abstract e dei testi completi. I criteri erano i seguenti:

1. Le opere che utilizzavano il termine AI literacy ma si concentravano su un argomento diverso sono state omesse.
2. Gli studi riguardanti altri tipi di alfabetizzazione, in particolare le competenze di lettura e scrittura, non sono stati considerati.
3. Editoriali e libri non sono stati inclusi a causa della loro natura non sottoposta a peer review.

4. Gli articoli che menzionavano "AI literacy" ma riguardavano l'applicazione dell'IA in campi specifici non correlati all'istruzione sono stati esclusi.
5. Le pubblicazioni in lingue diverse dall'inglese o dall'italiano sono state escluse per evitare errori di interpretazione dovuti a barriere linguistiche.

Il percorso di selezione degli studi è iniziato con l'identificazione di 632 record da diverse banche dati: 211 da Web of Science, 293 da SCOPUS, 102 da ERIC e 26 da IEEE Xplore. Da questo pool, i record sono stati preselezionati, portando alla rimozione di 309 record duplicati, ovvero circa il 48% del pool iniziale. Ciò ha lasciato 323 record da esaminare, circa il 51% del numero iniziale. Dopo lo screening, 199 record sono stati esclusi, il che rappresenta circa il 61% dei record. Questo ha portato a 124 rapporti (38,4%) che sono stati recuperati per l'analisi. Tuttavia, 4 rapporti non sono stati recuperati, rappresentando circa il 3,6% di quelli cercati per il recupero. Dei 120 rapporti recuperati e valutati per l'idoneità, un certo numero è stato escluso per vari motivi: 15 per il primo criterio (13,9% dei rapporti valutati), 10 per il secondo (9,3%), e così via, lasciando 87 studi che sono stati inclusi nella revisione in base a questa valutazione.

Due ricercatori hanno poi esaminato indipendentemente gli articoli selezionati per verificarne la pertinenza agli obiettivi dello studio attuale con il supporto dell'applicazione web Rayyan (Ouzzani et al., 2016; rayyan.ai), che è specificamente progettata per assistere nell'esecuzione di revisioni della letteratura. In particolare, sulla base dei criteri di inclusione predefiniti, i due ricercatori hanno analizzato gli abstract, e per gli studi più dubbi anche il testo completo, impiegando un metodo di selezione cieca. Rayyan è stato utilizzato per facilitare il processo di screening evidenziando le parole chiave principali dai criteri di inclusione ed esclusione. Le selezioni individuali sono state confrontate e il livello di accordo tra i valutatori, quantificato utilizzando il coefficiente kappa di Cohen, era $\kappa = 0,78$ (Cohen, 1960). Questo valore indica un livello sostanziale di accordo tra le decisioni dei ricercatori, come classificato da Landis & Koch (1977). Il numero finale degli articoli inclusi nella revisione sistematica è stato $N = 87$.



**All records were excluded manually. No automation tools were used at this stage.

Figura 1

1.2.3 Estrazione e codifica dei dati

Gli articoli identificati sono stati esaminati e discussi attraverso un approccio di ricerca QUAN/QUAL, affrontando tre domande di ricerca che esplorano lo sviluppo, la definizione e il contenuto tematico dell'AI literacy nelle pubblicazioni accademiche. Iniziando con la prima domanda di ricerca (RQ1), l'indagine ha tracciato la progressione cronologica e la dispersione geografica della letteratura sull'AI literacy e i suoi target demografici. Una revisione sistematica ha supportato la metodologia, durante la quale l'estrazione dei metadati ha fornito la necessaria impalcatura quantitativa. I metadati estratti, con l'aiuto di statistiche descrittive, sono culminati nella creazione della tabella 1, che espone le informazioni fondamentali degli articoli esaminati, facendo luce sulla traiettoria delle pubblicazioni sull'AI literacy nel tempo, sulla loro genesi globale e sui loro soggetti di analisi.

Per la seconda domanda di ricerca (RQ2), lo studio si è concentrato sulla delineazione di come l'AI literacy è definita all'interno della comunità accademica. Per dissezionare e comprendere le definizioni prevalenti di AI literacy tra gli studi raccolti, tutti gli articoli sono stati ispezionati e le definizioni sono state estratte. Questa procedura è culminata nella Tabella 3, che include citazioni dirette dagli articoli o una loro sintesi.

Per la terza domanda di ricerca (RQ3), è stato implementato un approccio a doppia strategia per l'analisi. Da un lato, è stata condotta un'analisi di rete bibliometrica per esplorare i principali cluster tematici e le loro interrelazioni, evidenziando la mappatura completa del campo di ricerca. D'altro lato, l'analisi è stata simultaneamente completata da un'analisi tematica, che ha comportato la codifica dei dati per identificare i temi prevalenti. Mentre la selezione degli articoli è stata eseguita da due ricercatori indipendenti, la validità delle analisi è stata garantita dalla triangolazione dei risultati sia dell'analisi bibliometrica che di quella tematica. I risultati complessivi sono stati incapsulati nella Tabella 2.

L'analisi di rete bibliometrica ha elucidato lo scheletro tematico del campo tramite un'analisi di co-occorrenza (Van Eck & Waltman, 2007) delle parole chiave più frequenti (100) all'interno degli 87 articoli identificati. La co-occorrenza delle parole più comuni assume che le parole chiave nello stesso contesto abbiano strette relazioni semantiche, permettendo la scoperta di termini chiave e la loro composizione in comunità. Nell'osservare reti complesse, si dice che una rete ha una struttura di comunità se i nodi della rete possono essere facilmente raggruppati in insiemi di nodi in modo tale che ciascun insieme di nodi sia densamente connesso internamente (Fortunato, 2010). Il corpus è stato costruito da un dataset creato automaticamente contenente i metadati e l'intero testo degli articoli selezionati, utilizzando uno script Python sviluppato dall'autore. Quindi, utilizzando tecniche di Natural Language Processing (NLP) (ad esempio, rimozione delle stop words, stemming), specificamente applicate dagli autori, il dataset è stato ripulito da parole non utili per l'analisi (pronomi, preposizioni, congiunzioni). Infine, l'analisi di rete bibliometrica è stata condotta su questo dataset con l'aiuto del software VOSviewer. VOSviewer è un programma sviluppato per costruire e visualizzare in dettaglio mappe bibliometriche. Può, ad esempio, essere utilizzato per costruire mappe di autori o riviste basate su dati di co-citazione o per costruire mappe di parole chiave basate su dati di co-occorrenza. In questo caso, questo approccio è stato utilizzato per identificare concetti e terminologie chiave degli articoli esaminati (Van Eck & Waltman, 2010). La rete di co-occorrenza contiene ancora informazioni ricche nella sua struttura topologica (come argomenti costituiti da un gruppo di nodi vicini e relazioni tra questi argomenti). Tali informazioni diventano più significative all'aumentare della scala dei dati. Per evitare cluster non informativi, la soglia di presenza è stata impostata a 60 menzioni. I dati sono stati poi visualizzati in VOSviewer per rilevare le comunità correlate (Fig. 2).

Questa analisi quantitativa è stata completata da un'analisi tematica induttiva attraverso la quale ogni articolo è stato sottoposto a meticolosa codifica e classificazione all'interno del software Atlas.ti (Atlas.ti, 2023). La doppia analisi ha prodotto una mappa concettuale delle interrelazioni tra le parole chiave e un quadro sintetizzato dei temi principali, presentato nella Tabella 4. Infine, è stata sviluppata una tabella di confronto dati per estrarre contenuti rilevanti dagli studi. I campi identificati hanno tentato di catturare: 1) l'identità della ricerca (Autori, Titolo, Anno), 2) il focus della ricerca (un "Breve Riassunto" è stato redatto per ogni articolo per incapsulare i suoi contributi e risultati principali, garantendo una comprensione dell'essenza del lavoro), 3) il target dell'articolo (per identificare il pubblico o gruppo soggetto principale a cui ogni studio mirava a rivolgersi, rivelando l'ambito e l'applicabilità della ricerca), 4) i temi principali affrontati (identificati attraverso la letteratura sono stati codificati per scoprire schemi ricorrenti, preoccupazioni e prospettive all'interno

del dominio dell'AI literacy). L'output di questo processo è stata una tabella per la registrazione dei dati creata al fine di catturare informazioni pertinenti dagli studi.

Questo approccio metodologico rigoroso e trasparente ha consentito di analizzare in modo sistematico la letteratura esistente sull'AI literacy, identificando tendenze, temi e definizioni chiave, e fornendo una solida base per lo sviluppo del framework concettuale presentato nella seconda parte di questo capitolo.

1.2.4 Risultati dell'analisi

L'analisi della distribuzione temporale e geografica degli 87 studi selezionati ha rivelato tendenze significative nello sviluppo della ricerca sull'AI literacy. Come illustrato nella Tabella 1, è evidente un notevole incremento delle pubblicazioni nell'arco di tempo considerato. L'anno 2023 ha registrato un aumento significativo, rappresentando il 43% delle pubblicazioni totali, mentre il 2022 ha contribuito con il 28%. L'anno 2021 rimane significativo, con il 23% del totale. Gli anni precedenti, come il 2020 e il 2019, hanno entrambi contribuito con il 2.5%, mentre il 2016 con il 1%.

Questa accelerazione nelle pubblicazioni, con un'impennata particolarmente evidente nel triennio 2021-2023, riflette il crescente interesse accademico e professionale per l'AI literacy, in parallelo con la rapida diffusione e integrazione dell'IA in diversi settori della società. È interessante notare come questo trend corrisponda all'emergere di applicazioni di IA sempre più sofisticate e accessibili, come i modelli di linguaggio di grandi dimensioni, che hanno reso l'IA più rilevante e visibile sia per gli esperti che per il grande pubblico. La crescita esponenziale degli studi sull'AI literacy suggerisce che questa area di ricerca è ancora in una fase di sviluppo dinamico e in espansione, con ampie prospettive di ulteriore crescita nei prossimi anni.

Per quanto riguarda la distribuzione geografica, la Cina si posiziona al primo posto con 27 pubblicazioni, che rappresentano il 31% del totale. Questo è seguito da vicino dagli Stati Uniti con 20 pubblicazioni (23%) e dalla Germania con 12 (14%). Altri paesi come Corea, Giappone, Taiwan, Belgio, Italia, Austria, Inghilterra, Spagna e Grecia contribuiscono anche al corpus di ricerca, con ogni paese che rappresenta dal 2% al 5% delle pubblicazioni totali. Inoltre, un gruppo eterogeneo di paesi tra cui Australia, Canada, Finlandia, Iran, Norvegia, Romania e Turchia contribuisce con un articolo ciascuno, rappresentando l'1% del totale.

La leadership della Cina nelle pubblicazioni sull'AI literacy può essere attribuita al forte sostegno governativo e a significativi investimenti nella ricerca e nell'educazione all'IA. Il piano "Next Generation AI Development Plan" della Cina (State Council of China, 2017) ha promosso un importante focus sull'educazione all'IA e sull'alfabetizzazione come priorità strategica, risultando in un crescente corpus di ricerca mirato a migliorare l'AI literacy in vari settori (Chen et al., 2020). Analogamente, gli Stati Uniti, sede di aziende tecnologiche e istituzioni di ricerca leader, hanno dato priorità all'educazione all'IA, in particolare nell'istruzione superiore e nello sviluppo professionale, per garantire che la forza lavoro sia preparata per un futuro guidato dall'IA (National AI Initiative Act, 2020).

Fattori culturali e istituzionali giocano anche un ruolo nel modellare la ricerca sull'AI literacy. In Cina, l'approccio centralizzato del governo all'educazione e alla politica permette l'implementazione su larga scala dei programmi di AI literacy, mentre negli Stati Uniti, l'enfasi è più sull'innovazione e sulle iniziative guidate dal settore privato. Queste differenze influenzano non solo il volume della ricerca ma anche le aree di focalizzazione. Ad esempio, le pubblicazioni cinesi spesso enfatizzano l'AI literacy nell'istruzione K-12 come parte di una più ampia strategia nazionale per coltivare l'expertise tecnologica fin dalla giovane età (Chen et al., 2020), mentre le pubblicazioni statunitensi

tendono a concentrarsi sull'istruzione superiore e sullo sviluppo di competenze professionali (Long & Magerko, 2020).

In aggiunta a Cina e Stati Uniti, paesi europei come Germania e Regno Unito contribuiscono significativamente alla ricerca sull'AI literacy. Tuttavia, la distribuzione è più frammentata, riflettendo la natura decentralizzata dei sistemi educativi europei e i diversi approcci all'educazione all'IA in diversi paesi. Mentre le iniziative europee, come quelle guidate dall'UNESCO (2022), sottolineano l'importanza delle pratiche etiche di IA e della privacy dei dati, la ricerca dagli Stati Uniti e dalla Cina spesso dà priorità alla competenza tecnologica e alla preparazione della forza lavoro (UNESCO, 2022).

Riguardo al pubblico target, gli articoli si rivolgono prevalentemente a "ricercatori" e "educatori", con rispettivamente 30 e 28 articoli. I "professionisti" sono il successivo target più comune con 15 articoli. Gli "studenti universitari", con 14 pubblicazioni, e gli "studenti K-12", con 12, sono anch'essi target significativi. Il "pubblico generale" e gli "studenti K-5" sono target meno frequentemente affrontati ma comunque notevoli. Infine, abbiamo "non esperti" e "genitori" con rispettivamente 6 e 4 articoli.

Tabella 1

Caratteristiche	Frequenza	Percentuale
Nazione		
China	27	31%
US	20	23%
Germany	12	14%
Korea	4	5%
Japan	2	2%
Taiwan	2	2%
Belgium	3	3%
Italy	2	2%
Austria	2	2%
England	2	2%
Spain	2	2%
Greece	2	2%
Australia, Canada, Finland, Iran, Norway, Romania, Türkiye	1	1%
Anno di pubblicazione		
2023	38	43%
2022	25	28%
2021	19	23%
2020	2	2.5%
2019	2	2.5%
2016	1	1%
Audience di riferimento		
Ricercatori	30	
Educatori	28	
Professionisti	15	
Studenti universitari	14	
Studenti K-12	12	
Cittadinanza	10	

Studenti K-5	8
Non esperti	6
Genitori	4
Studenti K-8	4

Questa distribuzione dei pubblici target riflette la natura multidimensionale dell'AI literacy e la sua rilevanza in diversi contesti. La predominanza di ricercatori e educatori come pubblico target sottolinea il ruolo cruciale che questi gruppi giocano nello sviluppo e nell'implementazione di programmi di AI literacy. Gli studiosi riconoscono che per integrare efficacemente l'AI literacy nei curricula educativi e nella formazione professionale, è essenziale prima equipaggiare coloro che faciliteranno questi processi di apprendimento con le conoscenze e le competenze necessarie.

La significativa attenzione rivolta agli studenti universitari e K-12 sottolinea inoltre l'importanza di sviluppare competenze di AI literacy fin dalle prime fasi dell'istruzione e di continuare questa formazione attraverso l'istruzione superiore. Ciò riflette una comprensione dell'AI literacy come un continuum di apprendimento che dovrebbe evolversi e approfondirsi man mano che gli studenti avanzano nei loro percorsi educativi.

Sebbene in minor numero, gli articoli rivolti a non esperti, genitori e al pubblico generale evidenziano un riconoscimento crescente che l'AI literacy non è rilevante solo in contesti educativi formali, ma è anche essenziale per individui in vari ruoli e fasi della vita. Ciò è particolarmente importante considerando la pervasività dell'IA nella vita quotidiana e la necessità che tutti i cittadini sviluppino una comprensione di base di queste tecnologie per navigare efficacemente in un mondo sempre più guidato dall'IA.

La presenza di articoli rivolti specificamente a professionisti indica un riconoscimento della rapida trasformazione del panorama lavorativo a causa dell'IA e della conseguente necessità per i professionisti di aggiornare continuamente le loro conoscenze e competenze per rimanere rilevanti in un'economia in evoluzione.

L'analisi di rete bibliometrica e tematica degli 87 articoli selezionati ha rivelato diversi cluster di ricerca distinti e interconnessi che forniscono una panoramica completa del paesaggio dell'AI literacy. Questi cluster emergono intorno al nodo centrale dell'"AI literacy", indicando aree tematiche distinte e la loro interconnessione all'interno del corpo della ricerca.

Nella visualizzazione di VOSviewer (Figura 2), un cluster prominente che appare vividamente è legato ai contesti educativi, con termini come "bambino", "programma di AI literacy", "curriculum" in grande evidenza. Adiacente a questo c'è un cluster che enfatizza l'applicazione dell'AI literacy nell'istruzione superiore, evidenziato da termini come "studente universitario", "corso" e "programma". Inoltre, c'è un cluster che coinvolge gli aspetti più ampi dell'AI literacy. Questo è visibile attraverso nodi come "fiducia" e "privacy". Inoltre, c'è un cluster che indica un interesse per il settore professionale e sanitario, sottolineato da nodi come "educazione per adulti" e "professionale". Un altro cluster è composto da elementi relativi a test e scale per l'AI literacy con termini come "questionario", "valutazione", "test di AI literacy". C'è anche un cluster di indagini su come l'AI literacy impatta sulla vita civica e sui diritti individuali all'interno di un contesto comunitario composto da termini come "famiglia", "comunità" o "implicazioni sociali". Ancora, un altro cluster significativo consiste in termini relativi a "spiegazione" e "affidabilità".

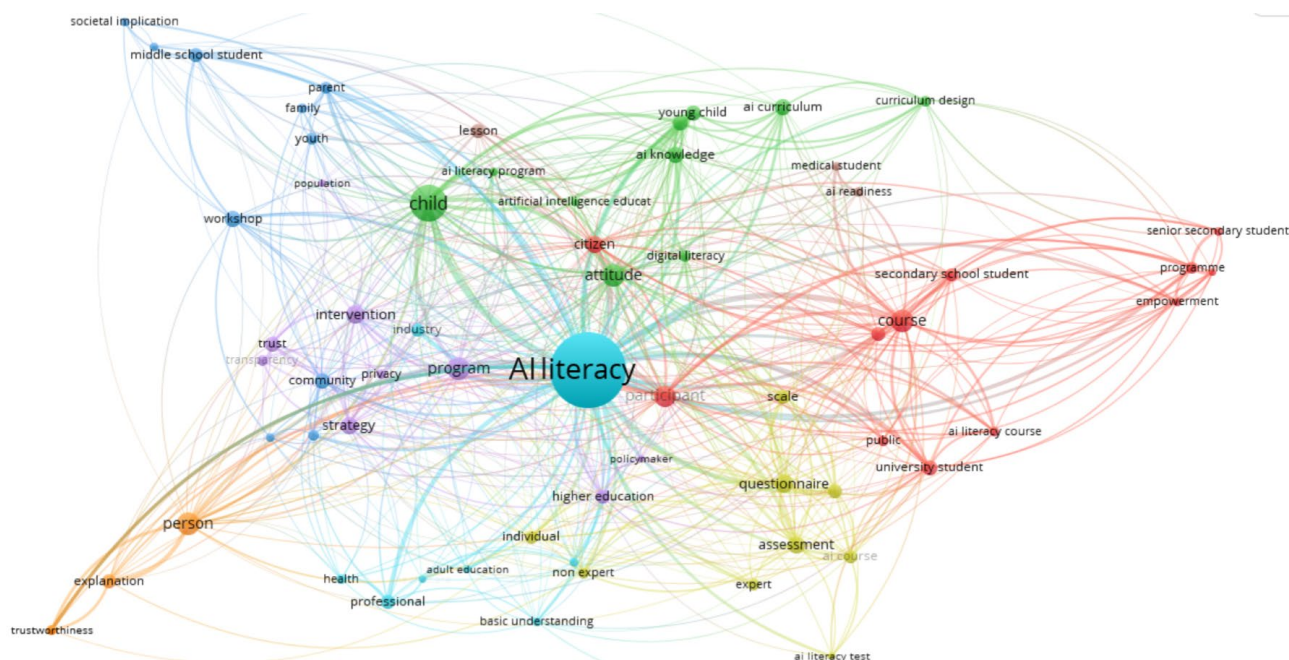


Figura 2

In aggiunta all'analisi bibliometrica, è stata condotta un'analisi tematica induttiva attraverso la codifica iterativa dei dati, che ha permesso l'identificazione e la categorizzazione dei temi mentre emergevano dalla letteratura. Ciò ha arricchito la comprensione degli aspetti sfumati dell'AI literacy. Dieci temi principali sono emersi da questa analisi, fornendo un quadro organizzato per comprendere i foci principali della ricerca sull'AI literacy:

Sviluppo di curriculum e interventi di AI literacy: Questo tema, affrontato da ricercatori come Kong et al. (2021), Laupichler et al. (2022), Park et al. (2023), Zhao et al. (2022) e altri, si concentra sulla creazione e implementazione di curricula e interventi per l'AI literacy. Questi studi si impegnano nel definire i confini dell'AI literacy stessa e nell'operationalizzarne i contenuti per l'insegnamento, offrendo informazioni preziose sulle componenti dell'alfabetizzazione stessa (ad esempio, quali argomenti dovrebbero essere insegnati nella dimensione critica).

Modelli e framework di AI literacy: In questo tema, studiosi come Ng et al. (2021), Carolus et al. (2023), Cuomo et al. (2022), Panciroli et al. (2023) e Pinski et al. (2023) propongono modelli e framework teorici per comprendere l'AI literacy. Questi modelli fungono da guide fondamentali per educatori e policymaker nella strutturazione dell'educazione all'IA. Specificamente, esaminano l'AI literacy in dimensioni distinte ma interrelate come quella cognitiva (che dettano la conoscenza necessaria per comprendere le tecnologie di IA, inclusa la logica algoritmica, l'analisi dei dati, e i principi del machine learning), affettiva (che indirizzano le considerazioni emotive ed etiche sollevate dall'IA, come le preoccupazioni sulla privacy, l'uso etico dell'IA, e la promozione di una mentalità responsabile verso gli impatti dell'IA sulla società), e socioculturale (che esplorano come l'AI literacy si interseca con contesti culturali e sociali, enfatizzando inclusività, accessibilità, e le implicazioni sociali dell'IA, incluse questioni di bias ed equità), offrendo un approccio strutturato per incorporare l'educazione all'IA attraverso diversi ambienti di apprendimento.

Implicazioni etiche e sociali dell'IA: Ricercatori come Hermann (2022), Yi (2021), Long et al. (2023), Lee (2021), Sakai (2023) e altri approfondiscono le complesse dimensioni etiche e sociali dell'IA. Il loro lavoro esamina come l'IA influenzi le norme sociali, i comportamenti individuali e le

considerazioni etiche. Questi studi non forniscono un intero framework per interpretare l'AI literacy, ma propongono invece linee guida per un uso responsabile dell'IA, discutendo questioni come la privacy, i bias e l'impatto sociale più ampio dell'adozione diffusa dell'IA. Esplorano anche come l'AI literacy possa includere una comprensione di questi aspetti etici e sociali.

AI literacy nell'istruzione K-12: Affrontato da ricercatori come Van Brummelen et al. (2021), Su et al. (2022), Steinbauer et al. (2021), Lee et al. (2021) e altri, questo tema enfatizza l'importanza di incorporare l'AI literacy nel sistema educativo K-12. Questi studi esplorano vari approcci per introdurre concetti di IA agli studenti più giovani, concentrandosi sullo sviluppo del curriculum, le metodologie didattiche e l'impatto dell'educazione all'IA sullo sviluppo cognitivo e sociale dei bambini. Gli autori sostengono un'esposizione precoce all'IA, preparando gli studenti per un futuro in cui l'IA gioca un ruolo significativo in molti aspetti della vita.

Strumenti di valutazione dell'AI literacy: Quest'area, esplorata da Laupichler et al. (2022), Wang et al. (2022) e Pinski et al. (2023), coinvolge lo sviluppo di strumenti e metodologie per valutare l'AI literacy. Questi strumenti mirano a valutare l'efficacia dei programmi di educazione all'IA e a misurare la comprensione dei concetti di IA da parte degli individui. La ricerca in questo campo contribuisce a stabilire benchmark e standard per l'AI literacy, fornendo a educatori e policymaker preziose informazioni sui punti di forza e le lacune nelle attuali iniziative educative sull'IA.

Integrazione dell'IA nell'istruzione e nei luoghi di lavoro: Questo tema, esplorato da autori come Cardon et al. (2023), Cetindamar et al. (2022), Hamburg et al. (2019) e Laupichler et al. (2022), si concentra su come l'intelligenza artificiale stia diventando parte integrante sia dei sistemi educativi che dei luoghi di lavoro. Questi studi esaminano come le tecnologie di IA stiano ridisegnando gli ambienti di apprendimento, le metodologie di insegnamento e le dinamiche del luogo di lavoro. Evidenziano la necessità di nuove competenze per integrare efficacemente l'IA in questi contesti, discutendo strategie per educatori e datori di lavoro per adattarsi a questa evoluzione tecnologica.

Ambienti di apprendimento familiare e informale: Investigato da Druga et al. (2022) e Long et al. (2023), questo tema esplora l'AI literacy in ambienti di apprendimento non tradizionali come i contesti familiari. Questi studi evidenziano il ruolo dei genitori e delle attività educative informali nel favorire l'AI literacy. Discutono di come le famiglie possano interagire con le tecnologie di IA, l'importanza della guida genitoriale nell'educazione all'IA dei bambini e il potenziale delle esperienze di apprendimento informale nello sviluppare una comprensione fondamentale dell'IA.

AI literacy per l'educazione nella prima infanzia: Ricercatori come Su et al. (2023), Yang (2022) e Zhao et al. (2022) si concentrano sull'AI literacy nell'educazione della prima infanzia. Il loro lavoro enfatizza l'importanza di introdurre concetti di IA ai giovani apprendenti in modo coinvolgente appropriato. Questo include l'esplorazione di strategie pedagogiche per insegnare concetti di IA di base, l'uso di strumenti educativi potenziati dall'IA e l'impatto dell'educazione precoce all'IA sull'apprendimento e lo sviluppo dei bambini.

IA nella sanità e nell'educazione medica: Laupichler et al. (2022), Perchik et al. (2023), Kang et al. (2023), Sakai (2023) e Mello-Thoms (2023) investigano l'applicazione e le implicazioni dell'IA nella sanità e nell'educazione medica. Il loro lavoro esplora come le tecnologie di IA stiano trasformando la diagnostica medica, la pianificazione del trattamento e la cura del paziente. Questi studi discutono anche l'integrazione dell'educazione all'IA nella formazione medica, enfatizzando la necessità per i

professionisti sanitari di essere competenti nell'uso e nella comprensione degli strumenti di IA in contesti clinici.

IA e digital literacy: Carolus et al. (2023) e Lee (2021) si concentrano sulla literacy richiesta per interazioni digitali efficaci e responsabili con i sistemi di IA. La loro ricerca affronta le competenze e la conoscenza necessarie per navigare e utilizzare tecnologie e piattaforme guidate dall'IA. Questo include la comprensione di come funzionano gli algoritmi di IA, il riconoscimento dei limiti e dei potenziali bias nei sistemi di IA e lo sviluppo di competenze di pensiero critico nel contesto delle interazioni digitali con l'IA.

La distribuzione di frequenza di questi dieci temi, come indicato nella Tabella 2, mostra una predominanza di ricerca su "Sviluppo di Curriculum e Interventi di AI literacy" (39 articoli) e "Modelli e Framework di AI literacy" (38 articoli), seguito da "Implicazioni Etiche e Sociali dell'IA" (29 articoli) e "AI literacy nell'Istruzione K-12" (26 articoli). I temi meno rappresentati includono "Strumenti di Valutazione dell'AI literacy" (14 articoli), "Integrazione dell'IA nell'Istruzione e nei Luoghi di Lavoro" (10 articoli), "Ambienti di Apprendimento Familiare e Informale" (8 articoli), "AI literacy per l'Educazione nella Prima Infanzia" (7 articoli), e sia "IA nella Sanità e nell'Educazione Medica" che "IA e Digital Literacy" (6 articoli ciascuno).

Questa distribuzione tematica riflette lo stato attuale della ricerca sull'AI literacy, con un forte focus sullo sviluppo di fondamenti teorici e applicazioni pratiche nel contesto educativo. La presenza significativa di ricerca sulle implicazioni etiche e sociali dell'IA sottolinea la crescente consapevolezza dell'importanza di questi aspetti nell'educazione all'IA. La minore rappresentazione di temi come l'AI literacy negli ambienti sanitari o familiari suggerisce aree emergenti che potrebbero beneficiare di ulteriore esplorazione in futuro.

Tabella 2

Tema	Autori	Frequenza
Sviluppo di interventi e curriculum dedicati all'AI literacy	Kong et al. (2021); Laupichler et al.; (2022); Panciroli et al. (2023); Park et al., (2023); Perchik et al. (2023); Southworth et al. (2023); Zhao et al., (2022)	39
Modelli e Framework sull'AI literacy	Carolus et al. (2023); Cuomo et al. (2022); Kong et al., (2021); Long & Magerko; Ng et al., (2021); Panciroli et al. (2023); Pinski et al. (2023)	38
Implicazioni etiche e sociali dell'AI	Hermann (2022); Lee, (2021).; Long et al (2023).; Lubin et al. (2021); Ng et al., (2021); Sakai (2023); Shih et al. (2021); Steinbauer et al. (2021); Yi (2021); Zhang et al. (2022)	29
AI literacy nell'educazione K-12	Casal-Otero et al. (2023); Eguchi (2021); Lee et al. (2021); Mertala et al. (2022); Ng et al., (2021); Steinbauer et al. (2021); Su et al. (2022); Van Brummelen et al., (2021)	26
Strumenti di valutazione per l'AI literacy	Laupichler et al., (2022); Pinski et al. (2023); Wang et al., (2022)	14
Integrazione di strumenti di AI nel luogo di lavoro	Cardon et al. (2023); Cetindamar et al. (2022); Hamburg et al. (2019); Laupichler et al., (2022)	10
Educazione informale all'AI	Druga et al. (2022); Long et al (2023).; Hemment et al. (2023); Kasinidou (2023)	8
AI literacy per l'early childhood Education	Su et al. (2022); Su & Yang (2023); Yang (2022); Zhao et al., (2022)	7

Tema	Autori	Frequenza
AI nella professione medica	Kang et al. (2023); Laupichler et al., (2022); Mello-Thoms (2023); Perchik et al. (2023); Sakai (2023)	6
AI e digital literacy	Carolus et al. (2023); Lee (2021); Liu & Xie (2021); Julie et al. (2020); Kang et al., (2023)	6

1.2.5 Definizione e concettualizzazione dell'AI literacy

Una delle domande di ricerca fondamentali della presente revisione sistematica riguarda come la letteratura concettualizza e differenzia le definizioni di AI literacy. L'analisi delle definizioni presenti negli articoli selezionati rivela una gamma diversificata di interpretazioni, sottolineando la complessità e la natura multiforme dell'AI literacy.

Dall'esame della letteratura, emerge che il termine "AI literacy" fu proposto per la prima volta nel 2016 da Kandlhofer et al., definendola come "la capacità di comprendere le tecniche e i concetti di base dietro l'intelligenza artificiale in diversi prodotti e servizi". Questa definizione iniziale poneva l'accento sulla comprensione concettuale e tecnica dell'IA, ma non includeva esplicitamente aspetti relativi all'uso critico o alle implicazioni etiche.

Con l'evoluzione del campo, le definizioni successive hanno iniziato ad ampliare il concetto di AI literacy per includere dimensioni più ampie. Ad esempio, Aoun (2017) ha definito l'AI literacy come "la capacità di realizzare e utilizzare l'intelligenza artificiale, comprendendone i concetti e gli usi", aggiungendo una componente di applicazione pratica alla definizione.

Un significativo arricchimento concettuale si osserva nella definizione proposta da Druga et al. (2019), che hanno introdotto la dimensione etica: "L'AI literacy può essere definita come la conoscenza e la comprensione delle funzioni di base dell'IA e come utilizzare le applicazioni dell'IA nella vita quotidiana in modo eticamente adeguato". Questa definizione segna un punto di svolta nell'evoluzione del concetto, riconoscendo che l'AI literacy va oltre la mera comprensione tecnica per includere considerazioni etiche sull'uso dell'IA.

Long & Magerko (2020) hanno ulteriormente ampliato il concetto, definendo l'AI literacy come un "insieme di competenze che consente agli individui di valutare criticamente le tecnologie di IA, di comunicare e collaborare efficacemente con l'IA, e di utilizzare l'IA in modo appropriato online, a casa e sul posto di lavoro". Questa definizione introduce esplicitamente le dimensioni della valutazione critica, della comunicazione e della collaborazione, sottolineando il ruolo attivo dell'individuo nell'interazione con l'IA.

Kim et al. (2021) hanno proposto un approccio tripartito all'AI literacy, comprendente conoscenza, abilità e atteggiamento. La conoscenza è definita come la capacità di comprendere i concetti fondamentali dell'IA (ad esempio, 'definizioni e tipi di IA', 'dati e machine learning', 'applicazioni'); l'abilità (o skill) comprende il pensiero computazionale o le competenze di programmazione; mentre l'atteggiamento (o attitude) si riferisce alla capacità di considerare collettivamente tutti gli aspetti dell'IA nella società, come l'"impatto sociale" e la "collaborazione con l'IA".

Kong & Zhang (2021) hanno discusso l'AI literacy nel contesto della preparazione alla carriera, enfatizzando la necessità di comprendere l'impatto dell'IA sui futuri mercati del lavoro e le competenze necessarie per prosperare in un mondo guidato dall'IA. La loro definizione include tre componenti: concetti di IA, implicazioni etiche e sociali, e futuro della carriera nell'IA.

Ng et al. (2021) hanno presentato un approccio strutturato, organizzando l'AI literacy in quattro concetti distinti: conoscere e comprendere l'IA (funzioni di base dell'IA e come utilizzare le applicazioni di IA), utilizzare e applicare l'IA (applicare la conoscenza, i concetti e le applicazioni dell'IA in diversi scenari), valutare e creare IA (valutare, stimare, prevedere, progettare), e etica dell'IA (equità, responsabilità, trasparenza, sicurezza).

Yi (2021) ha evidenziato gli aspetti culturali dell'AI literacy (comprendere le caratteristiche dell'IA e la società a cui si applica) e gli aspetti soggettivi (metacognizione come "competenza" per l'AI literacy). Questa prospettiva è unica nel collegare l'AI literacy all'adattabilità personale e culturale, suggerendo che l'AI literacy non sia solo un insieme di competenze tecniche, ma comporti anche la capacità di navigare e modellare la propria vita in mezzo alle trasformazioni portate dall'IA.

Deuze & Beckett (2022) e Hermann (2022) hanno aggiunto rispettivamente una lente normativa e pratica all'AI literacy, concentrandosi sulla sua applicazione e sul suo impatto. Laupichler et al. (2022) e Wang et al. (2022) hanno sottolineato l'importanza del processo decisionale informato riguardo all'uso e alle implicazioni dell'IA.

Infine, UNESCO (2022) e Weber et al. (2023) hanno contribuito con una visione completa, comprendente la data literacy e l'algorithm literacy come componenti chiave. Parallelamente all'AI literacy, la data literacy è emersa come componente critica, consentendo agli individui di interpretare, analizzare e prendere decisioni informate basate su dati guidati dall'IA. Riconoscendo ciò, l'UNESCO (2022) sottolinea il ruolo della data literacy come sottoinsieme essenziale dell'AI literacy, in particolare in relazione alla comprensione dei meccanismi del processo decisionale guidato dall'IA. Inoltre, il framework Learning Compass dell'OCSE posiziona la data literacy come una delle competenze fondamentali per gli studenti del XXI secolo, in particolare in un mondo in cui l'IA sta ridisegnando settori e paradigmi educativi.

Un'analisi comparativa di queste definizioni rivela due principali prospettive: (1) un modello basato sulla competenza tecnica, che si concentra su coding, pensiero computazionale e meccanismi dell'IA, e (2) un framework socio-etico, che dà priorità alle implicazioni dell'IA sugli individui e sulla società. Dall'esame di queste diverse definizioni, il nostro studio fornisce una comprensione strutturata di come l'AI literacy sia stata concettualizzata nel tempo, facendo luce su lacune e tendenze emergenti nella letteratura.

Costruendo su questo confronto, diventa evidente che l'AI literacy non riguarda solo la comprensione di come funzionano i sistemi di IA, ma anche la comprensione delle loro implicazioni etiche, sociali e professionali. L'interazione tra queste prospettive evidenzia la necessità di un approccio equilibrato che incorpori sia dimensioni tecniche che etiche, garantendo che gli individui siano equipaggiati per impegnarsi criticamente con l'IA in molteplici contesti.

La dimensione etica dell'AI literacy, come proposta da autori come Druga et al. (2019 & 2022) e Wong et al. (2020), sottolinea la crescente consapevolezza dell'impatto sociale più ampio dell'IA. Questo aspetto è fondamentale per guidare lo sviluppo e l'uso responsabile delle tecnologie di IA, affrontando preoccupazioni come la privacy dei dati, i bias algoritmici e l'implementazione etica dell'IA. L'inclusione di considerazioni etiche riflette un riconoscimento degli effetti potenzialmente di vasta portata dell'IA sulla società, rendendo necessario un approccio ben arrotondato e responsabile all'AI literacy.

Tabella 3

Autore	Definizione di AI literacy	Citato da
--------	----------------------------	-----------

Kandlhofer et al. (2016)	The ability to understand the basic techniques and concepts behind AI in different products and services	Cetindamar et al, (2022); Cuomo et al., (2022) Ng et al., (2021); Laupichler et al., (2022); Su et al., (2022); Yang (2022), Weber et al. (2023)
Aoun (2017)	The ability to realize and utilize AI by understanding its concepts and usage.	Yi (2021), Cuomo et al. (2022)
Druga et al (2019) Druga et al (2022)	AI literacy can be defined as knowing and understanding the basic functions of AI and how to use AI applications in everyday life ethically. AI literacies include the ability to read, work with, analyse and author with AI	Yang (2022)
Long & Magerko (2020)	Set of competencies that enables individuals to critically evaluate AI technologies; communicate and collaborate effectively with AI; and use AI as a tool online, at home, and in the workplace	Carolus (2023), Cetindamar et al., (2022); Cuomo et al (2022), Faruque et al., (2022) Kasinidou (2023) Laupichler et al., (2022); Park et al., (2023); Pinski & Benlian (2023), Su et al., (2022); Van Brummelen et al., (2021); Velandar et al, (2023); Yi (2021).
Wong et al 2020	Define AI literacy to have the components AI Concepts (Understand the basic AI), AI applications (Appreciate the real-world applications), AI Ethics and Safety (Describe the ethical challenges and safety issues)	Kong et al., 2021
Liu&Xie (2021)	Define AI literacy as composed by digital literacy, computational thinking, and programming ability.	-
UNICEF (2021)	AI literacy includes knowledge of the concepts, skills, and ethical considerations surrounding the creation and use of AI. It describes the ability to engage with AI as a skilled user but can also encompass knowledge of how to develop AI, understanding data biases, AI ethics, user rights	Su & Zhong (2022)
Kim et al., 2021	Summarized three competences to achieve AI literacy: AI Knowledge, AI Skill, and AI Attitude.	Su & Zhong (2022), Yau et al. (2022)
Kong & Zhang (2021)	Describe AI literacy with three components. AI concepts includes factual knowledge about AI and its concepts and technical details. Ethical and societal implications include the ability to understand consequences of the use of AI for society, and AI Career Future concerns the impact of AI on future careers.	Carolus (2023)
Kong et al. (2021)	AI literacy as understanding of AI concepts and competencies in using AI concepts for evaluation and using AI concepts for understanding the real world	Kong et al., (2022); Yi (2021)
Ng et al. (2021)	AI literacy can be organized into four concepts: (1) know and understand AI, (2) use and apply AI, (3) evaluate and create AI, and (4) AI Ethics	Carolus (2023), Kong et al, Pinski & Benlian (2023), Su et al., (2022); Yang (2022), Zhao et al., (2022)
Yi (2021)	AI literacy is an individual's ability to not only utilize AI, but to also critically recognize changing cultures. Furthermore, based on understanding AI, AI literacy allows the individual to design their own life. In other words, AI literacy is the basic ability to become a subjective human in the AI era	Kong et al., (2022); Velandar et al., (2023)
Cetindamar et al (2022)	A bundle of four core capabilities", namely technology-related, work-related, human-machine-related, and learning-related capabilities	Carolus (2023), Kong et al., (2022) Weber et al. (2023)
Deuze & Beckett (2022)	Artificial intelligence literacy is not simply knowing about AI, but also understanding and appreciating its normative dimension, as much as it is linked to impact and action: being able to identify ways to apply AI responsibly, creatively and efficiently. The three key components of AI literacy are: knowledge about artificial intelligence, the ability to recognize instances, skills to help, coach or teach others when strategically understanding, imagining, developing and implementing AI.	-

Hermann (2022)	Individuals' basic understanding of (a) how and which data are gathered; (b) the way data are combined or compared to draw inferences, create, and disseminate content; c) the own capacity to decide, act, and object; (d) AI's susceptibility to biases and selectivity; and (e) AI's potential impact in the aggregate.	Weber et al. (2023)
Laupichler et al (2022)	refers to the knowledge and understanding of AI that is necessary for individuals to participate in the broader discourse around AI and make informed decisions about its use and implications	Southworth et al. (2023)
Wang et al (2022)	AI literacy refers to the ability to properly identify, use, and evaluate AI-related products under the premise of ethical standards.	Carolus (2023), Weber et al. (2023)
UNESCO 2022	Competency about AI, including knowledge, understanding, skills, and value. AI literacy comprises both data literacy, or the ability to understand how AI collects, cleans, manipulates, and analyses data; and algorithm literacy, or the ability to understand how AI algorithms find patterns and connections in the data, which might be used for human-machine interactions.	Kong et al. 2022
Weber et al (2023)	Define AI literacy as a set of socio-technical competencies of humans that shape relevant types of human-AI interaction.	-

Dall'analisi delle definizioni, è possibile identificare quattro dimensioni fondamentali dell'AI literacy, che forniscono un framework concettuale per comprendere e sviluppare programmi educativi in questo ambito:

- **Dimensione conoscitiva:** Questa dimensione si riferisce alla comprensione dei concetti di base dell'IA, inclusi i suoi principi fondamentali, le diverse tipologie e le sue applicazioni. Include la conoscenza di cosa sia l'IA, come funzioni e quali siano le sue potenzialità e limitazioni.
- **Dimensione operativa:** Questa dimensione riguarda la capacità di utilizzare e applicare l'IA in diversi contesti. Include competenze pratiche come la programmazione, l'uso di strumenti di IA e la capacità di implementare soluzioni basate sull'IA per risolvere problemi.
- **Dimensione critica:** Questa dimensione implica la capacità di valutare criticamente i sistemi di IA, comprendendone i limiti, i bias potenziali e le implicazioni sociali. Include la capacità di analizzare e interpretare i risultati prodotti dall'IA e di comprendere i processi decisionali alla base.
- **Dimensione etica:** Questa dimensione si concentra sulla comprensione delle implicazioni etiche dell'IA e sulla capacità di prendere decisioni informate e responsabili riguardo al suo uso. Include la consapevolezza di questioni come la privacy, l'equità, la trasparenza e l'impatto sociale dell'IA.

Queste quattro dimensioni, emerse dall'analisi della letteratura, costituiscono la base per il framework concettuale che verrà presentato nella prossima sezione. L'integrazione di queste dimensioni in un approccio coerente all'AI literacy rappresenta un contributo significativo alla comprensione di come l'educazione all'IA debba essere strutturata per preparare gli individui a navigare efficacemente in un mondo sempre più plasmato dall'intelligenza artificiale.

1.2.6 Discussione e implicazioni

I risultati della nostra revisione sistematica della letteratura sull'AI literacy mettono in luce lo stato attuale della ricerca in questo campo emergente e le sue implicazioni per l'educazione, la pratica e la politica. In questa sezione, discutiamo i principali risultati e le loro implicazioni, evidenziando anche le lacune nella letteratura esistente e le direzioni per la ricerca futura.

I risultati relativi alla prima domanda di ricerca mostrano una dominanza di pubblicazioni provenienti dalla Cina e dagli Stati Uniti, attribuibile al forte sostegno governativo e a significativi investimenti nella ricerca e nell'educazione all'IA in entrambi i paesi. La spinta della Cina verso il dominio dell'IA, come delineato nel suo "Piano di sviluppo dell'IA di prossima generazione" (State Council of China, 2017), ha promosso un importante focus sull'educazione e l'alfabetizzazione all'IA come priorità strategica, risultando in un crescente corpus di ricerca mirato a migliorare l'AI literacy in vari settori (Chen et al., 2020). Analogamente, gli Stati Uniti, sede di aziende tecnologiche e istituzioni di ricerca leader, hanno dato priorità all'educazione all'IA, in particolare nell'istruzione superiore e nello sviluppo professionale, per garantire che la forza lavoro sia preparata per un futuro guidato dall'IA (National AI Initiative Act, 2020).

Fattori culturali e istituzionali giocano anche un ruolo nel modellare la ricerca sull'AI literacy. In Cina, l'approccio centralizzato del governo all'educazione e alla politica permette l'implementazione su larga scala dei programmi di AI literacy, mentre negli Stati Uniti, l'enfasi è più sull'innovazione e sulle iniziative guidate dal settore privato. Queste differenze influenzano non solo il volume della ricerca ma anche le aree di focalizzazione. Ad esempio, le pubblicazioni cinesi spesso enfatizzano l'AI literacy nell'istruzione K-12 come parte di una più ampia strategia nazionale per coltivare l'expertise tecnologica fin dalla giovane età (Chen et al., 2020), mentre le pubblicazioni statunitensi tendono a concentrarsi sull'istruzione superiore e sullo sviluppo di competenze professionali (Long & Magerko, 2020).

In aggiunta a Cina e Stati Uniti, paesi europei come Germania e Regno Unito contribuiscono significativamente alla ricerca sull'AI literacy. Tuttavia, la distribuzione è più frammentata, riflettendo la natura decentralizzata dei sistemi educativi europei e i diversi approcci all'educazione all'IA in diversi paesi. Mentre le iniziative europee, come quelle guidate dall'UNESCO (2022), sottolineano l'importanza delle pratiche etiche di IA e della privacy dei dati, la ricerca dagli Stati Uniti e dalla Cina spesso dà priorità alla competenza tecnologica e alla preparazione della forza lavoro (UNESCO, 2022).

Questa polarizzazione geografica nella ricerca sull'AI literacy ha implicazioni significative. Da un lato, riflette l'importanza strategica che i governi attribuiscono all'IA e alla preparazione dei cittadini per un futuro guidato dall'IA. Dall'altro, suggerisce la necessità di una collaborazione internazionale più forte per garantire che l'AI literacy venga sviluppata in modo inclusivo e equo a livello globale, piuttosto che concentrarsi solo in alcune regioni.

La seconda domanda di ricerca ha esplorato come la letteratura definisce l'AI literacy. I risultati rivelano una gamma diversificata di interpretazioni, sottolineando la complessità e la natura multiforme dell'AI literacy. Mentre alcuni studiosi (Kandlhofer et al., 2016; Aoun, 2017) enfatizzano la conoscenza fondamentale dell'IA e le sue applicazioni, altri (Druga et al., 2019; Ng et al., 2021) integrano considerazioni etiche, impatto sociale e la necessità di un impegno responsabile con l'IA.

Un'analisi comparativa evidenzia due prospettive primarie: (1) un modello basato sulla competenza tecnica, che si concentra su coding, pensiero computazionale e meccanismi dell'IA, e (2) un framework socio-etico, che dà priorità alle implicazioni dell'IA sugli individui e sulla società.

Esaminando queste diverse definizioni, il nostro studio fornisce una comprensione strutturata di come l'AI literacy sia stata concettualizzata nel tempo, facendo luce su lacune e tendenze emergenti nella letteratura.

Questa diversità di definizioni riflette la natura in evoluzione dell'IA stessa e gli approcci multidisciplinari necessari per comprenderla pienamente. Suggerisce anche la necessità di un framework integrato che riconosca sia gli aspetti tecnici che quelli socio-etici dell'AI literacy, come proposto nella seconda parte di questo capitolo.

Un'implicazione chiave di queste diverse concettualizzazioni è la necessità di approcci educativi flessibili che possano adattarsi a diversi contesti e obiettivi. Contesti educativi diversi – dall'istruzione primaria all'educazione degli adulti, dalla formazione tecnica all'istruzione liberale – possono richiedere enfasi diverse sulle varie dimensioni dell'AI literacy.

I risultati della nostra revisione sistematica evidenziano la natura diversificata e multiforme dell'AI literacy, rivelando la complessità di questo concetto attraverso diversi contesti educativi e sociali. L'analisi sottolinea l'importanza di sviluppare framework e curricula robusti adattati a diversi livelli educativi e bisogni sociali, riflettendo i vari approcci all'AI literacy emersi negli ultimi anni. Centrale a questo sforzo è la delineazione di modelli di AI literacy, che formano la spina dorsale di come l'AI literacy è compresa, insegnata e valutata. Questi framework possono essere ampiamente categorizzati in due orientamenti principali: framework tecnici, che si allineano strettamente con l'educazione STEM, e framework olistici, che sono più adatti ad approcci interdisciplinari.

I framework tecnici, come quelli proposti da Ng et al. (2021) e Su & Zhong (2022), enfatizzano il pensiero computazionale e le competenze tecniche (ad esempio, reti neurali o deep learning). Questi modelli si concentrano sulle dimensioni cognitive e operative dell'AI literacy, fornendo la base per comprendere come le tecnologie di IA funzionano a livello tecnico. Tali framework sono particolarmente efficaci in ambienti di apprendimento strutturati dove la competenza tecnica è di primaria importanza, come nelle materie STEM. Qui, gli studenti sviluppano le competenze necessarie per interagire con i sistemi di IA, inclusa la conoscenza di algoritmi, machine learning e analisi dei dati.

In contrasto, i framework olistici—come discusso da Yi (2021), Cuomo et al. (2022) e Hermann (2022) —offrono una prospettiva più ampia, integrando considerazioni etiche, sociali e culturali nell'AI literacy. Questi modelli si allineano più strettamente con curricula interdisciplinari, dove l'AI literacy non riguarda solo competenze tecniche ma anche la comprensione degli impatti sociali delle tecnologie di IA. Questo approccio si adatta bene anche all'interno di materie come gli studi sociali o le discipline umanistiche, dove gli studenti sono incoraggiati a riflettere sulle implicazioni etiche dell'IA, incluse questioni come privacy, bias e trasparenza. Promuovendo il pensiero critico, i framework olistici preparano gli studenti a impegnarsi con l'IA in modi che considerano i suoi effetti più ampi sulla società.

Passando dai framework allo sviluppo del curriculum, la sfida chiave è come operationalizzare questi diversi approcci all'AI literacy. Come evidenziato da Kong et al. (2021) e Laupichler et al. (2022), gli sforzi di sviluppo del curriculum devono essere abbastanza flessibili da accogliere sia framework tecnici che olistici, a seconda del contesto di apprendimento. Alcuni curricula, in particolare quelli allineati con le linee guida dell'UNESCO (2022), enfatizzano la necessità che l'AI literacy includa sia considerazioni etiche che competenze tecniche. Ad esempio, l'UNESCO sottolinea l'importanza della digital literacy come parte dell'AI literacy, garantendo che gli apprendenti non siano solo capaci di

utilizzare le tecnologie di IA ma anche criticamente consapevoli delle questioni etiche che queste tecnologie sollevano.

Una delle domande principali nello sviluppo del curriculum è come bilanciare queste diverse componenti. Ad esempio, i framework tecnici sono ben adatti a curricula orientati allo STEM, dove il focus è sullo sviluppo di competenze tecniche come la codifica, la manipolazione dei dati e il machine learning (Ng et al., 2021; Su & Zhong, 2022). Qui, il curriculum potrebbe includere progetti pratici che permettono agli studenti di progettare e implementare semplici sistemi di IA, favorendo una comprensione profonda dei meccanismi dietro le tecnologie di IA. In contrasto, curricula più interdisciplinari potrebbero integrare discussioni sull'etica dell'IA nei corsi di studi sociali, dove gli studenti esplorano come le tecnologie di IA influenzano privacy, governance ed equità (Long & Magerko, 2020; Cuomo et al., 2022). Questo approccio è in linea con l'enfasi dell'UNESCO sull'uso responsabile dell'IA e la necessità di sviluppare una comprensione ben arrotondata degli impatti sociali dell'IA (UNESCO, 2022).

L'integrazione dell'AI literacy nell'istruzione K-12 è un aspetto critico dello sviluppo del curriculum. Come notato da Van Brummelen et al. (2021) e Su et al. (2022), l'esposizione precoce ai concetti di IA può aiutare gli studenti a sviluppare le competenze fondamentali di cui avranno bisogno in un futuro sempre più plasmato dall'IA. Tuttavia, l'approccio all'AI literacy nell'istruzione K-12 varia significativamente a seconda che l'enfasi sia sulla competenza tecnica o su questioni etiche e sociali più ampie.

Ad esempio, in Finlandia, l'AI literacy è incorporata nei programmi STEM attraverso l'apprendimento basato su progetti, dove gli studenti sono incoraggiati ad applicare concetti di IA per risolvere problemi del mondo reale (Mertala et al., 2022). Questo approccio pratico si allinea bene con i framework tecnici, aiutando gli studenti a sviluppare competenze computazionali mentre favorisce creatività e problem-solving. In contrasto, la Corea del Sud ha integrato l'AI literacy nel suo curriculum nazionale K-12 con un forte focus sull'etica (Kim et al., 2021). Qui, gli studenti si impegnano in discussioni sulle implicazioni sociali dell'IA, come l'uso dell'IA nella sorveglianza o il suo ruolo negli algoritmi dei social media, incoraggiandoli a riflettere criticamente sull'influenza dell'IA nella vita quotidiana.

Nonostante i progressi nell'integrazione dell'AI literacy nell'istruzione, una delle principali lacune rimane la mancanza di una guida pratica su come insegnare l'etica dell'IA. Mentre c'è un ampio riconoscimento dell'importanza delle considerazioni etiche, come discusso da Druga et al. (2019 & 2022), molti curricula si concentrano ancora principalmente sulle competenze tecniche, con meno enfasi sulle strategie operative per integrare l'etica in classe. Questo squilibrio evidenzia la necessità di ulteriore ricerca e sviluppo di strumenti pedagogici che possano aiutare gli insegnanti ad affrontare questioni etiche insieme alla formazione tecnica.

In aggiunta all'istruzione formale, l'AI literacy è anche coltivata in ambienti di apprendimento informali, come i contesti familiari. La ricerca di Druga et al. (2022) e Long et al. (2023) evidenzia l'importanza dei genitori e dei caregiver nel favorire una comprensione di base delle tecnologie di IA al di fuori dell'aula. Queste esperienze di apprendimento informali sono cruciali per sviluppare l'AI literacy precoce, in particolare nei bambini più piccoli, e offrono un approccio complementare all'istruzione formale.

Infine, la crescente integrazione delle tecnologie di IA nel luogo di lavoro presenta una sfida significativa per l'AI literacy. Come notato da Cardon et al. (2023) e Cetindamar et al. (2022), l'AI literacy sta diventando una competenza essenziale per navigare negli ambienti di lavoro moderni,

dove l'IA è utilizzata per l'analisi dei dati, il processo decisionale e l'automazione. I lavoratori devono non solo essere competenti nell'uso degli strumenti di IA ma anche comprendere le implicazioni più ampie dell'IA per i loro ruoli e industrie. Datori di lavoro e educatori devono collaborare nello sviluppo di programmi di formazione che garantiscano che i lavoratori possano impegnarsi efficacemente con le tecnologie di IA, considerando anche le dimensioni etiche e sociali dell'IA nel luogo di lavoro.

In campi specializzati come la sanità, l'AI literacy è particolarmente critica. Studi di Laupichler et al. (2022) e Perchik et al. (2023) esplorano come l'IA stia trasformando la diagnostica medica e la pianificazione del trattamento, evidenziando la necessità per i professionisti sanitari di essere competenti nelle tecnologie di IA. Tuttavia, come con l'istruzione, c'è ancora una lacuna nell'integrazione delle considerazioni etiche nei programmi di formazione per i professionisti medici. Affrontare questa lacuna è essenziale per garantire che le tecnologie di IA siano utilizzate responsabilmente in contesti clinici.

Lo sviluppo di strumenti di valutazione è cruciale per misurare l'efficacia dei programmi di AI literacy. La ricerca di Laupichler et al. (2022) e Wang et al. (2022) ha mostrato che, mentre sono stati fatti sforzi per creare strumenti per valutare l'AI literacy, non c'è ancora consenso su cosa costituisca una valutazione completa. Strumenti di valutazione efficaci devono catturare sia le competenze tecniche che gli aspetti di pensiero critico ed etico dell'AI literacy, fornendo agli educatori i dati necessari per perfezionare e migliorare i curricula di AI literacy.

In definitiva, lo sviluppo dell'AI literacy è un processo multidimensionale che richiede l'integrazione di vari framework, modelli e strategie pedagogiche. Dalla prima infanzia alla formazione professionale, l'AI literacy deve comprendere sia competenze tecniche che considerazioni etiche per preparare gli individui alle sfide di un mondo guidato dall'IA. Mentre l'IA continua a ridisegnare industrie e società, è essenziale che i programmi di AI literacy evolvano in tandem, garantendo che gli apprendenti siano equipaggiati non solo per utilizzare le tecnologie di IA ma anche per comprendere e criticare le loro implicazioni più ampie. Concentrandosi sullo sviluppo di curricula completi, strumenti di valutazione efficaci e l'integrazione dell'AI literacy sia in ambienti di apprendimento formali che informali, possiamo creare una roadmap per l'AI literacy che garantisca che le generazioni future siano preparate a navigare in un panorama sempre più guidato dall'IA.

Nonostante il crescente interesse per l'AI literacy, la nostra revisione ha identificato diverse lacune significative nella letteratura esistente. In primo luogo, c'è una mancanza di ricerca empirica solida sull'efficacia di diversi approcci all'AI literacy. Mentre molti studi propongono framework teorici o curricula, relativamente pochi esaminano sistematicamente i risultati di questi interventi educativi. Sono necessarie più ricerche sperimentali e quasi-sperimentali per valutare quali approcci all'AI literacy sono più efficaci e in quali contesti.

In secondo luogo, c'è una significativa disparità geografica nella ricerca sull'AI literacy, con una predominanza di studi provenienti da Cina e Stati Uniti. Sono necessarie più ricerche da altre regioni, in particolare dall'Africa, dall'America Latina e da parti dell'Asia, per comprendere come l'AI literacy possa essere adattata a diversi contesti culturali, economici e politici. Questo è particolarmente importante data la natura pervasiva dell'IA e il suo potenziale impatto su scala globale.

In terzo luogo, mentre la dimensione etica dell'AI literacy è frequentemente menzionata, ci sono relativamente pochi studi dettagliati su come insegnare efficacemente l'etica dell'IA in diversi contesti educativi. Considerando l'importanza crescente delle considerazioni etiche nell'IA, questa rappresenta una significativa lacuna nella letteratura.

Infine, c'è una mancanza di ricerca longitudinale sull'AI literacy. La maggior parte degli studi fornisce istantanee di interventi o framework specifici, ma pochi esaminano come l'AI literacy si sviluppa nel tempo o come le diverse dimensioni dell'AI literacy interagiscano nel corso dello sviluppo di un individuo.

Alla luce di queste lacune, suggeriamo diverse direzioni per la ricerca futura:

Studi empirici sull'efficacia: Sono necessarie più ricerche che valutino sistematicamente l'efficacia di diversi approcci all'AI literacy, utilizzando metodologie robuste e misurazioni standardizzate.

Diversificazione geografica: La ricerca dovrebbe espandersi oltre i centri dominanti per includere prospettive da diverse regioni e contesti.

Approfondimento dell'etica dell'IA: Sono necessari più studi che esplorino specificamente come insegnare efficacemente gli aspetti etici dell'IA, in particolare in approcci di apprendimento formale.

Ricerca longitudinale: Gli studi dovrebbero esaminare come l'AI literacy si sviluppa nel tempo, dal ciclo della prima infanzia all'istruzione degli adulti.

Strumenti di valutazione standardizzati: Lo sviluppo e la validazione di strumenti di valutazione dell'AI literacy robusti e standardizzati rappresenterebbe un significativo contributo al campo.

Interdisciplinarietà: Maggiore ricerca interdisciplinare che attinga da campi come le scienze dell'educazione, l'informatica, la filosofia e le scienze sociali potrebbe fornire intuizioni più ricche sulla complessa natura dell'AI literacy.

La nostra revisione sistematica della letteratura sull'AI literacy ha fornito una panoramica completa dello stato attuale della ricerca in questo campo emergente. L'analisi ha rivelato un interesse crescente per l'AI literacy, con un significativo aumento delle pubblicazioni negli ultimi anni, in particolare dal 2021 al 2023. Questo trend riflette la crescente rilevanza dell'IA in vari settori della società e la conseguente necessità di sviluppare competenze che permettano agli individui di navigare efficacemente in un mondo sempre più guidato dall'IA.

La distribuzione geografica delle pubblicazioni ha evidenziato una predominanza di studi provenienti da Cina e Stati Uniti, seguiti da paesi europei come la Germania. Questa distribuzione riflette le diverse priorità nazionali e approcci all'IA, con la Cina che enfatizza l'integrazione dell'AI literacy nell'educazione K-12 come parte della sua strategia nazionale, e gli Stati Uniti che si concentrano maggiormente sull'istruzione superiore e lo sviluppo professionale. La diversità di approcci all'AI literacy in diverse regioni sottolinea l'importanza di considerare contesti culturali, politici ed economici specifici nello sviluppo di programmi di AI literacy.

L'analisi delle definizioni di AI literacy presenti in letteratura ha rivelato una gamma diversificata di interpretazioni, che va dalla comprensione tecnica dei meccanismi dell'IA a prospettive più ampie che includono dimensioni etiche, sociali e culturali. Questa diversità di interpretazioni riflette la natura complessa dell'AI literacy e la necessità di approcci che integrino sia competenze tecniche che consapevolezza critica delle implicazioni sociali dell'IA.

L'analisi tematica ha identificato dieci temi principali nella letteratura sull'AI literacy, tra cui lo sviluppo di curriculum e interventi, modelli e framework, implicazioni etiche e sociali, e applicazioni in contesti specifici come l'istruzione K-12, gli ambienti sanitari e gli ambienti familiari. La predominanza di temi relativi allo sviluppo di curriculum e ai framework teorici sottolinea la fase

ancora emergente di questo campo, con un forte focus sulla creazione di basi concettuali e pratiche per l'educazione all'IA.

Le implicazioni di questa revisione sono significative per educatori, ricercatori e policymaker. Per gli educatori, i risultati evidenziano la necessità di approcci integrati all'AI literacy che bilancino competenze tecniche e consapevolezza critica delle dimensioni sociali ed etiche dell'IA. Per i ricercatori, la revisione identifica importanti lacune nella letteratura esistente, in particolare la necessità di più ricerca empirica sull'efficacia degli interventi di AI literacy, studi che esplorino l'AI literacy in diverse regioni geografiche, e ricerche longitudinali sullo sviluppo dell'AI literacy nel tempo. Per i policymaker, i risultati sottolineano l'importanza di integrare l'AI literacy nei curricula educativi a tutti i livelli e di promuovere approcci che preparino i cittadini a partecipare criticamente e consapevolmente a un mondo sempre più guidato dall'IA.

In conclusione, mentre la ricerca sull'AI literacy è ancora in una fase relativamente emergente, c'è un crescente consenso sulla sua importanza come competenza fondamentale per il XXI secolo. La diversità di definizioni, approcci e tematiche presenti in letteratura riflette la natura complessa e multidimensionale dell'AI literacy, che richiede framework integrati che abbraccino sia dimensioni tecniche che socio-etiche. Il framework concettuale presentato nella prossima sezione si propone di rispondere a questa esigenza, offrendo un approccio strutturato all'AI literacy che integra le diverse prospettive emerse dalla letteratura.

1.3 Un quadro concettuale per l'AI literacy

1.3.1 Dal concetto di literacy all'AI literacy: evoluzione e trasformazione

Il concetto di literacy ha subito trasformazioni significative nel corso dell'ultimo secolo, particolarmente in risposta alla rapida digitalizzazione della società. Inizialmente concettualizzata come un insieme di abilità cognitive – comprendenti lettura, scrittura e aritmetica – la literacy era vista come una competenza tecnica neutrale che gli individui acquisivano attraverso l'educazione formale (Ong, 1982). In questa visione funzionalista, la literacy era legata ai processi istituzionali dell'educazione, dove gli individui non solo acquisivano abilità cognitive ma anche interiorizzavano norme e valori sociali che promuovevano la mobilità sociale e l'inclusione (Cappello, 2017).

Tuttavia, alla fine degli anni '70, questa comprensione della literacy iniziò ad essere messa in discussione. Studiosi all'interno del movimento New Literacies Studies (NLS), come Street (2003), proposero che la literacy non dovesse essere vista come un'abilità puramente cognitiva o ideologicamente neutra, ma piuttosto come una pratica sociale situata. In questa prospettiva, la literacy è intesa come un repertorio di pratiche che variano attraverso diversi contesti sociali, modellate dalle specifiche necessità, valori e convenzioni di particolari gruppi. La literacy non riguarda solo l'acquisizione di abilità tecniche, ma la partecipazione e il contributo alle specifiche pratiche culturali e sociali di una comunità (Gee, 2003). Questo cambiamento ha ridefinito la literacy come una pratica dinamica ed evolutiva, profondamente radicata nei contesti culturali, sociali e politici in cui è performata.

Alla fine degli anni '90, un secondo significativo cambiamento nel concetto di literacy si verificò, guidato dalla rapida proliferazione delle tecnologie digitali. La cosiddetta "svolta digitale" (Mills, 2010) segnò l'emergere di nuove literacies mediali o multiliteracies, un concetto sviluppato da studiosi come Cope & Kalantzis (2000) e Kress (2000). Le multiliteracies riflettono la crescente complessità della comunicazione negli ambienti digitali, dove gli individui devono impegnarsi con testi multimodali che combinano linguaggio scritto con elementi audio, visivi e spaziali. Queste nuove forme di literacy evidenziano la necessità per gli individui non solo di interpretare testi scritti

tradizionali ma anche di navigare e creare contenuti digitali e multimodali attraverso una varietà di piattaforme e contesti sociali (Lankshear & Knobel, 2003; Sefton-Green, 2007).

Questa duplice espansione della literacy – prima come pratica sociale e poi come competenza medialmente ampliata – ha avuto profonde implicazioni per come comprendiamo la literacy nell'era digitale. Come nota Ranieri (2019), la digital literacy oggi comprende un'ampia gamma di competenze, dalle basi dell'uso degli strumenti digitali a abilità più complesse come la valutazione critica dei contenuti digitali e la comprensione delle implicazioni etiche delle interazioni digitali. In questo senso, la digital literacy è andata oltre i compiti cognitivi di base per includere le dimensioni sociali, culturali ed etiche dell'interazione in un mondo digitalizzato.

Tuttavia, con le tecnologie di IA che continuano a rimodellare il panorama digitale, c'è un crescente riconoscimento che i concetti tradizionali di digital literacy non sono più sufficienti. L'ascesa dei sistemi guidati dall'IA introduce nuove sfide che vanno oltre l'ambito della digital literacy come tradizionalmente definita. I sistemi di IA, che spesso operano autonomamente e prendono decisioni basate su algoritmi complessi e grandi dataset, richiedono agli individui di sviluppare una comprensione più profonda di come queste tecnologie funzionano e quali implicazioni comportano per la società.

Costruendo sull'evoluzione della literacy come pratica sociale situata, l'AI literacy deve anch'essa essere compresa come una competenza sia tecnica che contestualmente radicata. Non è sufficiente per gli individui semplicemente sapere come usare gli strumenti di IA; devono anche sviluppare la consapevolezza critica necessaria per impegnarsi con le dimensioni etiche, sociali e culturali delle tecnologie di IA. Come sosteneva Street (2003) per la literacy tradizionale, l'AI literacy deve similmente essere situata all'interno dei specifici contesti sociali in cui le tecnologie di IA sono sviluppate, implementate e utilizzate.

Pertanto, proprio come la digital literacy si è espansa per accogliere la natura multimodale e connessa della comunicazione nell'era digitale, l'AI literacy deve espandersi per affrontare le complessità introdotte dai sistemi intelligenti e autonomi. Queste complessità includono non solo la comprensione tecnica di come operano i sistemi di IA ma anche la capacità di valutare criticamente i loro impatti sociali più ampi – come i bias algoritmici, le preoccupazioni sulla privacy e i dilemmi etici posti dai processi decisionali guidati dall'IA (Bostrom & Yudkowsky, 2014).

In sintesi, l'AI literacy emerge come una naturale evoluzione nella traiettoria della literacy nel mondo digitale, richiedendo un'integrazione di competenze tecniche, consapevolezza critica e sensibilità etico-sociale. Come vedremo nelle sezioni seguenti, questa concezione multidimensionale dell'AI literacy forma la base per il framework concettuale proposto, che mira a fornire una struttura comprensiva per lo sviluppo di programmi educativi che preparino gli individui a navigare efficacemente in un mondo sempre più plasmato dall'intelligenza artificiale.

1.3.2 Fondamenti teorici del framework dell'AI literacy

Il framework dell'AI literacy proposto in questo lavoro si basa su solide fondamenta teoriche, attingendo a diverse tradizioni di ricerca e incorporando intuizioni da campi come la digital literacy, l'educazione tecnologica, l'etica dell'IA e la pedagogia critica. Questa sezione esplora i principali pilastri teorici che sostengono il framework, evidenziando come questi diversi filoni di pensiero convergono per formare un approccio comprensivo all'AI literacy.

Un punto di partenza fondamentale per il nostro framework è il modello di competenza digitale sviluppato da Calvani, Fini e Ranieri (2009), che concettualizza la competenza digitale come un

costrutto multidimensionale comprendente tre dimensioni principali: la dimensione tecnologica, la dimensione cognitiva e la dimensione etica.

La dimensione tecnologica si concentra sulle abilità tecnico-operative necessarie per utilizzare efficacemente gli strumenti digitali. Enfatizza l'importanza dell'agilità nell'adattarsi a tecnologie in rapida evoluzione e la capacità di esplorare e adottare nuovi strumenti (Calvani et al., 2009). La dimensione tecnologica non implica meramente la conoscenza di strumenti specifici ma promuove anche una mentalità flessibile per navigare nel panorama digitale.

La dimensione cognitiva si riferisce al pensiero critico e alla valutazione delle informazioni, richiedendo agli individui non solo di elaborare informazioni ma di discernere la loro affidabilità, organizzare dati sistematicamente e creare modelli astratti da essi (Calvani et al., 2009). La dimensione cognitiva risuona con definizioni precedenti di digital literacy, come quella di Paul Gilster (1997), che enfatizzava la capacità di impegnarsi criticamente con i contenuti digitali e formulare giudizi informati sulle informazioni incontrate.

La dimensione etica si concentra sulle responsabilità sociali ed etiche associate agli ambienti digitali. Sottolinea la necessità di comprendere la netiquette e i diritti alla privacy, di garantire un comportamento rispettoso online e di navigare nelle sfide etiche poste dalle interazioni digitali (Calvani et al., 2009). Questa dimensione rispecchia la più ampia conversazione sull'etica digitale e il comportamento online, una questione centrale nei framework come la media literacy (Buckingham, 2003) e il consumo etico dei contenuti digitali.

L'approccio tridimensionale di Calvani, Fini e Ranieri alla competenza digitale fornisce un utile punto di partenza per concettualizzare l'AI literacy. Riconosce che la competenza nelle tecnologie digitali non è solo questione di abilità tecniche ma anche di capacità cognitive e consapevolezza etica. Tuttavia, mentre questo modello fornisce una solida base, deve essere adattato e ampliato per affrontare le specificità e le complessità uniche dell'IA. Le prossime sezioni esplorano come il nostro framework costruisce su questo modello, incorporando intuizioni da altre tradizioni teoriche e adattandosi alle caratteristiche distintive delle tecnologie di IA.

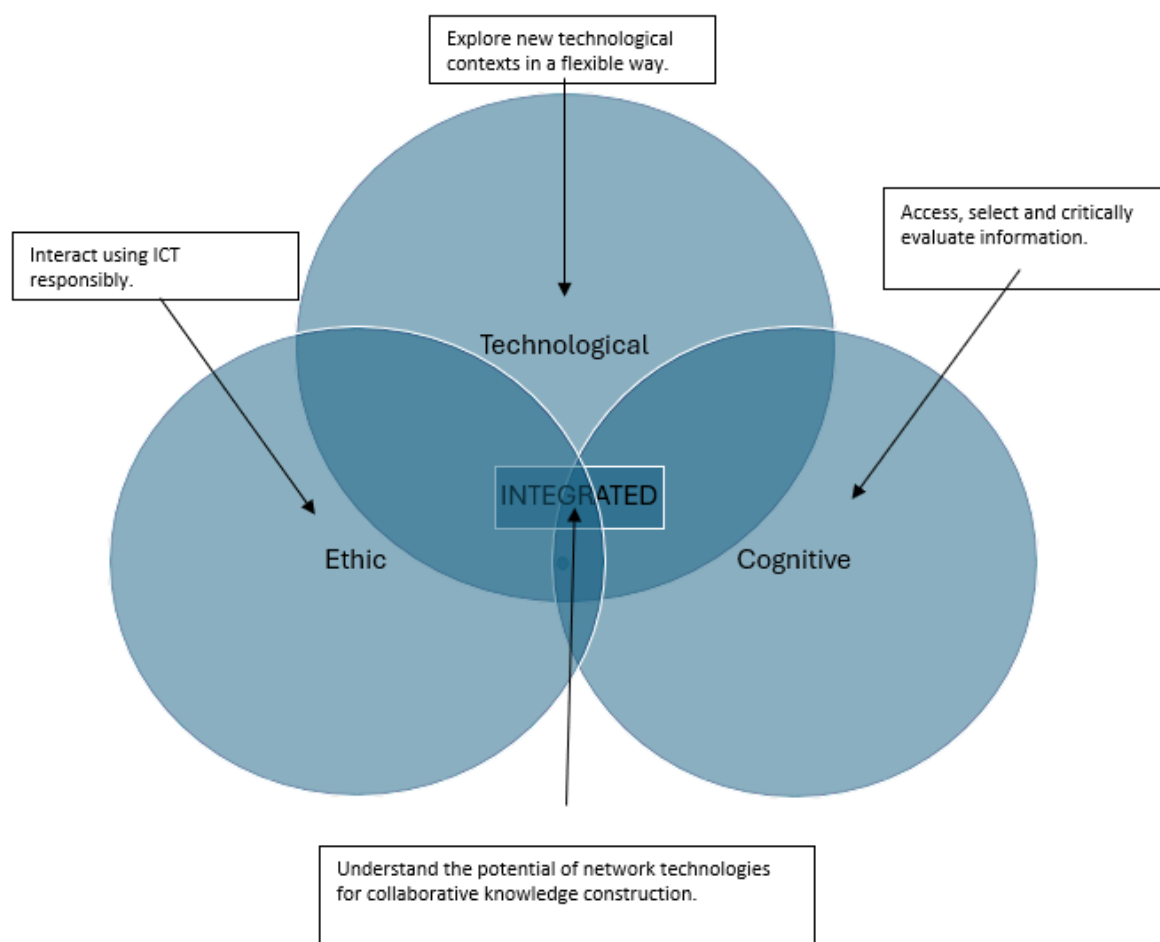


Figura 3

Negli ultimi anni, diversi ricercatori hanno iniziato a sviluppare framework specifici per l'AI literacy, fornendo preziose intuizioni per la nostra concettualizzazione. Long & Magerko (2020) offrono uno dei primi tentativi completi di definire l'AI literacy, descrivendola come un insieme di competenze che consentono agli individui di valutare criticamente le tecnologie di IA, comunicare e collaborare efficacemente con l'IA, e utilizzare l'IA in vari contesti. Il loro lavoro enfatizza non solo la comprensione tecnica dell'IA ma anche la capacità di interagire con essa in modi critici e produttivi.

Ng et al. (2021) hanno ulteriormente raffinato il concetto, proponendo un framework che organizza l'AI literacy in quattro aree chiave: conoscere e comprendere l'IA, utilizzare e applicare l'IA, valutare e creare IA, e l'etica dell'IA. Questo framework completo riconosce la natura multidimensionale dell'AI literacy e l'importanza di integrare considerazioni etiche nelle competenze tecniche.

Altri contributi significativi alla concettualizzazione dell'AI literacy includono il lavoro di Yi (2021), che enfatizza la dimensione culturale e soggettiva dell'AI literacy, e Deuze & Beckett (2022), che sottolineano la sua dimensione normativa. Questi diversi framework, pur differendo nei dettagli specifici, convergono nell'enfatizzare la natura multidimensionale dell'AI literacy e la necessità di un approccio che integri competenze tecniche, pensiero critico e consapevolezza etica.

Il nostro framework costruisce su questi contributi emergenti, sintetizzando le loro intuizioni in un approccio coerente e comprensivo all'AI literacy. In particolare, adattiamo la struttura tridimensionale di Calvani et al. (2009) per includere una quarta dimensione che affronta specificamente la natura

operativa dell'IA, riflettendo la crescente enfasi nella letteratura sull'importanza non solo di conoscere l'IA a livello concettuale ma anche di essere in grado di lavorare con essa in modo pratico e applicato.

Il nostro framework è inoltre informato da teorie critiche della tecnologia che sottolineano l'importanza di riconoscere le dimensioni sociali, politiche ed etiche delle tecnologie. Studiosi come Feenberg (1991, 2002) argomentano che le tecnologie non sono strumenti neutrali ma sono intrinsecamente intrecciate con valori, relazioni di potere e strutture sociali. Questa prospettiva critica è particolarmente rilevante per l'IA, data la sua crescente influenza su processi decisionali che impattano significativamente la vita delle persone.

Parallelamente, teorie pedagogiche critiche, in particolare il lavoro di Freire (1970, 1973) sulla pedagogia della coscientizzazione e dell'empowerment, informano il nostro approccio all'educazione all'IA. Queste teorie sottolineano l'importanza di sviluppare una coscienza critica che permetta agli individui di riconoscere e sfidare strutture e pratiche ingiuste, e di impegnarsi attivamente nella trasformazione della loro realtà sociale. Nel contesto dell'AI literacy, questo si traduce in un'enfasi sullo sviluppo della capacità degli apprendenti di riconoscere e affrontare questioni come i bias algoritmici, le disuguaglianze nell'accesso e nel controllo delle tecnologie di IA, e le implicazioni etiche dell'automazione e dei processi decisionali basati su algoritmi.

In sintesi, il framework dell'AI literacy proposto in questo lavoro integra intuizioni da diverse tradizioni teoriche, tra cui modelli di competenza digitale, framework emergenti dell'AI literacy, teorie critiche della tecnologia e teorie pedagogiche critiche. Questa base teorica multidisciplinare fornisce un fondamento solido per un approccio all'AI literacy che va oltre la mera competenza tecnica per abbracciare anche la comprensione critica, l'applicazione pratica e la consapevolezza etico-sociale.

1.3.3 Le quattro dimensioni del framework dell'AI literacy

Il framework dell'AI literacy proposto in questo lavoro si articola in quattro dimensioni fondamentali: conoscitiva, operativa, critica ed etica. Questa struttura quadridimensionale, che rappresenta un'evoluzione del modello tridimensionale della competenza digitale di Calvani et al. (2009), è stata sviluppata per rispondere alle specifiche esigenze e sfide poste dall'intelligenza artificiale. In questa sezione, esploriamo in dettaglio ciascuna di queste dimensioni, analizzando le competenze che esse comprendono e la loro rilevanza nel contesto dell'AI literacy.

La dimensione conoscitiva dell'AI literacy si concentra sulla comprensione dei fondamenti teorici e concettuali dell'intelligenza artificiale. Questa dimensione è essenziale perché fornisce la base di conoscenza necessaria per interagire efficacemente con le tecnologie di IA e per sviluppare una comprensione critica delle loro capacità e limitazioni. Seguendo la concettualizzazione di Ng et al. (2021), questa dimensione comprende la conoscenza e la comprensione dell'IA, ovvero le funzioni di base dell'IA e come utilizzare le applicazioni di IA. Più specificamente, può essere articolata in diverse componenti:

Conoscenza delle definizioni e dei tipi di IA: Questa componente include la comprensione di cosa sia l'IA, la differenza tra IA forte e debole, e la distinzione tra diverse tecniche e approcci all'IA, come il machine learning, le reti neurali, e i sistemi basati su regole. Come sottolineato da Kandlhofer et al. (2016), questa comprensione di base è fondamentale per distinguere l'IA da altre tecnologie digitali.

Comprensione dei principi del machine learning e dell'elaborazione dei dati: Questa componente riguarda la comprensione di come i sistemi di IA apprendono dai dati, come riconoscono pattern e come utilizzano questi pattern per fare previsioni o prendere decisioni. Include anche la conoscenza

dei diversi tipi di machine learning (supervisionato, non supervisionato, per rinforzo) e dei loro rispettivi approcci all'apprendimento.

Conoscenza delle applicazioni dell'IA: Questa componente si riferisce alla familiarità con le diverse applicazioni dell'IA in vari settori, come la sanità, l'educazione, la finanza e l'intrattenimento. Include la consapevolezza di come l'IA viene applicata in settori specifici e dei potenziali benefici e limitazioni di queste applicazioni.

Comprensione delle capacità e dei limiti dell'IA: Questa componente riguarda la conoscenza di ciò che l'IA può e non può fare, delle sue attuali capacità e delle sfide che deve ancora superare. Include la comprensione di concetti come il problema dell'allineamento, la generalizzazione, e le limitazioni dell'apprendimento dai dati.

La dimensione conoscitiva dell'AI literacy è stata enfatizzata da numerosi studiosi nel campo. Ad esempio, Long & Magerko (2020) includono nella loro definizione di AI literacy la capacità di "comprendere i concetti fondamentali dell'IA", mentre Kim et al. (2021) parlano di "AI Knowledge" come una delle tre competenze fondamentali per l'AI literacy. Questa dimensione è particolarmente importante in quanto fornisce le basi concettuali necessarie per le altre dimensioni dell'AI literacy.

È importante notare che la dimensione conoscitiva dell'AI literacy non richiede necessariamente una comprensione tecnica approfondita o abilità di programmazione avanzate. Piuttosto, si concentra su una comprensione concettuale di alto livello dei principi e dei meccanismi dell'IA, accessibile a un pubblico non specializzato. Come suggeriscono Burgsteiner et al. (2016), è più importante per i non esperti capire "cosa può fare l'IA" piuttosto che "come funziona l'IA" a un livello tecnico dettagliato.

La dimensione operativa dell'AI literacy si concentra sulle abilità pratiche necessarie per utilizzare e interagire con le tecnologie di IA. Questa dimensione va oltre la comprensione concettuale per includere la capacità di applicare la conoscenza dell'IA in contesti reali. Seguendo la concettualizzazione di Ng et al. (2021), questa dimensione comprende l'uso e l'applicazione dell'IA, ovvero l'applicazione della conoscenza, dei concetti e delle applicazioni dell'IA in diversi scenari.

La dimensione operativa può essere articolata in diverse componenti:

Utilizzo di strumenti e applicazioni di IA: Questa componente si riferisce alla capacità di utilizzare efficacemente vari strumenti e applicazioni basati sull'IA, come assistenti virtuali, sistemi di raccomandazione, e strumenti di analisi dei dati. Include la conoscenza di come interagire con queste tecnologie, come interpretare i loro output, e come sfruttare le loro funzionalità per raggiungere obiettivi specifici.

Competenze di base in programmazione e data science: Mentre non tutti gli utenti di IA necessitano di competenze avanzate di programmazione, una comprensione di base di concetti come gli algoritmi, le strutture di dati e i principi di data science può essere utile per interagire più efficacemente con le tecnologie di IA. Come suggeriscono Kim et al. (2021), questa componente include il "pensiero computazionale e la programmazione", ovvero la capacità di costruire semplici applicazioni di intelligenza artificiale.

Applicazione dell'IA per risolvere problemi: Questa componente si riferisce alla capacità di identificare problemi o sfide che potrebbero beneficiare dell'applicazione dell'IA e di selezionare e implementare soluzioni appropriate basate sull'IA. Include la capacità di valutare quando l'IA è la

soluzione più adatta a un determinato problema e come integrare l'IA in un processo di risoluzione dei problemi più ampio.

Collaborazione con sistemi di IA: Questa componente, enfatizzata da Long & Magerko (2020), si riferisce alla capacità di lavorare efficacemente con sistemi di IA, di comprendere come questi sistemi possono complementare le capacità umane, e di sviluppare modalità di interazione efficaci con l'IA. Include la comprensione dei ruoli rispettivi degli umani e dell'IA in un sistema collaborativo e di come massimizzare i benefici di questa collaborazione.

La dimensione operativa dell'AI literacy è stata enfatizzata da diversi studiosi nel campo. Ad esempio, Druga et al. (2019) parlano della capacità di "utilizzare in modo etico le applicazioni dell'IA nella vita di tutti i giorni", mentre Long & Magerko (2020) includono nella loro definizione di AI literacy la capacità di "usare l'IA come uno strumento online, a casa e sul posto di lavoro". Questa dimensione è particolarmente importante in quanto consente agli individui di applicare la loro comprensione concettuale dell'IA a situazioni pratiche, trasformando la conoscenza teorica in competenze applicate.

È importante notare che la dimensione operativa dell'AI literacy non mira a trasformare tutti gli utenti in esperti di IA o programmatori. Piuttosto, si concentra sullo sviluppo di un livello di competenza pratica che consenta agli individui di utilizzare in modo efficace ed efficiente le tecnologie di IA disponibili, di adattarsi a nuove tecnologie man mano che emergono, e di identificare opportunità per l'applicazione dell'IA nella loro vita personale e professionale.

La dimensione critica dell'AI literacy si concentra sulla capacità di analizzare, valutare e interpretare criticamente le tecnologie di IA, i loro output e il loro impatto sulla società. Questa dimensione è essenziale per sviluppare un'interazione consapevole e informata con l'IA, che vada oltre l'uso passivo per includere una comprensione più profonda delle implicazioni e delle limitazioni di queste tecnologie.

La dimensione critica può essere articolata in diverse componenti:

Valutazione della qualità e dell'affidabilità dei sistemi di IA: Questa componente si riferisce alla capacità di giudicare criticamente la qualità, l'accuratezza e l'affidabilità dei sistemi di IA e dei loro output. Include la comprensione dei potenziali errori, bias e limitazioni dei sistemi di IA, e la capacità di valutare se un determinato sistema è adatto a uno scopo specifico. Come sottolineano Su & Zhong (2022), questa componente è fondamentale per un uso responsabile dell'IA.

Comprensione dei bias algoritmici e delle questioni di equità: Questa componente riguarda la capacità di riconoscere e comprendere i bias che possono essere presenti nei sistemi di IA, sia a causa dei dati su cui sono stati addestrati sia a causa delle decisioni prese durante il loro sviluppo. Include la comprensione di come questi bias possono portare a risultati iniqui o discriminatori, e la capacità di identificare strategie per mitigare questi bias. Questa componente è stata evidenziata da ricercatori come Strauß (2021), che sottolinea l'importanza di una "AI literacy critica per l'uso costruttivo della tecnologia IA".

Interpretazione e spiegabilità dell'IA: Questa componente si riferisce alla capacità di comprendere e interpretare come i sistemi di IA arrivano a determinate conclusioni o decisioni. Include la comprensione delle sfide legate alla "black box" dell'IA, in cui i processi decisionali non sono trasparenti, e la capacità di richiedere e comprendere spiegazioni per le decisioni dell'IA. Come

suggeriscono Hermann (2022) e Carolus et al. (2023), questa componente è cruciale per lo sviluppo di fiducia nei sistemi di IA.

Analisi dell'impatto sociale dell'IA: Questa componente riguarda la capacità di analizzare e comprendere l'impatto più ampio dell'IA sulla società, includendo questioni come la sorveglianza, la privacy, l'autonomia, e i cambiamenti nel mercato del lavoro. Include la capacità di considerare le implicazioni a lungo termine dell'adozione diffusa dell'IA e di partecipare a dibattiti informati su come gestire questi cambiamenti. Questa componente è stata enfatizzata da ricercatori come Yi (2021) e Long et al. (2023), che sottolineano l'importanza di comprendere l'IA nel suo contesto sociale e culturale.

La dimensione critica dell'AI literacy è stata enfatizzata da numerosi studiosi nel campo. Ad esempio, Long & Magerko (2020) includono nella loro definizione di AI literacy la capacità di "valutare criticamente le tecnologie di IA", mentre Ng et al. (2021) parlano di "valutare e creare IA" come uno dei quattro concetti chiave dell'AI literacy. Questa dimensione è particolarmente importante in quanto consente agli individui di andare oltre l'uso passivo dell'IA per sviluppare una comprensione più profonda e critica di queste tecnologie e del loro impatto.

È importante notare che la dimensione critica dell'AI literacy non implica necessariamente una posizione negativa o scettica nei confronti dell'IA. Piuttosto, promuove un'interazione riflessiva e informata con l'IA, che riconosce sia i suoi potenziali benefici che le sue limitazioni e rischi. Come suggeriscono Shih et al. (2021), l'obiettivo è sviluppare un "discernimento critico" che permetta agli individui di prendere decisioni informate sull'uso, lo sviluppo e la regolamentazione dell'IA.

La dimensione etica dell'AI literacy si concentra sulla comprensione e l'applicazione di principi e valori etici nell'interazione con le tecnologie di IA. Questa dimensione riconosce che l'IA solleva questioni etiche significative, dalla privacy e autonomia alla responsabilità e equità, e che gli individui devono essere equipaggiati per navigare queste questioni in modo responsabile e consapevole.

La dimensione etica può essere articolata in diverse componenti:

Comprensione dei principi etici dell'IA: Questa componente si riferisce alla conoscenza dei principi etici fondamentali relativi all'IA, come l'equità, la trasparenza, la responsabilità, la privacy e la beneficenza. Include la comprensione di come questi principi possono essere applicati nel contesto dell'IA e di come possono talvolta entrare in conflitto tra loro. Come sottolineano Ng et al. (2021) e UNESCO (2022), questa componente fornisce un quadro concettuale per valutare le implicazioni etiche dell'IA.

Consapevolezza delle questioni di privacy e sicurezza: Questa componente riguarda la comprensione di come l'IA può impattare sulla privacy degli individui, sia attraverso la raccolta e l'analisi di dati personali sia attraverso la creazione di profili dettagliati degli utenti. Include la conoscenza di strategie per proteggere la propria privacy nel contesto dell'IA e la comprensione delle implicazioni più ampie per la privacy sociale. Come suggeriscono Laupichler et al. (2022) e Wang et al. (2022), questa componente è fondamentale per un uso consapevole dell'IA.

Valutazione delle questioni di responsabilità e accountability: Questa componente si riferisce alla capacità di comprendere chi è responsabile per le decisioni prese dai sistemi di IA e come questa responsabilità dovrebbe essere distribuita tra sviluppatori, utenti e altre parti interessate. Include la comprensione dei meccanismi per garantire l'accountability dell'IA e delle sfide legate all'attribuzione

di responsabilità in sistemi complessi e autonomi. Come sottolineano Deuze & Beckett (2022) e Floridi (2018), questa componente è cruciale per la governance responsabile dell'IA.

Impegno con questioni di giustizia sociale e equità: Questa componente riguarda la comprensione di come l'IA può influenzare la giustizia sociale e l'equità, sia positivamente che negativamente. Include la capacità di valutare criticamente l'impatto dell'IA su diverse comunità e gruppi sociali, di identificare potenziali disuguaglianze o ingiustizie perpetuate dai sistemi di IA, e di partecipare a dibattiti su come l'IA può essere sviluppata e utilizzata in modi che promuovano, piuttosto che minino, la giustizia sociale. Questa componente è stata enfatizzata da ricercatori come Long et al. (2023) e Hermann (2022), che sottolineano l'importanza di considerare le implicazioni sociali più ampie dell'IA.

La dimensione etica dell'AI literacy è stata enfatizzata da numerosi studiosi nel campo. Ad esempio, Druga et al. (2019) includono nella loro definizione di AI literacy la capacità di "utilizzare le applicazioni dell'IA nella vita quotidiana in modo eticamente adeguato", mentre Wong et al. (2020) parlano di "AI Ethics and Safety" come una delle tre componenti principali dell'AI literacy. Questa dimensione è particolarmente importante in quanto riconosce che l'IA non è solo una questione tecnica ma anche una questione profondamente etica e sociale, che richiede una considerazione attenta di valori, principi e responsabilità.

È importante notare che la dimensione etica dell'AI literacy non implica necessariamente l'aderenza a un particolare insieme di valori o principi etici. Piuttosto, promuove la capacità di riflettere criticamente sulle implicazioni etiche dell'IA, di articolare i propri valori in relazione all'IA, e di impegnarsi in discussioni informate e rispettose su questioni etiche complesse. Come suggeriscono Shih et al. (2021) e Yi (2021), l'obiettivo è sviluppare individui che siano non solo tecnicamente competenti ma anche eticamente consapevoli e responsabili nell'età dell'IA.

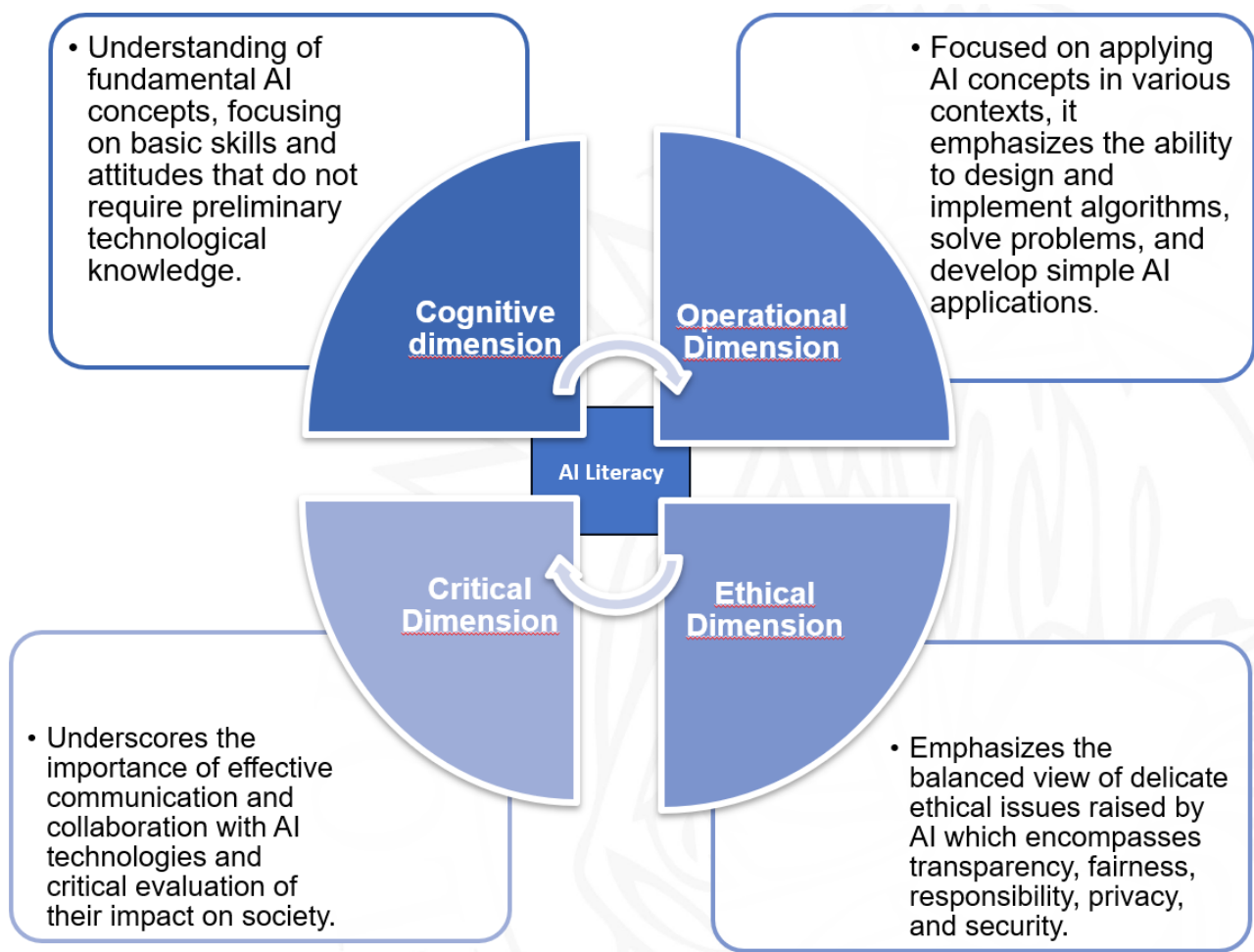


Figura 4

1.3.4 Altri framework e modelli

Il framework quadridimensionale per l'AI literacy proposto in questo lavoro si inserisce in un panorama più ampio di modelli e framework che cercano di concettualizzare le competenze necessarie per interagire efficacemente con l'intelligenza artificiale. In questa sezione, confrontiamo il nostro framework con altri modelli esistenti, evidenziando similarità, differenze e contributi distintivi.

Il framework di Long & Magerko (2020) è uno dei primi tentativi completi di definire l'AI literacy, descrivendola come un insieme di competenze che consentono agli individui di valutare criticamente le tecnologie di IA, comunicare e collaborare efficacemente con l'IA, e utilizzare l'IA in vari contesti. Questo framework identifica otto competenze chiave dell'AI literacy:

1. Riconoscere l'IA
2. Comprendere l'intelligenza
3. Interdisciplinarietà
4. Rappresentazione dei dati
5. Riconoscere il processo decisionale dell'IA
6. Promuovere la trasparenza

7. Pensiero critico

8. Apprendere dall'IA

Confrontando questo framework con il nostro, possiamo osservare alcune similarità significative. Entrambi i framework riconoscono l'importanza di una comprensione concettuale dell'IA (corrispondente alla nostra dimensione conoscitiva), della capacità di utilizzare l'IA in modo pratico (dimensione operativa), e della valutazione critica dell'IA (dimensione critica). Inoltre, competenze come "promuovere la trasparenza" e "pensiero critico" risuonano con la nostra dimensione etica.

Tuttavia, ci sono anche differenze notevoli. Il framework di Long & Magerko è strutturato intorno a competenze specifiche piuttosto che a dimensioni più ampie, e non distingue esplicitamente tra competenze operative e critiche. Inoltre, mentre include considerazioni etiche, queste non sono articolate come una dimensione distinta dell'AI literacy. Il nostro framework, d'altra parte, dà pari peso alle quattro dimensioni e fornisce un quadro più ampio per organizzare e comprendere le diverse competenze necessarie per l'AI literacy.

Ng et al. (2021) propongono un framework che organizza l'AI literacy in quattro aree chiave: conoscere e comprendere l'IA, utilizzare e applicare l'IA, valutare e creare IA, e l'etica dell'IA. Questo framework è particolarmente allineato con il nostro, condividendo una struttura a quattro dimensioni che comprende aspetti conoscitivi, operativi, valutativi ed etici.

Le prime due aree del framework di Ng et al., "conoscere e comprendere l'IA" e "utilizzare e applicare l'IA", corrispondono strettamente alle nostre dimensioni conoscitiva e operativa. Similmente, "valutare e creare IA" e "l'etica dell'IA" risuonano con le nostre dimensioni critica ed etica. Questa convergenza suggerisce un emergente consenso nella letteratura sulla natura multidimensionale dell'AI literacy e sull'importanza di includere sia competenze tecniche che considerazioni etiche e critiche.

Tuttavia, mentre il framework di Ng et al. si allinea strettamente con il nostro in termini di struttura, ci sono differenze nell'articolazione specifica delle competenze all'interno di ciascuna area e nell'enfasi relativa posta su diverse componenti. Ad esempio, il nostro framework pone una maggiore enfasi sulla dimensione critica, articolando in modo più dettagliato le competenze necessarie per valutare criticamente i sistemi di IA, i loro output e il loro impatto sociale.

Kim et al. (2021) propongono un approccio tripartito all'AI literacy, comprendente conoscenza, abilità e atteggiamento. La conoscenza è definita come la capacità di comprendere i concetti fondamentali dell'IA; l'abilità comprende il pensiero computazionale o le competenze di programmazione; mentre l'atteggiamento si riferisce alla capacità di considerare collettivamente tutti gli aspetti dell'IA nella società.

Confrontando questo framework con il nostro, possiamo osservare che la "conoscenza" di Kim et al. corrisponde alla nostra dimensione conoscitiva, mentre la loro "abilità" si allinea con la nostra dimensione operativa. Il loro "atteggiamento" comprende elementi sia della nostra dimensione critica che di quella etica, suggerendo un approccio più integrato a questi aspetti.

Una differenza significativa è che il framework di Kim et al. non distingue esplicitamente tra pensiero critico ed etica, includendoli entrambi nella categoria più ampia di "atteggiamento". Il nostro framework, d'altra parte, tratta questi come dimensioni distinte ma interconnesse dell'AI literacy, riconoscendo che il pensiero critico e la consapevolezza etica, sebbene correlati, richiedono competenze e approcci educativi distinti.

Oltre ai framework specificamente focalizzati sull'AI literacy, il nostro modello può essere confrontato con approcci più ampi alla competenza digitale e alla media literacy. Ad esempio, il DigComp 2.1 dell'Unione Europea (Carretero et al., 2017) identifica cinque aree di competenza digitale: informazione e data literacy, comunicazione e collaborazione, creazione di contenuti digitali, sicurezza, e problem solving. Mentre queste aree non sono specificamente focalizzate sull'IA, molte delle competenze che includono sono rilevanti per l'AI literacy, e c'è una considerevole sovrapposizione con le nostre dimensioni, in particolare per quanto riguarda l'informazione e data literacy (dimensione conoscitiva), la creazione di contenuti digitali (dimensione operativa), e la sicurezza (dimensione etica).

Similmente, il framework P21 per le competenze del XXI secolo (P21, 2019) include pensiero critico, creatività, comunicazione e collaborazione come "competenze di apprendimento e innovazione" chiave, molte delle quali sono rilevanti per l'AI literacy. Il nostro framework incorpora queste competenze, in particolare all'interno delle dimensioni operativa e critica, ma le contestualizza specificamente nel dominio dell'IA.

1.3.5 Operazionalizzazione e implementazione del framework

Il framework quadridimensionale per l'AI literacy proposto in questo lavoro fornisce una base concettuale solida per lo sviluppo di iniziative educative volte a promuovere la comprensione e l'uso consapevole dell'intelligenza artificiale. In questa sezione, esploriamo come questo framework può essere operazionalizzato e implementato in vari contesti educativi, delineando principi guida, approcci pedagogici e considerazioni pratiche.

L'implementazione efficace del framework dell'AI literacy dovrebbe essere guidata da diversi principi fondamentali:

Approccio olistico e integrato: Le quattro dimensioni del framework dovrebbero essere trattate come componenti interconnesse di un approccio olistico all'AI literacy, piuttosto che come moduli separati. L'educazione all'IA dovrebbe mirare a sviluppare contemporaneamente competenze conoscitive, operative, critiche ed etiche, riconoscendo le interrelazioni tra queste dimensioni. Come suggeriscono Cuomo et al. (2022) e Ng et al. (2021), un approccio integrato è essenziale per una comprensione completa e significativa dell'IA.

Adattabilità a diversi contesti e pubblici: L'implementazione del framework dovrebbe essere flessibile e adattabile a diversi contesti educativi, livelli di istruzione e gruppi target. Diverse popolazioni possono richiedere diverse enfasi all'interno del framework; ad esempio, studenti più giovani potrebbero beneficiare di un maggiore focus sugli aspetti conoscitivi di base, mentre studenti universitari o professionisti potrebbero richiedere una maggiore attenzione alle dimensioni operative e critiche.

Apprendimento esperienziale e basato su progetti: L'AI literacy dovrebbe essere promossa attraverso approcci di apprendimento esperienziale e basato su progetti, che consentano agli studenti di impegnarsi attivamente con le tecnologie di IA in contesti reali o simulati. Come sottolineano Long et al. (2023) e Van Brummelen et al. (2021), l'apprendimento pratico è particolarmente efficace per lo sviluppo di competenze di AI literacy.

Interdisciplinarietà e contestualizzazione: L'educazione all'IA dovrebbe essere interdisciplinare, attingendo da campi come l'informatica, la filosofia, la sociologia, l'etica e la psicologia per fornire una comprensione ricca e contestualizzata dell'IA. Come suggerisce Yi (2021), l'AI literacy non

riguarda solo competenze tecniche ma anche la comprensione dell'IA nel suo contesto sociale, culturale e storico.

Attenzione all'accessibilità e all'inclusione: L'implementazione del framework dovrebbe essere attenta alle questioni di accessibilità e inclusione, assicurando che l'educazione all'IA sia disponibile e rilevante per diverse popolazioni, indipendentemente dal background socioeconomico, dall'identità culturale o dalle abilità. Come sottolineano Druga et al. (2019) e Long et al. (2023), è importante evitare di riprodurre o amplificare disuguaglianze esistenti attraverso l'educazione all'IA.

Ciascuna dimensione del framework dell'AI literacy può beneficiare di approcci pedagogici specifici, progettati per sviluppare efficacemente le competenze associate:

Dimensione conoscitiva:

- Lezioni strutturate sui concetti fondamentali dell'IA
- Dimostrazioni di diverse tecnologie e applicazioni di IA
- Discussioni e dibattiti su casi di studio di IA
- Visualizzazioni e simulazioni per illustrare principi complessi

Dimensione operativa:

- Laboratori pratici con strumenti e applicazioni di IA
- Progetti di gruppo che richiedono l'uso o lo sviluppo di soluzioni basate sull'IA
- Compiti autentici che richiedono la collaborazione con sistemi di IA
- Apprendimento basato su problemi che richiede l'applicazione di conoscenze dell'IA

Dimensione critica:

- Analisi critica di prodotti e servizi basati sull'IA
- Discussioni sui bias algoritmici e sulle loro implicazioni
- Esercizi di valutazione della qualità e dell'affidabilità dei sistemi di IA
- Studio di casi di fallimenti o limitazioni dell'IA

Dimensione etica:

- Discussioni e dibattiti sui dilemmi etici relativi all'IA
- Analisi di casi di studio su questioni etiche come la privacy, l'autonomia e l'equità
- Giochi di ruolo che richiedono decisioni etiche in scenari relativi all'IA
- Progetti di impatto sociale che esplorano come l'IA può essere utilizzata per affrontare sfide sociali

La valutazione efficace dell'AI literacy richiede un approccio multidimensionale che allinei le strategie di valutazione con le diverse dimensioni del framework. Alcune possibili strategie di valutazione includono:

- **Test di conoscenza** per valutare la comprensione concettuale dell'IA (dimensione conoscitiva)
- **Compiti pratici e valutazioni basate sulle prestazioni** per valutare le competenze operative
- **Analisi critica e progetti di ricerca** per valutare le capacità di pensiero critico e valutazione
- **Riflessioni scritte e discussioni** per valutare la comprensione etica e la consapevolezza sociale

Laupichler et al. (2022), Wang et al. (2022) e Pinski et al. (2023) hanno sviluppato strumenti specifici per la valutazione dell'AI literacy, che possono essere adattati e integrati nelle strategie di valutazione basate sul nostro framework. È importante che la valutazione sia allineata con gli obiettivi educativi, autentica, e fornisca opportunità per il feedback e la riflessione.

L'implementazione del framework dell'AI literacy dovrebbe essere sensibile a vari fattori contestuali che possono influenzare l'efficacia degli approcci educativi:

Età e livello di istruzione: Gli approcci all'AI literacy dovrebbero essere adattati al livello di sviluppo e alle esperienze educative precedenti degli studenti. Ad esempio, l'educazione all'IA per bambini della scuola primaria potrebbe concentrarsi su concetti di base e esperienze ludiche, mentre l'educazione per studenti universitari potrebbe includere analisi più approfondite e progetti più complessi.

Background disciplinare: Gli studenti provenienti da diversi background disciplinari possono richiedere diversi punti di ingresso all'AI literacy. Ad esempio, studenti con background in informatica potrebbero beneficiare di una maggiore enfasi sulle dimensioni critiche ed etiche, mentre studenti con background in scienze umane potrebbero richiedere un'introduzione più graduale agli aspetti tecnici dell'IA.

Risorse disponibili: L'implementazione pratica del framework sarà influenzata dalle risorse disponibili, inclusi hardware, software, materiali didattici e tempo. In contesti con risorse limitate, potrebbero essere necessari approcci più creativi o semplificati per promuovere l'AI literacy.

- **Contesto culturale e sociale:** L'educazione all'IA dovrebbe essere sensibile al contesto culturale e sociale in cui avviene, considerando norme, valori e pratiche locali. Come suggerisce Yi (2021), l'AI literacy ha una dimensione culturale significativa che dovrebbe essere riconosciuta e integrata negli approcci educativi.

1.3.6 Sfide e prospettive future

L'implementazione efficace del framework dell'AI literacy proposto in questo lavoro affronta diverse sfide significative, che dovranno essere affrontate per realizzare pienamente il potenziale di questo approccio. Allo stesso tempo, emergono prospettive promettenti per lo sviluppo futuro e l'evoluzione dell'AI literacy. In questa sezione finale, esploriamo queste sfide e prospettive, delineando direzioni per la ricerca e la pratica future.

L'attuazione pratica del framework dell'AI literacy deve affrontare diverse sfide:

Rapida evoluzione tecnologica: L'IA è un campo in rapida evoluzione, con nuove tecnologie, applicazioni e implicazioni che emergono costantemente. Mantenere il framework e i relativi materiali educativi aggiornati con questi sviluppi rappresenta una sfida significativa. Come

suggeriscono Ng et al. (2021) e Long & Magerko (2020), l'AI literacy deve essere concettualizzata come un processo continuo piuttosto che come un endpoint fisso.

Formazione degli educatori: Molti educatori hanno una limitata familiarità con l'IA e possono sentirsi impreparati a integrare l'AI literacy nei loro contesti di insegnamento. La formazione efficace degli educatori è quindi una componente critica per l'implementazione del framework. Come sottolineano Su et al. (2022) e Steinbauer et al. (2021), è necessario sviluppare approcci comprensivi alla formazione degli educatori nell'AI literacy.

Barriere all'accesso: Non tutti gli studenti hanno uguale accesso alle tecnologie, risorse e opportunità necessarie per sviluppare l'AI literacy. Queste disuguaglianze possono amplificare divisioni esistenti ed escludere popolazioni già marginali da opportunità formative cruciale. Affrontare queste barriere all'accesso richiede un impegno consapevole per l'equità e l'inclusione nell'educazione all'IA.

Complessità concettuale: Alcuni concetti di IA possono essere concettualmente complessi e difficili da comunicare a un pubblico non specializzato, in particolare bambini o individui con background non tecnici. Come suggeriscono Druga et al. (2019) e Yang (2022), è importante sviluppare approcci pedagogici che rendano questi concetti accessibili senza sacrificare l'accuratezza o la profondità.

Tensioni etiche: La dimensione etica dell'AI literacy può coinvolgere tensioni tra diversi valori, principi e prospettive. Navigare queste tensioni in modo rispettoso e produttivo rappresenta una sfida pedagogica significativa. Come sottolineano Hermann (2022) e Shih et al. (2021), l'educazione all'etica dell'IA dovrebbe promuovere il dibattito informato piuttosto che imporre una particolare visione morale.

Nonostante queste sfide, emergono diverse prospettive promettenti per il futuro sviluppo dell'AI literacy:

Integrazione curricolare: Una direzione promettente è l'integrazione più sistematica dell'AI literacy nei curricula educativi a tutti i livelli, dalla scuola primaria all'istruzione superiore e all'educazione degli adulti. Come suggeriscono Kong et al. (2021) e Su et al. (2022), l'AI literacy può essere integrata sia come componente standalone sia come elemento trasversale in varie discipline.

Approcci interdisciplinari: Il futuro dell'AI literacy probabilmente vedrà un'enfasi crescente su approcci interdisciplinari che connettono l'IA con altre aree di conoscenza e pratica. Questo può includere ponti tra l'IA e le scienze umane, le arti, le scienze sociali e altre discipline STEM, offrendo percorsi diversi verso l'AI literacy che risuonano con diversi interessi e background.

Co-creazione con gli studenti: Approcci emergenti all'AI literacy enfatizzano il ruolo attivo degli studenti come co-creatori piuttosto che come consumatori passivi di conoscenza. Come suggeriscono Long et al. (2023) e Druga et al. (2022), coinvolgere gli studenti nella progettazione e valutazione di esperienze di apprendimento dell'IA può portare ad approcci più rilevanti e coinvolgenti.

Apprendimento informale e community-based: Mentre gran parte dell'attenzione si è concentrata sull'educazione formale, c'è un crescente riconoscimento dell'importanza dell'apprendimento informale e community-based nell'AI literacy. Iniziative come musei, biblioteche, club e attività familiari possono giocare un ruolo significativo nel rendere l'AI literacy più accessibile e rilevante per diverse popolazioni.

Sviluppo di strumenti e risorse innovative: Continuare a sviluppare strumenti e risorse innovative specificamente progettati per promuovere l'AI literacy rappresenta un'area promettente per la ricerca e la pratica future. Questi possono includere simulazioni interattive, piattaforme di apprendimento adattivo, giochi educativi e ambienti di apprendimento immersivi che rendono l'AI più tangibile e accessibile.

Ricerca empirica: C'è un bisogno continuo di ricerca empirica robusta sull'efficacia di diversi approcci all'AI literacy, su come le diverse dimensioni dell'AI literacy interagiscono e si sviluppano nel tempo, e su come l'AI literacy influenza atteggiamenti, comportamenti e risultati in vari contesti. Questa ricerca può fornire una base più solida per lo sviluppo e il raffinamento del framework e delle pratiche educative associate.

Il framework quadridimensionale per l'AI literacy presentato in questo capitolo offre un approccio comprensivo e strutturato alla concettualizzazione e alla promozione dell'AI literacy in un'era caratterizzata dalla crescente pervasività e influenza dell'intelligenza artificiale. Costruendo sulla base di modelli precedenti di competenza digitale e sui contributi emergenti alla concettualizzazione dell'AI literacy, il framework identifica quattro dimensioni fondamentali – conoscitiva, operativa, critica ed etica – che insieme costituiscono un approccio olistico all'educazione all'IA.

La dimensione conoscitiva fornisce la base concettuale necessaria per comprendere cosa sia l'IA, come funzioni e quali siano le sue capacità e limitazioni. La dimensione operativa si concentra sulle abilità pratiche necessarie per utilizzare, applicare e interagire con le tecnologie di IA in modo efficace ed efficiente. La dimensione critica enfatizza la capacità di analizzare, valutare e interpretare criticamente le tecnologie di IA, i loro output e il loro impatto sulla società. La dimensione etica, infine, riconosce l'importanza di considerazioni etiche, sociali e valoriali nell'interazione con l'IA, promuovendo un uso responsabile e consapevole di queste tecnologie.

Questo framework rappresenta un contributo significativo al campo emergente dell'AI literacy, offrendo un quadro concettuale coerente che può guidare lo sviluppo di curricula educativi, materiali didattici, approcci pedagogici e strumenti di valutazione. Il framework è sufficientemente flessibile da poter essere adattato a diversi contesti educativi, livelli di istruzione e gruppi target, pur mantenendo una coerenza concettuale che facilita la comunicazione e la collaborazione tra ricercatori, educatori e policymaker interessati alla promozione dell'AI literacy.

Allo stesso tempo, il framework riconosce le sfide significative associate all'implementazione pratica dell'AI literacy, incluse la rapida evoluzione tecnologica, la necessità di formazione degli educatori, le barriere all'accesso, la complessità concettuale e le tensioni etiche. Affrontare queste sfide richiederà un impegno continuo da parte di diverse parti interessate e una volontà di adattare e raffinare gli approcci all'AI literacy in risposta alle emergenti necessità, opportunità e intuizioni.

Guardando al futuro, il framework dell'AI literacy proposto in questo lavoro offre un punto di partenza promettente per ulteriori ricerche e sviluppi in questo campo. L'integrazione curricolare, gli approcci interdisciplinari, la co-creazione con gli studenti, l'apprendimento informale e community-based, lo sviluppo di strumenti e risorse innovative, e la ricerca empirica robusta rappresentano tutte direzioni promettenti per l'evoluzione dell'AI literacy.

In ultima analisi, il valore dell'AI literacy risiede nella sua capacità di empowerment. In un mondo sempre più plasmato dall'intelligenza artificiale, l'AI literacy fornisce agli individui le conoscenze, le abilità e le disposizioni necessarie per navigare, comprendere, utilizzare e contribuire criticamente a questo paesaggio tecnologico in evoluzione. Il framework quadridimensionale presentato in questo

capitolo mira a contribuire a questo obiettivo, promuovendo un approccio all'AI literacy che è comprensivo, critico, etico e orientato all'azione.

CAPITOLO 2: Insegnare l'AI literacy: dai contenuti alle tecniche didattiche

La straordinaria espansione dell'Intelligenza Artificiale (IA) nel tessuto quotidiano delle nostre vite rappresenta una trasformazione profonda, paragonabile per portata ad altre rivoluzioni tecnologiche del passato, ma, soprattutto, è caratterizzata da una velocità di diffusione e pervasività senza precedenti. Si tratta, come si è visto, di una trasformazione che ha investito ambiti produttivi, economici o della comunicazione e, di conseguenza, anche il mondo dell'educazione, ponendo sfide inedite e richiedendo lo sviluppo di nuove competenze per muoversi con consapevolezza in un presente e un futuro sempre più mediati da sistemi intelligenti. È in tale contesto, che è emersa con forza quell'esigenza di AI literacy – documentata dalla letteratura scientifica che abbiamo esaminato nel capitolo precedente – intesa non come mero addestramento tecnico, bensì come acquisizione di una competenza complessa adatta ad abilitare individui e comunità a comprendere, utilizzare e valutare criticamente le tecnologie di cui ora disponiamo. L'impegno è notevole e ricade inevitabilmente sui sistemi formativi, chiamati a integrare questa nuova istanza educativa nei propri curricula, a partire dai primi cicli scolastici.

Il presente capitolo si colloca esattamente in questa prospettiva, proponendosi di delineare un percorso organico e pedagogicamente fondato per l'insegnamento dell'AI literacy. L'intento è quello di muovere da un quadro concettuale robusto – il framework quadridimensionale precedentemente elaborato (Cuomo et al., 2022) – per tradurlo in una proposta curricolare flessibile e adattabile, e successivamente esplorare in profondità lo sviluppo di strategie didattiche specifiche, capaci di promuovere attivamente le diverse sfaccettature dell'AI literacy. L'approccio adottato intende superare la dicotomia tra sapere teorico e pratica educativa, proponendo un impianto teorico-pratico radicato nei principi dell'instructional design e nelle teorie dell'apprendimento più accreditate in ambito educativo, quali il costruttivismo, l'apprendimento situato e la pedagogia critica. Si tratta di fornire non solo una "cassetta degli attrezzi" metodologica, ma anche una cornice di senso che ne giustifichi l'utilizzo e ne orienti l'applicazione verso finalità educative elevate, quali lo sviluppo del pensiero autonomo, della responsabilità civica e della capacità di agire in modo informato e consapevole nel mondo contemporaneo.

La trattazione che segue ha l'obiettivo non solo di descrivere il "cosa" insegnare (i contenuti del curriculum), ma soprattutto il "come" (le strategie didattiche) e il "perché" (le ragioni pedagogiche e le finalità formative), in linea con le esigenze di approfondimento e riflessione critica proprie di una tesi di dottorato. Le riflessioni e le proposte operative qui presentate si basano sull'analisi e sull'espansione sistematica dei materiali di ricerca e delle esperienze formative precedentemente condotte (Cuomo et al., 2022; Ranieri, Cuomo, Biagini, 2024), con l'ambizione di offrire un contributo utile alla progettazione di percorsi di AI literacy efficaci, significativi e realmente implementabili nei diversi contesti educativi, dalla scuola dell'infanzia alla formazione universitaria e continua.

2.1 Dal framework ad un curriculum per l'AI literacy: progettare per la comprensione profonda

La costruzione di un curriculum efficace per l'AI literacy non può prescindere da un solido impianto teorico che ne orienti le scelte contenutistiche, metodologiche e valutative. Il framework (Cuomo et al., 2022) rappresenta tale fondamento concettuale. Esso offre una mappa per navigare la complessità dell'AI literacy, identificandone le componenti essenziali (conoscenze, abilità, atteggiamenti) e le loro

interrelazioni dinamiche. Tuttavia, un framework, per quanto robusto, non è ancora un curriculum. Il passaggio dalla mappa concettuale alla progettazione di un percorso di apprendimento concreto richiede un processo intenzionale di traduzione pedagogica e didattica, guidato dai principi dell'instructional design e informato da una chiara visione dei processi di apprendimento che si intendono promuovere. Si tratta di trasformare le dimensioni astratte del framework in moduli formativi coerenti, obiettivi di apprendimento specifici e attività didattiche significative.

2.1.1 Principi metodologici per la trasposizione del framework in curriculum

Adottare un approccio sistematico alla progettazione curricolare, come quello suggerito da modelli di instructional design quali ADDIE (Analysis, Design, Development, Implementation, Evaluation), permette di strutturare questo processo in modo rigoroso e riflessivo. La fase di Analysis (analisi dei bisogni, del contesto, del target) ha già portato all'identificazione della necessità di un'AI literacy multidimensionale e alla definizione del framework stesso come risposta a tale bisogno. La fase cruciale di Design, che qui approfondiamo, implica la definizione precisa degli obiettivi di apprendimento per ciascuna dimensione (cosa dovrebbero sapere, saper fare e come dovrebbero atteggiarsi i discenti al termine del percorso?), la selezione e sequenziazione dei contenuti (quali concetti chiave, quali esempi, quali strumenti?), la scelta delle strategie didattiche più appropriate (come facilitare l'apprendimento attivo e significativo?) e la pianificazione delle modalità di valutazione (come verificare il raggiungimento degli obiettivi in modo autentico?).

La trasformazione di un modello teorico in una proposta curricolare richiede dunque l'applicazione di principi metodologici rigorosi che garantiscano coerenza e validità al processo. Nel passaggio dal framework quadridimensionale dell'AI literacy a un curriculum strutturato, sono stati adottati alcuni principi fondamentali che hanno guidato l'intero processo di progettazione.

Il primo principio metodologico punta sull'importanza di mantenere un'aderenza rigorosa alla struttura quadridimensionale del framework originario. Le quattro dimensioni identificate da Cuomo et al. (2022) costituiscono, infatti, l'ossatura concettuale intorno alla quale organizzare i contenuti curricolari. Tale principio ha permesso di preservare l'integrità concettuale del framework, evitando riduzioni o distorsioni che avrebbero potuto comprometterne la validità. Come hanno sottolineato Ng et al. (2021), la multidimensionalità dell'AI literacy non è semplicemente una scelta organizzativa, ma riflette la complessità intrinseca del fenomeno dell'intelligenza artificiale e delle competenze necessarie per comprenderla e utilizzarla consapevolmente.

Il secondo principio metodologico concerne l'adattamento dei contenuti ai diversi livelli scolastici. Nel progettare un curriculum di AI literacy che possa essere adottato dalla scuola primaria fino alla secondaria di secondo grado, è stato necessario calibrare i contenuti in base alle capacità cognitive e alle conoscenze pregresse degli studenti. Questo principio si allinea con quanto evidenziato nel documento UNESCO (2022) "K-12 AI curricula: a mapping of government-endorsed AI curricula", che propone una progressione di contenuti in base allo sviluppo cognitivo degli studenti. Tale approccio non implica una semplificazione delle dimensioni del framework, ma piuttosto una loro declinazione in termini di complessità concettuale e operativa appropriata alle diverse età dei discenti.

Il terzo principio metodologico riguarda l'interconnessione tra le dimensioni. Sebbene il framework presenti quattro dimensioni distinte, nella realtà dell'apprendimento tali dimensioni sono intrinsecamente correlate. Nel progettare il curriculum, si è cercato di evidenziare e valorizzare queste connessioni, creando percorsi di apprendimento che integrassero le diverse dimensioni. Questo approccio si fonda sulla convinzione che l'AI literacy non possa essere sviluppata attraverso compartimenti stagni, ma richieda una visione olistica che permetta agli studenti di cogliere le

relazioni tra conoscenze, abilità operative, capacità critiche e consapevolezza etica (Long & Magerko, 2020).

Il quarto principio metodologico concerne l'ancoraggio delle attività didattiche a paradigmi pedagogici consolidati. Nel trasformare il framework in curriculum, si è fatto riferimento a teorie dell'apprendimento basate sul costruttivismo (Fosnot, 2013), sul socio-costruttivismo (Vygotsky, 1978) e sull'apprendimento situato (Lave & Wenger, 1991). Questa prospettiva vede l'apprendimento non come un travaso di conoscenze da un docente esperto a un discente passivo, ma come un processo attivo, sociale e situato di costruzione di significato. I discenti arrivano con le proprie preconcezioni e teorie implicite sull'IA (spesso mutate dai media o dall'esperienza quotidiana) perciò il curriculum deve offrire esperienze che permettano loro di esplicitare, confrontare, mettere in discussione e rielaborare tali conoscenze attraverso l'interazione con i pari, con il docente-facilitatore e con artefatti culturali e tecnologici. Ciò implica la necessità di progettare ambienti di apprendimento ricchi di opportunità per l'esplorazione, la sperimentazione, il dialogo argomentato e la riflessione condivisa.

L'apprendimento situato rafforza questa visione, sottolineando che l'apprendimento è più efficace quando è ancorato a contesti e pratiche autentiche. Questo per ribadire che le competenze di AI literacy non vanno acquisite in astratto, ma sviluppate attraverso l'analisi di scenari reali in cui l'IA opera (dalle raccomandazioni di Netflix alla diagnostica medica), la manipolazione diretta (o simulata) di strumenti di IA (come Teachable Machine o piattaforme di coding) e la risoluzione di problemi concreti che richiedono l'applicazione di tali competenze (es. valutare l'affidabilità di un'informazione generata da un chatbot o progettare un semplice algoritmo per un compito specifico).

La pedagogia critica (Giroux, 2011; Kellner & Share, 2007) fornisce una lente essenziale per orientare il curriculum verso finalità emancipatrici. L'AI literacy, in questa chiave, non si limita alla comprensione tecnica o all'uso funzionale, ma include la capacità di decostruire criticamente le narrazioni dominanti sull'IA, di riconoscere le dinamiche di potere e le disuguaglianze che essa può generare o amplificare (come i bias algoritmici), e di sviluppare una coscienza etica e una capacità di azione civica per promuovere un uso dell'IA più giusto, equo e al servizio del benessere collettivo. Il curriculum, quindi, deve integrare momenti di analisi critica dei media, di discussione su temi eticamente sensibili e di riflessione sul ruolo dell'IA nella società e nella democrazia.

Infine, il quinto principio metodologico riguarda l'apertura e la flessibilità del curriculum. Data la rapida evoluzione dell'intelligenza artificiale e delle sue applicazioni, un curriculum di AI literacy deve necessariamente mantenere un carattere aperto e flessibile, capace di adattarsi alle trasformazioni tecnologiche e sociali. Questo principio ha portato a concepire il curriculum non come un percorso rigidamente definito, ma come una struttura modulare che può essere aggiornata e adattata in funzione dei contesti specifici di implementazione e dell'evoluzione tecnologica.

Questi cinque principi metodologici hanno guidato il processo di trasformazione del framework teorico in un curriculum strutturato, garantendo coerenza concettuale, adeguatezza pedagogica e flessibilità applicativa.

2.1.2 Progressione delle competenze: verso una padronanza multilivello

Un curriculum efficace non può presentare tutte le competenze allo stesso livello, ma deve prevedere una progressione graduale che rispetti lo sviluppo cognitivo dei discenti e costruisca la complessità in modo incrementale. Ispirandoci alla tassonomia di Bloom rivista (Anderson & Krathwohl, 2001), che articola l'apprendimento su sei livelli cognitivi (Ricordare, Comprendere, Applicare, Analizzare, Valutare, Creare), possiamo delineare una traiettoria evolutiva per le competenze chiave dell'AI literacy, fornendo una bussola per differenziare obiettivi e attività.

La progressione verticale del curriculum di AI literacy è stata organizzata tenendo conto delle fasi di sviluppo cognitivo degli studenti, dalla scuola primaria alla secondaria di secondo grado. Questa scelta si allinea con gli studi di Piaget (1964) sullo sviluppo cognitivo e con l'approccio di Vygotskij (1978) alla zona di sviluppo prossimale, entrambi fondamentali per comprendere come gli studenti costruiscano progressivamente concetti complessi. Nel contesto specifico dell'AI literacy, la progressione verticale si traduce in un incremento graduale della complessità concettuale, della profondità analitica nonché dell'autonomia operativa.

Per le diverse dimensioni del framework, possiamo articolare la progressione delle competenze come segue:

Dimensione conoscitiva (ricordare, comprendere, applicare):

Ricordare: Identificare esempi comuni di IA (assistenti vocali, filtri social), nominare figure chiave (Turing), richiamare eventi storici salienti (il seminario di Dartmouth). Livello Primaria: Riconoscere oggetti "smart".

Comprendere: Spiegare con parole proprie cos'è un algoritmo, la differenza tra IA debole e forte, le caratteristiche distintive dell'IA rispetto all'intelligenza umana, il ruolo dei dati. Livello Secondaria I: Spiegare perché un navigatore usa l'IA.

Applicare: Usare le conoscenze acquisite per identificare l'uso dell'IA in un nuovo dispositivo, spiegare un concetto base a un compagno, classificare diversi tipi di IA. Livello Secondaria II: Applicare la conoscenza storica per contestualizzare un dibattito attuale sull'IA.

Dimensione operativa (comprendere, applicare, analizzare):

Comprendere: Spiegare il significato di "dataset", "addestramento", "classificazione". Descrivere i passaggi logici di un semplice algoritmo.

Applicare: Seguire o creare istruzioni per attività di coding unplugged, usare un'applicazione IA per uno scopo specifico (es. traduzione), fornire dati di addestramento a un modello semplice (es. Teachable Machine). Livello Primaria/Secondaria I: Usare Scratch per creare una semplice animazione basata su regole.

Analizzare: Scomporre un problema per identificarne i passaggi risolutivi (pensiero computazionale), analizzare come un dataset specifico influenza l'output di un modello di classificazione, confrontare l'efficacia di diversi algoritmi semplici per lo stesso compito. Livello Secondaria II: Analizzare il funzionamento di un sistema di raccomandazione.

Dimensione critica (analizzare, valutare, creare):

Analizzare: Identificare stereotipi o bias in un testo o immagine generati dall'IA, riconoscere le tecniche di persuasione usate da un chatbot, analizzare le fonti di un'informazione fornita dall'IA.

Valutare: Giudicare l'affidabilità di un output dell'IA sulla base di criteri definiti, valutare l'importanza della trasparenza di un algoritmo in un contesto specifico (es. medico vs. gioco), confrontare criticamente diverse prospettive sull'impatto sociale dell'IA. Livello Università: Valutare la robustezza metodologica di uno studio sull'efficacia di un sistema IA in educazione.

Creare: Progettare un'attività didattica che promuova il pensiero critico sull'IA, modificare un prompt per ottenere risultati più accurati o meno distorti da un LLM, proporre soluzioni per mitigare un bias identificato in un sistema.

Dimensione etica (comprendere, applicare, analizzare, valutare):

Comprendere: Descrivere i principali principi etici dell'IA (privacy, equità, ecc.), spiegare perché un certo uso dell'IA pone un dilemma morale.

Applicare: Utilizzare i principi etici per analizzare un caso di studio (es. il giocattolo che registra), identificare le possibili conseguenze etiche di un'azione legata all'IA.

Analizzare: Scomporre un dilemma etico complesso nelle sue componenti (valori in conflitto, stakeholder coinvolti, possibili azioni), analizzare come i dati personali vengono raccolti e usati da una specifica applicazione.

Valutare: Giudicare diverse soluzioni a un dilemma etico sulla base di criteri morali argomentati, valutare l'adequatezza delle attuali normative sulla privacy rispetto alle sfide dell'IA, formulare un giudizio ponderato sull'opportunità di delegare certe decisioni all'IA.

Questa progressione, che si sviluppa sia all'interno di ciascuna dimensione sia trasversalmente attraverso i livelli cognitivi, suggerisce l'adozione di un curriculum "a spirale" (Bruner, 1960). I concetti e le competenze fondamentali (es. algoritmo, dato, bias, privacy) vengono introdotti precocemente in forme semplici e concrete, per poi essere ripresi e approfonditi a livelli di scolarità successivi, con crescente complessità concettuale e richiesta di prestazioni cognitive più elevate. Ad esempio, il concetto di bias può essere introdotto alla primaria attraverso storie che mostrano stereotipi, analizzato alla secondaria I grado tramite esperimenti con Teachable Machine, e discusso alla secondaria II grado in relazione a complessi sistemi decisionali e alle loro implicazioni sociali e legali.

Per rappresentare in modo chiaro questa strutturazione curricolare, è stata elaborata una matrice che incrocia le dimensioni dell'AI literacy con i livelli scolastici (Tabella 2.1).

Tabella 2.1. Matrice della progressione curricolare per dimensioni e livelli scolastici

Livello scolastico	Dimensione conoscitiva	Dimensione operativa	Dimensione critica	Dimensione etica
Scuola primaria	Concetti di base dell'AI attraverso metafore e analogie	Interazione guidata con semplici applicazioni di AI	Osservazione delle differenze tra decisioni umane e automatizzate	Introduzione a semplici questioni etiche attraverso narrazioni
Scuola secondaria di I grado	Comprensione dei principi di funzionamento dell'AI e del machine learning	Utilizzo consapevole di applicazioni di AI e sperimentazione di semplici processi di training	Analisi guidata delle potenzialità e dei limiti dell'AI	Riflessione sulle implicazioni etiche dell'AI nella vita quotidiana

Scuola secondaria di II grado	Comprensione approfondita dei diversi approcci all'AI e delle loro applicazioni	Utilizzo critico degli strumenti di AI e progettazione di semplici sistemi	Valutazione autonoma dell'impatto dell'AI in vari contesti sociali	Analisi critica delle questioni etiche legate all'AI e sviluppo di framework valutativi
--	---	--	--	---

La matrice evidenzia come, per ciascuna dimensione, i contenuti e le competenze evolvano in termini di complessità e profondità attraverso i diversi livelli scolastici. Questa progressione non è semplicemente lineare, ma riflette una spirale di apprendimento in cui concetti e competenze vengono ripresi e approfonditi in contesti progressivamente più complessi.

La sequenzialità orizzontale, invece, si manifesta nelle connessioni tra le diverse dimensioni all'interno di ciascun livello scolastico. Ad esempio, nella scuola primaria, l'introduzione ai concetti di base dell'AI (dimensione conoscitiva) è strettamente collegata all'interazione guidata con semplici applicazioni (dimensione operativa). Entrambe queste dimensioni preparano il terreno per l'osservazione delle differenze tra decisioni umane e automatizzate (dimensione critica) e per l'introduzione a semplici questioni etiche (dimensione etica). Tali connessioni sono esplicitate nei percorsi didattici e nelle attività proposte, al fine di promuovere una visione integrata dell'AI literacy.

Come si può vedere, la progressione e la sequenzialità del curriculum di AI literacy non sono concepite come strutture rigide, ma come linee guida flessibili che possono essere riformulate in base ai contesti specifici di implementazione, alle conoscenze pregresse degli studenti e agli obiettivi educativi specifici. È una flessibilità necessaria per garantire che il curriculum rimanga rilevante e significativo anche nel campo dell'IA che sappiamo in rapida evoluzione.

2.1.3 Selezione e organizzazione dei contenuti per le quattro dimensioni

La selezione e l'organizzazione dei contenuti per ciascuna delle quattro dimensioni del framework rappresenta un passaggio cruciale nella trasformazione del modello teorico in un curriculum operativo. Questo processo richiede una riflessione approfondita sui concetti fondamentali dell'intelligenza artificiale, sulle competenze essenziali da sviluppare e sulle modalità più efficaci per promuovere un apprendimento significativo.

Il curriculum si articola idealmente in quattro moduli principali, che rispecchiano le dimensioni del framework, ma che dovrebbero essere concepiti come interconnessi e permeabili, piuttosto che come compartimenti stagni.

Modulo 1: dimensione conoscitiva – comprendere l'intelligenza artificiale

Per la **dimensione conoscitiva** (si veda 1.3.3 per la sua descrizione approfondita), la selezione dei contenuti si è basata sull'identificazione dei concetti fondamentali necessari per comprendere cosa sia l'intelligenza artificiale, come funzioni e quali siano le sue principali applicazioni.

Questo modulo può beneficiare dell'apporto delle scienze cognitive per esplorare i modelli della mente umana e confrontarli con quelli dell'IA nonché della linguistica per analizzare come l'IA processa e genera il linguaggio naturale. Attività interdisciplinari potrebbero includere la creazione di linee del tempo comparative tra sviluppo tecnologico e pensiero filosofico sull'intelligenza, o l'analisi linguistica di testi prodotti da IA.

La selezione dei contenuti per questa dimensione ha tenuto conto della necessità di bilanciare rigore concettuale e accessibilità, adattando i concetti alla capacità di comprensione degli studenti dei diversi livelli scolastici. Per la scuola primaria, ad esempio, i concetti sono presentati attraverso metafore, analogie e narrazioni, mentre per la scuola secondaria si adotta un approccio progressivamente più formale e strutturato.

Modulo 2: dimensione operativa – interagire con e attraverso l'Intelligenza Artificiale

Per la **dimensione operativa** (si veda 1.3.3 per la sua descrizione approfondita), la selezione dei contenuti si è concentrata sulle competenze pratiche necessarie per interagire con l'intelligenza artificiale, utilizzarla in modo efficace e, a livelli più avanzati, partecipare al suo sviluppo.

Accanto a matematica, statistica, scienze, tecnologia e lingue, questo modulo può integrare l'educazione fisica (per le attività di Body Coding), l'arte (per la Pixel Art o l'uso di IA generative per creare immagini), la musica (per analizzare sistemi di raccomandazione musicale o IA che compongono musica). Un progetto interdisciplinare potrebbe riguardare la creazione di un semplice "robot" (anche unplugged) che esegue compiti legati a una disciplina specifica (es. classificare foglie in scienze, seguire percorsi geometrici in matematica).

L'organizzazione di questi contenuti ha seguito un principio di progressione dalla fruizione passiva alla produzione attiva, dalla semplice interazione con applicazioni esistenti alla progettazione e realizzazione di semplici sistemi di AI. Questa progressione riflette il passaggio da un approccio basato sull'utente a uno basato sul creatore, in linea con l'evoluzione delle competenze digitali proposte da vari framework internazionali (Ferrari et al., 2014).

Modulo 3: dimensione critica – valutare l'intelligenza artificiale e i suoi impatti

Per la **dimensione critica** (si veda 1.3.3 per la sua descrizione approfondita), la selezione dei contenuti si è focalizzata sullo sviluppo di capacità analitiche e valutative necessarie per comprendere le potenzialità e i limiti dell'intelligenza artificiale, valutarne l'affidabilità e comprenderne l'impatto sociale.

Oltre a educazione ai media, sociologia, educazione civica e psicologia, questo modulo si presta a integrazioni con la storia (analizzando come le tecnologie passate abbiano generato bias o impatti sociali inattesi), l'economia (valutando l'impatto dell'automazione sul lavoro), e il diritto (discutendo casi di discriminazione algoritmica). Attività interdisciplinari aggiuntive potrebbero consistere nell'analizzare criticamente la rappresentazione dell'IA nel cinema o nella letteratura, o nel condurre piccole inchieste sull'uso dei dati in diverse piattaforme online.

L'organizzazione di questi contenuti ha seguito un principio di progressione dalla semplice osservazione e descrizione all'analisi critica e alla valutazione autonoma. Tale progressione è particolarmente importante per lo sviluppo del pensiero critico, che richiede un percorso graduale dalla consapevolezza alla valutazione critica (Facione, 2011).

Modulo 4: Dimensione etica – agire in modo responsabile nell'era dell'Intelligenza Artificiale

Infine, per la **dimensione etica** (si veda 1.3.3 per la sua descrizione approfondita), la selezione dei contenuti si è concentrata sulle questioni etiche legate all'utilizzo dell'intelligenza artificiale, sui principi e i valori che dovrebbero guidarne lo sviluppo e l'applicazione, e sulle implicazioni morali delle decisioni automatizzate.

Qui l'interdisciplinarietà è massima. Oltre a filosofia morale, diritto, economia, educazione civica e scienze ambientali, possono contribuire la letteratura (esplorando dilemmi morali in opere di

fantascienza), la religione (offrendo prospettive etiche diverse), l'arte (come forma di espressione e riflessione su temi etici). Progetti interdisciplinari potrebbero riguardare la stesura di una "carta etica" per l'uso dell'IA a scuola, o l'organizzazione di un dibattito pubblico su un tema etico rilevante per la comunità.

L'organizzazione di questi contenuti ha seguito un principio di progressione dalla semplice sensibilizzazione alle questioni etiche alla capacità di analizzare dilemmi complessi e formulare giudizi etici articolati. Questa progressione riflette lo sviluppo del ragionamento morale, che evolve da una visione basata su regole concrete a una comprensione più astratta e principale dell'etica (Kohlberg, 1984).

La Tabella 2.2 sintetizza i contenuti selezionati per ciascuna dimensione e la loro articolazione nelle diverse aree tematiche.

Tabella 2.2. Contenuti selezionati per le quattro dimensioni dell'AI literacy

Dimensione	Area tematica	Contenuti principali
Conoscitiva	Definizioni e tipi di AI	- Definizioni di intelligenza artificiale (AI debole/forte) - Breve storia dell'AI - Tipologie e paradigmi di AI
	Dati e machine learning	- Dati: raccolta, elaborazione, rappresentazione - Principi di machine learning - Reti neurali e deep learning
	Applicazioni dell'AI	- Computer vision e riconoscimento delle immagini - Elaborazione del linguaggio naturale - Robotica e sistemi autonomi
Operativa	Utilizzo delle applicazioni di AI	- Interfacce di AI e modalità di interazione - Utilizzo di assistenti virtuali e chatbot - Sistemi di raccomandazione e personalizzazione
	Pensiero computazionale e programmazione	- Principi del pensiero computazionale - Algoritmi e flow chart - Programmazione visuale e testuale
	Analisi e gestione dei dati	- Raccolta e organizzazione dei dati - Visualizzazione e interpretazione dei dati - Pulizia e trasformazione dei dati
Critica	Interpretazione e valutazione dell'AI	- Comprensione dei risultati dell'AI - Bias e limitazioni dei sistemi di AI - Valutazione dell'affidabilità dell'AI
	Comprensione dell'impatto sociale dell'AI	- AI e trasformazione del lavoro - AI e disuguaglianze sociali - AI e democrazia
	Alfabetizzazione ai dati e all'informazione	- Valutazione critica delle fonti di informazione - Comprensione dei processi di dataficazione - Infodemia e disinformazione

Etica	Principi etici fondamentali dell'AI	- Trasparenza e spiegabilità - Giustizia e non discriminazione - Privacy e protezione dei dati
	Dilemmi etici nell'applicazione dell'AI	- Sorveglianza e controllo sociale - Decisioni automatizzate in contesti critici - Responsabilità e accountability
	Prospettive etiche e governance dell'AI	- Prospettive deontologiche- Regolamentazione e governance dell'AI - Partecipazione civica e democratica

La selezione e l'organizzazione dei contenuti per ciascuna dimensione non rappresenta un processo concluso, ma piuttosto un punto di partenza per lo sviluppo di percorsi didattici specifici. I contenuti proposti possono e devono essere adattati in base ai contesti specifici di implementazione, alle conoscenze pregresse degli studenti e agli obiettivi educativi particolari.

La flessibilità del curriculum è essenziale. Nuovamente occorre notare che non si tratta di un percorso rigido, ma di una cornice adattabile. Il docente può scegliere quali moduli approfondire maggiormente, quali strategie privilegiare, come inserire i temi dell'AI literacy nel tessuto delle discipline esistenti. L'originalità della proposta, pur allineandosi alle raccomandazioni internazionali come quelle UNESCO (2022), risiede in questa sua strutturazione pedagogicamente robusta, basata su un framework multidimensionale e su principi di apprendimento attivo, critico e situato, e nella sua concreta implementabilità, testimoniata dall'esperienza formativa descritta nel Capitolo 4, che ne valida l'efficacia nel promuovere una comprensione profonda e trasformativa dell'AI literacy in futuri insegnanti, figure chiave per preparare le nuove generazioni alle sfide del futuro.

2.1.4 Connessioni interdisciplinari e integrazione curricolare

La natura intrinsecamente multidisciplinare dell'intelligenza artificiale richiede che un curriculum di AI literacy non sia concepito come un'entità isolata, ma come un sistema aperto che stabilisce connessioni significative con le diverse discipline scolastiche. L'integrazione curricolare rappresenta una sfida cruciale per l'implementazione efficace dell'AI literacy nel sistema educativo, poiché richiede di superare la tradizionale compartimentazione del sapere in favore di un approccio più olistico e interconnesso.

Come sottolineato dall'UNESCO (2022), l'insegnamento dell'AI literacy non dovrebbe configurarsi come l'aggiunta di una nuova materia al curriculum già denso delle scuole, ma piuttosto come un'opportunità di arricchimento e trasformazione delle discipline esistenti attraverso l'integrazione di concetti, metodologie e strumenti legati all'intelligenza artificiale. Questa prospettiva si allinea con l'approccio delle competenze trasversali proposto da vari framework educativi internazionali (OECD, 2018; European Commission, 2019), che enfatizzano l'importanza di sviluppare competenze che attraversano i confini disciplinari tradizionali.

Le connessioni interdisciplinari del curriculum di AI literacy possono essere articolate su tre livelli principali: concettuale, metodologico e applicativo.

A livello concettuale, l'AI literacy stabilisce connessioni significative con diverse discipline. Con la matematica, condivide concetti fondamentali come algoritmi, logica, probabilità e statistica, che sono essenziali per comprendere i meccanismi di funzionamento dell'intelligenza artificiale. Con l'informatica, condivide i principi del pensiero computazionale, della programmazione e della

gestione dei dati. Con le scienze, condivide l'approccio sperimentale, l'analisi dei dati e la modellizzazione. Infine, con le discipline umanistiche, condivide la riflessione critica, l'analisi etica e la comprensione dell'impatto sociale delle tecnologie.

A livello metodologico, l'AI literacy si integra con le metodologie didattiche di varie discipline. Ad esempio, l'approccio basato sull'indagine (*inquiry-based learning*), tipico delle scienze, può essere efficacemente applicato all'esplorazione dell'intelligenza artificiale. Similmente, la discussione critica e il dibattito, metodologie tipiche delle discipline umanistiche, possono essere utilizzati per affrontare le questioni etiche e sociali legate all'AI. Il problem-solving, metodologia trasversale a molte discipline, trova nell'AI un campo di applicazione particolarmente ricco e stimolante.

A livello applicativo, l'AI literacy offre numerose opportunità di integrazione con contenuti disciplinari specifici. Ad esempio, nell'insegnamento delle lingue, si possono esplorare le applicazioni dell'AI per la traduzione automatica o la generazione di testi; nelle scienze, è possibile analizzare come l'AI viene utilizzata per la modellizzazione climatica o la ricerca biomedica; nella geografia, si possono esaminare le applicazioni dell'AI per l'analisi di dati spaziali o la previsione di fenomeni ambientali.

La Tabella 2.3 illustra alcune delle principali connessioni interdisciplinari del curriculum di AI literacy, evidenziando per ciascuna disciplina le connessioni concettuali, metodologiche e applicative.

Tabella 2.3. Connessioni interdisciplinari del curriculum di AI literacy

Disciplina	Connessioni concettuali	Connessioni metodologiche	Connessioni applicative
Matematica	- Algoritmi e procedure - Logica e ragionamento - Probabilità e statistica	- Problem-solving - Modellizzazione - Analisi dei dati	- Sistemi di apprendimento supervisionato - Algoritmi di classificazione - Reti neurali
Informatica	- Pensiero computazionale - Programmazione - Strutture dati	- Sviluppo iterativo - Debugging - Progettazione modulare	- Sviluppo di applicazioni di AI - Analisi e visualizzazione dei dati - Sistemi di raccomandazione
Scienze	- Metodo scientifico - Modelli e teorie - Riconoscimento di pattern	- Inquiry-based learning - Sperimentazione - Analisi dei dati	- Modellizzazione climatica - Ricerca biomedica - Biodiversità e conservazione
Lingue	- Linguaggio e comunicazione - Semantica e sintassi - Narrativa e retorica	- Analisi testuale - Produzione creativa - Dibattito	- Traduzione automatica - Generazione di testi - Analisi del sentiment

Storia e Scienze Sociali	- Evoluzione tecnologica - Impatto sociale - Dilemmi etici	- Analisi delle fonti - Dibattito - Studio di caso	- Analisi di dati storici - Previsione di tendenze sociali - Sistemi di supporto alle decisioni
Arte e Musica	- Creatività e innovazione - Estetica e percezione - Espressione e comunicazione	- Progettazione creativa - Sperimentazione - Analisi critica	- Arte generativa - Composizione musicale assistita - Sistemi di riconoscimento visivo

L'integrazione curricolare dell'AI literacy può essere facilitata attraverso diverse strategie, come l'approccio per progetti interdisciplinari, l'apprendimento basato su problemi reali, e la collaborazione tra docenti di diverse discipline. Queste strategie permettono di superare la frammentazione disciplinare e di promuovere un apprendimento più olistico e significativo.

Un esempio concreto di integrazione curricolare è rappresentato dal progetto "Cambiamenti climatici ed impatto sociale" che integra concetti di AI literacy con matematica, geografia, scienze e italiano. In questo progetto, gli studenti esplorano come l'intelligenza artificiale può essere utilizzata per analizzare dati climatici, prevedere tendenze future e supportare decisioni in ambito ambientale. Attraverso questo approccio interdisciplinare, gli studenti non solo sviluppano competenze specifiche di AI literacy, ma approfondiscono anche la loro comprensione dei cambiamenti climatici e delle loro implicazioni sociali.

L'integrazione curricolare dell'AI literacy presenta inoltre sfide significative, come la necessità di formazione degli insegnanti, la riorganizzazione dei tempi e degli spazi di apprendimento, e la revisione delle modalità di valutazione. Queste sfide richiedono un impegno sistematico a livello di politiche educative, formazione docenti e ripensamento delle pratiche didattiche.

In conclusione, le connessioni interdisciplinari e l'integrazione curricolare rappresentano un aspetto fondamentale del curriculum di AI literacy, che, ricordiamo ancora, non può essere concepito come un'entità isolata, ma come un sistema aperto che stabilisce relazioni significative con le diverse discipline scolastiche. Questa prospettiva interdisciplinare non solo arricchisce il curriculum di AI literacy, ma contribuisce anche a rinnovare l'insegnamento delle discipline tradizionali, promuovendo un approccio più consona con la richiesta di sapere del mondo contemporaneo.

2.2 Sviluppo di strategie didattiche per l'AI literacy: metodologie per un apprendimento attivo e consapevole

Se il curriculum definisce la mappa del percorso formativo, le strategie didattiche rappresentano i veicoli e le modalità con cui tale percorso viene effettivamente compiuto. Per promuovere un'AI literacy profonda e significativa, che vada oltre la semplice acquisizione di nozioni, è indispensabile adottare metodologie attive, partecipative e riflessive, in grado di coinvolgere gli studenti nella costruzione della conoscenza, nello sviluppo del pensiero critico e nella maturazione di una consapevolezza etica. Attingendo al patrimonio della ricerca pedagogica e alle esperienze concrete documentate nei materiali di riferimento (Ranieri, Cuomo, Biagini, 2024), questo paragrafo esplora in dettaglio un ventaglio di strategie didattiche particolarmente promettenti per coltivare le diverse dimensioni dell'AI literacy, illustrandone la ratio, l'implementazione pratica, gli adattamenti per età e le potenzialità cooperative.

2.2.1 Approcci pedagogici per l'insegnamento dell'AI literacy

Prima di addentrarci nelle specifiche strategie didattiche per ciascuna dimensione dell'AI literacy, è opportuno delineare gli approcci pedagogici che hanno guidato la loro elaborazione. Questi approcci costituiscono la cornice teorica all'interno della quale si collocano le diverse attività e metodologie proposte.

Il primo approccio pedagogico di riferimento è il costruttivismo, che concepisce l'apprendimento come un processo attivo di costruzione della conoscenza da parte del discente (Piaget, 1964; Vygotskij, 1978). Secondo questa prospettiva, gli studenti non sono recipienti passivi di informazioni, ma attori che elaborano attivamente significati a partire dalle proprie esperienze e interazioni con l'ambiente. Nel contesto dell'AI literacy, l'approccio costruttivista si traduce in attività che incoraggiano gli studenti a esplorare attivamente l'intelligenza artificiale, a formulare ipotesi sul suo funzionamento, a testare queste ipotesi attraverso l'interazione con sistemi di AI e a rielaborare le proprie concezioni in base ai risultati di queste esperienze.

Il secondo approccio di riferimento è l'apprendimento esperienziale, teorizzato da Dewey (1938) e sviluppato da Kolb (1984), che enfatizza il ruolo dell'esperienza diretta e della riflessione su di essa nel processo di apprendimento. Secondo questo approccio, l'apprendimento efficace si sviluppa attraverso un ciclo che comprende l'esperienza concreta, l'osservazione riflessiva, la concettualizzazione astratta e la sperimentazione attiva. Applicato all'AI literacy, l'apprendimento esperienziale si traduce in attività che permettono agli studenti di interagire direttamente con sistemi di AI, di riflettere su queste interazioni, di elaborare concetti e teorie a partire da esse e di applicare questi concetti in nuove situazioni.

Il terzo approccio è l'apprendimento per scoperta, proposto da Bruner (1961), che sottolinea l'importanza di coinvolgere gli studenti in processi di esplorazione e scoperta autonoma, piuttosto che nella semplice ricezione di conoscenze già strutturate. Secondo questa prospettiva, l'apprendimento è più significativo ed efficace quando gli studenti scoprono da sé i concetti e le relazioni, attraverso l'osservazione, l'analisi e la sperimentazione. Nel contesto dell'AI literacy, l'apprendimento per scoperta si traduce in attività che pongono gli studenti di fronte a problemi, situazioni o fenomeni legati all'intelligenza artificiale, incoraggiandoli a esplorare, formulare ipotesi e scoprire autonomamente principi e concetti.

Il quarto approccio è l'apprendimento basato su problemi (*problem-based learning*), che situa l'apprendimento all'interno di contesti problematici reali e significativi (Savery & Duffy, 1995). Secondo questa prospettiva, gli studenti apprendono più efficacemente quando sono chiamati a risolvere problemi autentici e complessi, che richiedono l'integrazione di conoscenze e competenze diverse. Applicato all'AI literacy, l'apprendimento basato su problemi si traduce in attività che presentano agli studenti problemi reali legati all'intelligenza artificiale, come l'analisi di bias negli algoritmi, la progettazione di sistemi di AGLI etici o la valutazione dell'impatto dell'AI in contesti specifici.

Infine, il quinto approccio è l'apprendimento collaborativo, che enfatizza l'importanza dell'interazione sociale e della costruzione condivisa di significati nel processo di apprendimento (Johnson & Johnson, 1999). Secondo questa prospettiva, l'apprendimento è facilitato dalla collaborazione tra pari, dal confronto di idee e dalla negoziazione di significati. Nel contesto dell'AI literacy, l'apprendimento collaborativo si traduce in attività che incoraggiano gli studenti a lavorare insieme, a discutere concetti e problemi legati all'intelligenza artificiale, a condividere intuizioni e a costruire collettivamente conoscenze e comprensioni.

Questi approcci pedagogici non sono mutualmente esclusivi, ma si integrano e si completano a vicenda, offrendo una cornice teorica ricca e articolata per lo sviluppo di strategie didattiche efficaci per l'AI literacy. La loro combinazione permette di rispondere alla complessità e alla multidimensionalità dell'intelligenza artificiale, coinvolgendo gli studenti in processi di apprendimento attivi, significativi e socialmente situati.

Le strategie didattiche sviluppate per l'AI literacy attingono a questi approcci pedagogici, adattandoli alle specificità di ciascuna dimensione del framework e alle caratteristiche cognitive ed esperienziali delle diverse fasce d'età. Questa modulazione permette di proporre percorsi didattici che sono al contempo concettualmente rigorosi, pedagogicamente fondati e significativamente rilevanti per gli studenti.

2.2.2 Strategie didattiche per la dimensione conoscitiva

La dimensione conoscitiva dell'AI literacy si focalizza sulla comprensione di cosa sia l'intelligenza artificiale, come funzioni e quali siano le sue principali applicazioni. Le strategie didattiche per questa dimensione mirano a sviluppare negli studenti una comprensione concettuale solida e articolata dell'intelligenza artificiale, adattata al loro livello di sviluppo cognitivo e alle loro conoscenze pregresse.

La costruzione di una solida base conoscitiva sull'IA, che includa la sua storia, le sue definizioni e i suoi concetti fondamentali, richiede approcci che superino la lezione frontale tradizionale e stimolino l'interrogazione attiva e la costruzione condivisa di significato.

Philosophy for children (P4C): L'arte di interrogare l'intelligenza

Ratio pedagogica

Come ampiamente esplorato in Ranieri, Cuomo, Biagini (2024, Cap. 2.2.1), la P4C si configura come un ambiente privilegiato per l'indagine concettuale. La sua ratio pedagogica risiede nella trasformazione della classe in una comunità di ricerca filosofica, dove il pensiero complesso (riflessivo, critico, creativo, valoriale, come definito da Lipman, 2003) viene coltivato attraverso il dialogo strutturato su domande autenticamente aperte. Applicata all'IA, questa metodologia permette di affrontare collettivamente questioni fondamentali e spesso controintuitive: "Cosa distingue un essere pensante da una macchina che simula il pensiero?", "Possiamo considerare 'intelligente' un algoritmo che supera un umano a scacchi ma non sa preparare un caffè?", "Le definizioni di IA riflettono più le capacità delle macchine o le nostre aspettative su di esse?".

Implementazione dettagliata

Il processo inizia con la presentazione di uno stimolo accuratamente selezionato per la sua capacità di generare dissonanza cognitiva o curiosità (un breve racconto di fantascienza, un'immagine ambigua di un robot, un video che mostra un'interazione uomo-macchina sorprendente, un semplice oggetto tecnologico). Segue una fase di raccolta delle domande individuali o a piccoli gruppi, facilitata da tecniche come il brainstorming silenzioso o l'uso di post-it (fisici o virtuali). La comunità, guidata dal facilitatore, seleziona democraticamente la domanda-focus per la sessione. Il cuore dell'attività è il dialogo filosofico, in cui il facilitatore non fornisce risposte ma rilancia la discussione con domande maieutiche ("Cosa intendi per...?", "Sei d'accordo con...?", "Quali potrebbero essere le conseguenze di questa idea?"), incoraggiando l'ascolto attivo, la formulazione di ipotesi, la ricerca di ragioni e la critica costruttiva delle argomentazioni altrui. La sessione si conclude con una riflessione metacognitiva sul processo di pensiero e sulle eventuali evoluzioni delle proprie idee.

Adattamento per età

Alla primaria, gli stimoli sono ancorati al vissuto (giocattoli, cartoni animati) e le domande vertono su concetti come "vivo/non vivo", "pensante/non pensante", "amico/strumento". Il dialogo è più guidato e si utilizzano supporti visivi. Alla secondaria, si possono introdurre testi più astratti (Turing, Asimov, definizioni istituzionali di IA) e affrontare concetti come coscienza artificiale, IA debole/forte, implicazioni del Test di Turing. L'attività di progettazione di una sessione P4C, come quella realizzata nella sperimentazione (Cap. 4.1.3), diventa un esercizio fondamentale per futuri insegnanti, richiedendo loro di selezionare stimoli, formulare domande guida e anticipare possibili traiettorie del dialogo per diverse fasce d'età.

Modalità cooperativa e valutazione

La P4C è intrinsecamente cooperativa. Online, strumenti come lavagne collaborative (Miro, Jamboard) o piattaforme di discussione strutturata possono supportare il processo. La valutazione è formativa e si basa sull'osservazione della qualità della partecipazione al dialogo (pertinenza degli interventi, capacità argomentativa, ascolto attivo, rispetto delle regole della comunità di ricerca) piuttosto che sulla "correttezza" delle opinioni espresse.

Questa attività presenta uno straordinario potenziale per lo sviluppo del pensiero critico, autonomo e dialogico. Tuttavia, richiede una formazione specifica e approfondita del facilitatore per gestire il dialogo in modo non direttivo ma produttivo, evitando la dispersione o la prevaricazione.

Coding Unplugged: algoritmi con il corpo

Ratio pedagogica

Questa famiglia di attività (Ranieri, Cuomo, Biagini, 2024, Cap. 2.2.2), ispirata al costruzionismo di Papert (1980), introduce i concetti fondamentali dell'informatica e del pensiero computazionale senza l'uso di computer. La ratio pedagogica è quella di rendere concreti, visibili e manipolabili i processi algoritmici attraverso il corpo, il movimento, il disegno e il gioco, favorendo una comprensione intuitiva e incarnata.

Implementazione dettagliata

Le attività possono variare:

- **Body Coding:** gli studenti diventano "processori" che eseguono istruzioni date da un compagno "programmatore" per compiere azioni nello spazio (es. costruire una torre di blocchi, seguire un percorso).
- **Pixel Art:** utilizzando griglie e colori, gli studenti decodificano o creano immagini basate su istruzioni numeriche o simboliche, comprendendo la logica della rappresentazione digitale bitmap.
- **Robotica Unplugged:** un bambino (il "robot") esegue le istruzioni fornite dai compagni (i "programmatori") per navigare un percorso, magari disseminato di ostacoli imprevisti che richiedono strategie di "debugging" collettivo (es. introdurre un'istruzione "fai un passo indietro" o "aggira l'ostacolo").
- **Sorting Networks:** gli studenti stessi diventano gli elementi da ordinare e si muovono attraverso una rete disegnata a terra che rappresenta un algoritmo di ordinamento (es. bubble sort), visualizzando fisicamente il processo.

- **Pair Programming Unplugged:** due studenti collaborano per risolvere un puzzle logico (es. Sudoku, Mastermind) o per scrivere su carta le istruzioni per un compito complesso (es. preparare una ricetta), alternandosi nei ruoli di "driver" (chi scrive/agisce) e "navigator" (chi osserva, suggerisce, controlla).

Adattamento per età

La versatilità è massima. Alla primaria, si privilegiano attività motorie, comandi semplici, griglie piccole per la pixel art. Alla secondaria I grado, si introducono concetti più complessi come cicli, condizioni (SE... ALLORA...), variabili semplici, algoritmi di ordinamento base. Alla secondaria II grado, si possono affrontare algoritmi più sofisticati o usare il coding unplugged come metafora per discutere di complessità computazionale o limiti degli algoritmi.

Modalità cooperativa e valutazione

Molte attività sono naturalmente cooperative (programmare il robot, pair programming). La valutazione formativa si basa sull'osservazione della corretta esecuzione/formulazione delle istruzioni, della capacità di identificare e correggere errori (debugging), della collaborazione all'interno del gruppo e della verbalizzazione dei processi logici utilizzati.

Queste attività sono estremamente inclusive (non richiedono competenze tecnologiche pregresse né dispositivi) e favoriscono l'apprendimento attivo, ludico e incarnato. Tuttavia, possono essere percepite come "infantili" da studenti più grandi se non ne viene esplicitato chiaramente il collegamento con i concetti informatici fondamentali e la loro valenza metaforica.

La Tabella 2.4 sintetizza le principali strategie didattiche per la dimensione conoscitiva dell'AI literacy, articolate per livello scolastico.

Tabella 2.4. Strategie didattiche per la dimensione conoscitiva per livello scolastico

Livello scolastico	Strategie didattiche	Esempi di attività
Scuola primaria	- Metafore e analogie - Narrazioni - Drammatizzazione e gioco di ruolo - Philosophy for Children	- "È intelligente un...?": discussione guidata sul concetto di intelligenza - "La storia del robot che impara": narrazione e illustrazione - "Facciamo il robot": gioco di ruolo in cui un bambino interpreta un robot
Scuola secondaria di I grado	- Apprendimento basato sull'indagine - Simulazioni e modelli semplificati - Analisi comparativa - Pensiero computazionale	- "Come funziona...?": indagine su un'applicazione di AI - Programmazione con Scratch per simulare comportamenti di AI - "AI vs Umano": confronto tra intelligenza artificiale e umana - Attività di "unplugged coding" per comprendere gli algoritmi
Scuola secondaria di II grado	- Studio di caso - Ricerca e presentazione - Partecipazione a conferenze o seminari - Analisi	- Analisi di un sistema di AI in un contesto specifico - Presentazione su un approccio specifico all'AI - Partecipazione a webinar o

	di articoli scientifici o divulgativi	seminari tenuti da esperti - Analisi critica di articoli su temi legati all'AI
--	---------------------------------------	--

Le strategie didattiche per la dimensione conoscitiva dell'AI literacy si evolvono dunque in funzione del livello di sviluppo cognitivo degli studenti, passando da approcci più concreti e ludici nella scuola primaria, a metodologie più strutturate e analitiche nella scuola secondaria di secondo grado. Questa progressione permette di costruire gradualmente una comprensione solida e articolata dell'intelligenza artificiale, che costituisce la base per lo sviluppo delle altre dimensioni dell'AI literacy.

2.2.3 Strategie didattiche per la dimensione operativa

La dimensione operativa dell'AI literacy si concentra sulle competenze pratiche necessarie per interagire con l'intelligenza artificiale, utilizzarla in modo efficace e, a livelli più avanzati, partecipare al suo sviluppo. Le strategie didattiche per questa dimensione mirano a sviluppare negli studenti la capacità di utilizzare strumenti e applicazioni di AI, di comprendere e applicare i principi del pensiero computazionale, e di analizzare e gestire dati.

Comprendere i meccanismi operativi dell'IA, in particolare del machine learning, richiede un passaggio dalla teoria alla pratica, sperimentando direttamente come le macchine "imparano" dai dati.

Trial and Error con piattaforme educative: addestrare l'IA per prove e errori

Ratio pedagogica

Questa strategia (Ranieri, Cuomo, Biagini, 2024, Cap. 3.2.1) sfrutta strumenti digitali progettati per rendere accessibile l'esperienza dell'addestramento di modelli di machine learning. La ratio pedagogica è quella di far vivere agli studenti un processo analogo a quello degli scienziati dei dati, basato su cicli iterativi di fornitura di dati, addestramento del modello, test, valutazione dei risultati e affinamento dei dati o del modello stesso, comprendendo l'importanza cruciale dei dati e il concetto di "apprendimento" della macchina.

Implementazione dettagliata

Si utilizzano piattaforme come "AI per gli Oceani" (Code.org) o "Teachable Machine" (Google). In "AI per gli Oceani", gli studenti "insegnano" all'IA a distinguere tra pesci e rifiuti marini mostrandole esempi etichettati. Osservano come l'IA classifica nuovi oggetti e possono aggiungere altri esempi per migliorarne l'accuratezza. In "Teachable Machine", gli studenti possono creare modelli di classificazione personalizzati (es. riconoscere diversi tipi di frutta, diversi suoni, diverse pose corporee) fornendo campioni attraverso la webcam o il microfono. Sperimentano direttamente come la quantità, la varietà e la qualità dei dati di addestramento influenzino le prestazioni del modello, scoprendo magari anche come introdurre involontariamente dei bias (es. addestrando il modello solo con mele rosse, non riconoscerà quelle verdi). La riflessione guidata dal docente è fondamentale per collegare l'esperienza pratica ai concetti teorici di apprendimento supervisionato, dataset, etichette, accuratezza, bias.

Adattamento per età

"AI per gli Oceani" è ideale per primaria e secondaria I grado per la sua interfaccia semplice e il contesto motivante. "Teachable Machine" è più potente e flessibile, adatta dalla secondaria I grado in su. La complessità del compito di classificazione può essere graduata: dalla distinzione tra due oggetti semplici alla classificazione di categorie più numerose o sfumate.

Modalità cooperativa e valutazione

Gli studenti possono lavorare in piccoli gruppi per decidere quali dati raccogliere, come etichettarli, come testare il modello e come interpretare i risultati. La valutazione formativa può avvenire tramite l'osservazione del processo, l'analisi dei modelli creati e la discussione guidata sulla comprensione dei concetti chiave.

Questa strategia rende estremamente concreto e coinvolgente il concetto astratto di machine learning e stimola la creatività nel trovare applicazioni. Tuttavia, richiede dispositivi con webcam/microfono e connessione internet. È cruciale la guida del docente per evitare che l'attività rimanga a livello di "gioco" tecnologico senza una reale comprensione dei principi sottostanti.

Gamification per la Predizione: Sfidare l'Incertezza

Ratio pedagogica

L'uso di meccaniche ludiche (Ranieri, Cuomo, Biagini, 2024, Cap. 3.2.3) può rendere più accessibile e motivante la comprensione dei processi di stima e predizione basati sui dati, che sono al cuore di molte applicazioni IA (dalle previsioni meteo ai sistemi di raccomandazione). La ratio pedagogica è quella di trasformare un compito cognitivamente impegnativo (analizzare dati, identificare pattern, formulare ipotesi probabilistiche) in una sfida coinvolgente, fornendo feedback immediati e stimolando la competizione o la collaborazione.

Implementazione dettagliata

Si creano scenari di gioco in cui gli studenti devono prendere decisioni o fare previsioni in condizioni di incertezza, basandosi su set di dati forniti. In "Un picnic in montagna" (Ranieri, Cuomo, Biagini, 2024, Cap. 3.2.3), i gruppi analizzano dati sugli avvistamenti passati di animali selvatici in diverse aree (caratterizzate da variabili come tipo di vegetazione, presenza umana, rumore) per predire la probabilità di incontrarli in una nuova area. In "Zombie Outbreak" (Code.org), analizzano dati sulla diffusione di un'epidemia in diverse città (popolazione, densità, collegamenti) per prevederne l'evoluzione. Gli studenti devono esplicitare le loro strategie di analisi (fare medie, cercare correlazioni, valutare similarità tra aree) e giustificare le loro predizioni. Il confronto con i dati reali (preparati dal docente) e l'assegnazione di punteggi (per corrispondenza esatta, entro un certo margine) forniscono il feedback ludico. Si possono aggiungere classifiche, badge per strategie particolarmente brillanti, o missioni aggiuntive (es. prevedere non solo dove, ma anche quando).

Adattamento per età

Questi scenari sono adatti per la secondaria di I e II grado. La complessità dei dati e delle variabili può essere modulata. Per la primaria, si possono creare giochi da tavolo o digitali più semplici che introducano il concetto di probabilità e pattern recognition in contesti narrativi (es. "Indovina dove apparirà il tesoro sulla mappa, basandoti sugli indizi precedenti").

Modalità cooperativa e valutazione

Le attività si prestano bene al lavoro di gruppo, stimolando la discussione sulle strategie di analisi e predizione. La valutazione formativa osserva la capacità di leggere e interpretare i dati, di identificare pattern, di formulare ipotesi motivate e di riflettere criticamente sul grado di incertezza delle previsioni. Il debriefing finale è essenziale per collegare l'esperienza ludica ai concetti di data analysis, modellizzazione e predizione nell'IA.

Questa strategia presenta un elevato potenziale di engagement e motivazione e rende tangibile il processo predittivo. Tuttavia, c'è il rischio di eccessiva semplificazione dei modelli reali di IA; è quindi necessaria una progettazione accurata dello scenario ludico e un forte ancoraggio ai concetti chiave durante il debriefing.

La Tabella 2.5 sintetizza le principali strategie didattiche per la dimensione operativa dell'AI literacy, articolate per livello scolastico.

Tabella 2.5. Strategie didattiche per la dimensione operativa per livello scolastico

Livello scolastico	Strategie didattiche	Esempi di attività
Scuola primaria	- Apprendimento attraverso il gioco - Sperimentazione diretta con applicazioni di AI - Programmazione visuale - Coding unplugged	- "AI e gli oceani": classificazione di pesci e rifiuti - Interazione con "Quick, Draw!" di Google - Creazione di storie interattive con ScratchJr - "Il robot e il labirinto": attività di coding senza computer
Scuola secondaria di I grado	- Gamification - Apprendimento basato su progetti - Programmazione a coppie - Uso di tool specifici per l'AI	- "Un picnic in montagna": prevedere con la gamification - Creazione di un semplice chatbot con Scratch - Attività di pair programming per lo sviluppo di algoritmi - Creazione di modelli di ML con "Machine Learning for Kids"
Scuola secondaria di II grado	- Progettazione e sviluppo di applicazioni di AI - Analisi e manipolazione di dataset reali - Partecipazione a hackathon o competizioni - Uso di piattaforme avanzate per l'AI	- Sviluppo di un'applicazione di analisi dei dati ambientali - Analisi di dataset climatici per previsioni future - Partecipazione a un hackathon su temi sociali - Creazione di modelli di riconoscimento con Google Teachable Machine

Le strategie didattiche per la dimensione operativa dell'AI literacy evolvono dunque in funzione del livello di sviluppo cognitivo e delle competenze tecniche degli studenti, passando da attività ludiche e concrete nella scuola primaria, a progetti complessi e significativi nella scuola secondaria di secondo grado. Questa progressione permette di costruire gradualmente competenze operative solide e articolate, che costituiscono un elemento essenziale dell'AI literacy.

2.2.4 Strategie didattiche per la dimensione critica

La dimensione critica dell'AI literacy si focalizza sullo sviluppo di capacità analitiche e valutative necessarie per comprendere le potenzialità e i limiti dell'intelligenza artificiale, valutarne l'affidabilità e comprenderne l'impatto sociale. Le strategie didattiche per questa dimensione mirano a sviluppare negli studenti la capacità di interpretare e valutare criticamente i sistemi di AI, di comprenderne l'impatto sociale e di sviluppare una alfabetizzazione ai dati e all'informazione.

Sviluppare un pensiero critico sull'IA implica andare oltre la superficie, interrogandosi sui processi interni (per quanto possibile), valutando l'affidabilità degli output e riconoscendo i limiti e i rischi intrinseci.

Inquiry based learning (IBL) con IA generativa: l'indagine come antidoto all'accettazione passiva

Ratio pedagogica

L'interazione con Large Language Models (LLM) come ChatGPT, se non guidata, rischia di promuovere un consumo acritico di informazioni. L'IBL (Ranieri, Cuomo, Biagini, 2024, Cap. 4.2.1) offre un quadro metodologico per trasformare questa interazione in un'opportunità di apprendimento critico. La ratio pedagogica è quella di spostare il focus dalla risposta dell'IA al processo di interrogazione, verifica e validazione condotto dallo studente, promuovendo l'information literacy e la consapevolezza epistemologica.

Implementazione dettagliata

Il ciclo di indagine guida il processo:

Orientamento: si parte da una domanda autentica o un problema complesso (es. "Quali sono le cause del cambiamento climatico e come l'IA può contribuire a mitigarlo?").

Concettualizzazione: gli studenti definiscono le sotto-domande chiave, attivano le preconoscenze e formulano ipotesi iniziali, potendo usare ChatGPT come strumento di brainstorming preliminare ("Quali sono i principali fattori...?").

Investigazione: è la fase cruciale. Gli studenti interrogano ChatGPT con prompt specifici e mirati, ma sono guidati a trattare le risposte come ipotesi da verificare. Devono attivamente cercare fonti autorevoli e diversificate (articoli scientifici, report istituzionali, libri) per confrontare, corroborare o confutare le informazioni fornite dall'IA. Si analizzano le risposte dell'IA per identificare possibili bias, omissioni, eccessive semplificazioni o "allucinazioni" (informazioni inventate).

Conclusione: gli studenti sintetizzano le informazioni raccolte da tutte le fonti, costruendo una propria risposta argomentata al problema iniziale, che integri criticamente i contributi dell'IA con quelli delle altre fonti.

Discussione: i risultati vengono condivisi e discussi collettivamente, riflettendo non solo sul contenuto, ma anche sul processo di ricerca e sulle sfide dell'interazione critica con l'IA generativa.

Adattamento per età

Il ciclo completo è più adatto per secondaria II grado e contesti universitari. Per la secondaria I grado, si possono proporre indagini più circoscritte: verificare la veridicità di un fatto specifico riportato da ChatGPT, confrontare due diverse spiegazioni (una dell'IA, una da un libro di testo) dello stesso fenomeno, analizzare il tono o lo stile di un testo generato dall'IA. L'insegnante può fornire una checklist per la valutazione delle fonti e modellare (tramite think aloud) come formulare prompt efficaci e come analizzare criticamente le risposte.

Modalità cooperativa e valutazione

I gruppi possono dividersi i compiti di ricerca (chi interroga l'IA, chi cerca fonti tradizionali), confrontare i risultati e collaborare alla stesura della sintesi critica. La valutazione formativa si

concentra sulla qualità dei prompt, sulla capacità di ricerca e valutazione delle fonti, sulla profondità dell'analisi critica delle risposte dell'IA e sulla qualità della sintesi finale.

Questa strategia sviluppa competenze essenziali per il XXI secolo (information literacy, pensiero critico, problem solving complesso) e trasforma l'uso di strumenti potenti come ChatGPT da rischio a opportunità educativa. Tuttavia, richiede una buona pianificazione dell'indagine, accesso a fonti diversificate e una forte guida metacognitiva da parte del docente.

Problem solving e explainability: aprire la scatola nera (o capire perché è chiusa)

Ratio pedagogica

Affrontare problemi concreti legati al funzionamento "interno" dell'IA (Ranieri, Cuomo, Biagini, 2024, Cap. 4.2.2) è cruciale per sviluppare una comprensione critica della cosiddetta "explainability" (XAI), ovvero la capacità (o l'incapacità) di spiegare come un sistema IA arrivi a una certa decisione. La ratio pedagogica è quella di rendere tangibili le sfide della trasparenza, dell'interpretabilità e dei bias intrinseci, stimolando un atteggiamento investigativo e critico verso gli algoritmi.

Implementazione dettagliata

Si possono usare diverse tipologie di problemi:

Analisi di errori/risultati anomali: si presenta un caso reale o simulato in cui un sistema IA ha prodotto un output palesemente errato o controintuitivo (es. la previsione di neve estiva, un sistema di riconoscimento facciale che fallisce su certi gruppi etnici, una raccomandazione assurda). Gli studenti, in gruppo, devono ipotizzare le possibili cause (problemi nei dati di input, bias nel dataset di addestramento, limiti intrinseci del modello algoritmico, errata configurazione) e, se possibile, verificare le loro ipotesi (es. analizzando i dati se disponibili, o ricercando informazioni sul funzionamento di quel tipo di sistema).

Costruzione di modelli trasparenti vs opachi: si chiede agli studenti di risolvere lo stesso problema (es. classificare e-mail come spam/non spam) usando prima un modello intrinsecamente interpretabile (come un albero decisionale semplice, che possono costruire anche manualmente o con tool visuali) e poi interagendo con un modello più complesso (una "scatola nera" simulata) di cui devono cercare di inferire il funzionamento osservando gli input e gli output. Questo confronto aiuta a comprendere il trade-off tra accuratezza e interpretabilità.

Confronto theory-driven vs data-driven: ad esempio in fisica, si potrebbe confrontare la traiettoria di un proiettile calcolata con le leggi della cinematica (theory-driven, spiegabile) con una predizione basata su un modello ML addestrato su migliaia di lanci precedenti (data-driven, potenzialmente più accurato ma meno interpretabile). In biologia, si potrebbe confrontare la classificazione di specie basata su chiavi dicotomiche (theory-driven) con una classificazione basata su clustering automatico di caratteristiche (data-driven). Questi esempi aiutano a capire i diversi approcci dell'IA e le sfide della spiegabilità dei modelli basati sui dati.

Attività ludiche sull'inferenza: l'attività "Scopri la regola" (Ranieri, Cuomo, Biagini, 2024) per la primaria è un'introduzione semplice al concetto che le azioni di un sistema (il "robot") sono determinate da una logica interna che non è immediatamente visibile e va inferita.

Adattamento per età

"Scopri la regola" per la primaria. Costruzione di alberi decisionali semplici e analisi di errori evidenti per la secondaria I grado. Analisi di casi di bias più complessi, confronto Theory/Data driven, interazione con modelli "black box" simulati per la secondaria II grado e oltre.

Modalità cooperativa e valutazione

Il lavoro di gruppo è ideale per l'analisi dei problemi, la formulazione di ipotesi e la discussione delle spiegazioni. La valutazione formativa osserva la capacità di analisi logica, di formulazione di ipotesi causali, di comprensione dei concetti di bias e trasparenza, e la consapevolezza dei limiti dell'interpretabilità dell'IA.

Questa strategia demistifica l'IA mostrando che non è "magia", ma risultato di processi (a volte opachi) e sviluppa pensiero analitico e investigativo. Tuttavia, può richiedere al docente competenze tecniche per predisporre i problemi o guidare l'analisi di modelli complessi.

La Tabella 2.6 sintetizza le principali strategie didattiche per la dimensione critica dell'AI literacy, articolate per livello scolastico.

Tabella 2.6. Strategie didattiche per la dimensione critica per livello scolastico

Livello scolastico	Strategie didattiche	Esempi di attività
Scuola primaria	- Esercizi di confronto - Racconti e discussioni guidate - Mappe concettuali semplici - Giochi di ruolo	- Confronto tra risposte di assistenti virtuali e umani - Discussione sul racconto "Il robot che confonde gatti e conigli" - Mappa delle differenze tra intelligenza umana e artificiale - Gioco di ruolo "Il negozio automatizzato"
Scuola secondaria di I grado	- Analisi di casi di studio semplificati - Apprendimento basato sulla critica - Media education critica - Dibattiti strutturati	- Analisi di un sistema di raccomandazione musicale - Attività sull'explainability dei sistemi di AI - Analisi dell'influenza degli algoritmi sui social media - Dibattito "L'AI nel sistema educativo"
Scuola secondaria di II grado	- Inquiry-based learning - Analisi di bias e disuguaglianze - Critical data literacy - Problem solving critico	- "ChatGPT come campo di allenamento per il pensiero critico" - Analisi di un sistema di assunzione basato sull'AI - Esame critico della raccolta e dell'uso dei dati - Analisi di un sistema di decision-making

Le strategie didattiche per la dimensione critica dell'AI literacy evolvono dunque in funzione del livello di sviluppo cognitivo e delle capacità analitiche degli studenti, passando da semplici esercizi di osservazione e confronto nella scuola primaria, a complesse analisi critiche e valutative nella scuola secondaria di secondo grado. Questa progressione permette di costruire gradualmente capacità critiche solide e articolate, essenziali per un'interazione consapevole con l'intelligenza artificiale.

2.2.5 Strategie didattiche per la dimensione etica

La dimensione etica dell'AI literacy si focalizza sulla consapevolezza delle questioni etiche legate all'intelligenza artificiale, sulla comprensione dei principi etici fondamentali e sulla capacità di analizzare dilemmi etici complessi. Le strategie didattiche per questa dimensione mirano a sviluppare negli studenti la capacità di riflettere sulle implicazioni etiche dell'AI, di comprenderle i principi che dovrebbero guidarne lo sviluppo e l'applicazione, e di formulare giudizi etici informati e ponderati.

L'educazione all'etica dell'IA non può limitarsi alla trasmissione di norme o principi, ma deve coltivare la capacità di ragionamento morale, di deliberazione collettiva e di assunzione di responsabilità in situazioni complesse e spesso inedite.

Debate strutturato: palestra di argomentazione etica

Ratio pedagogica

Questa metodologia (Ranieri, Cuomo, Biagini, 2024, Cap. 5.2.1), basata sul confronto regolamentato di tesi contrapposte, è particolarmente efficace per esplorare la natura controversa delle questioni etiche legate all'IA. La ratio pedagogica è duplice: da un lato, sviluppare competenze argomentative (ricercare evidenze, costruire ragionamenti validi, confutare tesi avverse, comunicare efficacemente); dall'altro, promuovere la comprensione della complessità etica, riconoscendo la plausibilità di diverse posizioni morali e la necessità di un confronto rispettoso e basato su ragioni.

Implementazione dettagliata

Si parte dalla definizione di una mozione chiara, bilanciata e realmente dibattibile (es. "L'uso di sistemi di riconoscimento facciale da parte delle forze dell'ordine dovrebbe essere vietato"; "Le aziende sviluppatrici di IA dovrebbero essere legalmente responsabili per i danni causati dai loro prodotti"; "È eticamente giustificabile usare chatbot IA per fornire supporto psicologico?"). La classe viene divisa in squadre (Pro e Contro), che hanno tempo per ricercare informazioni, preparare i discorsi argomentativi e le possibili confutazioni. Il dibattito si svolge secondo un formato preciso (es. Lincoln-Douglas, Karl Popper) con tempi contingentati per ogni intervento (discorsi costruttivi, confutazioni, repliche). Una giuria (composta dal docente e/o da studenti) valuta la performance sulla base di criteri espliciti (qualità delle argomentazioni, uso delle evidenze, efficacia comunicativa, rispetto delle regole). Segue un debriefing collettivo in cui si riflette non solo sull'esito, ma soprattutto sul processo argomentativo e sulle diverse prospettive emerse.

Adattamento per età

Il formato completo è adatto per la secondaria di I e II grado. La complessità della mozione e la profondità della ricerca richiesta vanno calibrate. Per la primaria, si possono organizzare forme più semplici di "discussione strutturata" o "gioco di ruolo" su storie che presentano conflitti di valori legati alla tecnologia (es. condividere o no una foto online, aiutare un amico a barare in un videogioco).

Modalità cooperativa e valutazione

La preparazione del dibattito avviene in squadra ed è un momento cruciale di apprendimento cooperativo. La valutazione può essere individuale (sulla performance del singolo oratore) e/o di squadra, basata su griglie di osservazione condivise. È importante valutare il processo argomentativo più che la "vittoria" nel dibattito.

Questa strategia è eccellente per sviluppare pensiero critico-etico, competenze argomentative e comunicative, capacità di lavorare in gruppo e di considerare prospettive diverse. Tuttavia, richiede tempo per la preparazione, una buona scelta della mozione e una gestione ferma ma equilibrata del dibattito da parte del moderatore.

Case study (studio di caso): ancorare l'etica alla realtà

Ratio pedagogica

L'analisi di casi concreti (Ranieri, Cuomo, Biagini, 2024, Cap. 5.2.2) permette di calare i principi etici astratti in situazioni problematiche specifiche, rendendo la riflessione morale più tangibile e significativa. La ratio pedagogica è quella di sviluppare il giudizio pratico (la *phronesis* aristotelica), la capacità di analizzare situazioni complesse identificando i fattori eticamente rilevanti, di considerare le conseguenze delle diverse azioni possibili e di formulare decisioni ponderate e giustificate.

Implementazione dettagliata

Si presenta un caso ben documentato, reale o verosimile, che presenti un chiaro dilemma etico legato all'IA (es. il giocattolo "intelligente" che pone problemi di privacy e manipolazione infantile; l'algoritmo di ammissione scolastica che discrimina studenti da contesti svantaggiati - Ranieri, Cuomo, Biagini, 2024, Cap. 5.2.2; il "docente robot" che solleva questioni di responsabilità e relazione educativa; un caso di diagnosi medica errata da parte di un sistema IA; l'uso di droni autonomi in contesti bellici). Gli studenti, lavorando in piccoli gruppi, leggono attentamente il caso, identificano i fatti principali, gli stakeholder coinvolti, i valori etici in conflitto (es. efficienza vs equità, sicurezza vs libertà, innovazione vs precauzione), le possibili opzioni di azione e le loro prevedibili conseguenze a breve e lungo termine. Ogni gruppo elabora una proposta di soluzione o una linea d'azione motivata, argomentandola sulla base di principi etici (quelli discussi nel modulo 4 del curriculum) e delle specificità del caso. Le proposte vengono poi presentate e discusse in plenaria, confrontando le diverse analisi e soluzioni.

Adattamento per età

La complessità del caso e la profondità dell'analisi richiesta vanno graduate. Per la primaria, si possono usare storie illustrate o brevi racconti che presentino dilemmi semplici legati all'uso di internet o dei videogiochi (condivisione, cyberbullismo, tempo di utilizzo). Per la secondaria I grado, casi come quello del giocattolo o semplici scenari di bias (es. un sistema di raccomandazione che propone solo certi tipi di contenuti). Per la secondaria II grado e oltre, casi più complessi che richiedono di considerare aspetti legali, economici e sociali più ampi.

Modalità cooperativa e valutazione: L'analisi del caso in gruppo è fondamentale per far emergere diverse interpretazioni e prospettive. La valutazione formativa si basa sulla qualità dell'analisi del caso (identificazione degli elementi rilevanti), sulla coerenza dell'applicazione dei principi etici, sulla profondità della considerazione delle conseguenze e sulla capacità di argomentare la soluzione proposta. Si può valutare sia il processo di discussione di gruppo che il prodotto finale (la presentazione della soluzione).

Potenzialità/criticità: Collega strettamente la riflessione etica a contesti applicativi concreti. Sviluppa problem solving etico, capacità di analisi complessa e perspective-taking. Richiede la disponibilità di casi di studio ben costruiti, aggiornati e sufficientemente dettagliati per permettere un'analisi approfondita.

Scuola primaria

Nella scuola primaria, le strategie didattiche per la dimensione etica privilegiano l'introduzione a semplici questioni etiche attraverso narrazioni, giochi e discussioni guidate, che permettono ai bambini di iniziare a riflettere su temi come la privacy, la responsabilità e la fairness in modo accessibile e coinvolgente.

Una strategia particolarmente efficace a questo livello è l'uso di storie e favole etiche che presentano situazioni o dilemmi legati all'intelligenza artificiale in modo adatto ai bambini. Ad esempio, una storia potrebbe narrare di un robot giocattolo che registra le conversazioni dei bambini e le condivide con l'azienda produttrice, sollevando questioni di privacy. Come proposto da Ranieri e colleghi (Ranieri, Cuomo, Biagini, 2024), questo può costituire lo spunto per uno studio di caso alla primaria, in cui i bambini sono incoraggiati a riflettere su domande come: "Cosa pensi del fatto che il robot possa registrare le conversazioni? È giusto o sbagliato?", "Come ti sentiresti sapendo che il robot condivide le informazioni con l'azienda?", "Come possiamo fare in modo che il robot rispetti la privacy dei bambini?". Attraverso queste storie e le discussioni che ne seguono, i bambini iniziano a sviluppare una prima consapevolezza delle questioni etiche legate all'IA.

Un'altra strategia efficace è l'uso di circle time etico, un momento dedicato alla discussione collettiva su temi etici, in cui i bambini sono incoraggiati a esprimere le proprie opinioni, ascoltare quelle degli altri e riflettere insieme. Ad esempio, un circle time potrebbe focalizzarsi sulla domanda "È giusto che un robot prenda decisioni al posto nostro?", stimolando i bambini a riflettere su temi come autonomia, responsabilità e fiducia. Questa attività permette ai bambini di sviluppare non solo consapevolezza etica, ma anche capacità di ascolto, rispetto delle opinioni altrui e articolazione delle proprie idee.

L'uso di giochi di ruolo etici rappresenta un'ulteriore strategia preziosa per la dimensione etica nella scuola primaria. Questi giochi permettono ai bambini di esplorare situazioni etiche complesse in modo concreto e coinvolgente, assumendo diversi punti di vista e riflettendo sulle conseguenze delle diverse scelte. Ad esempio, un gioco di ruolo potrebbe simulare una classe in cui viene introdotto un robot insegnante, con alcuni bambini che interpretano studenti entusiasti, altri che sono preoccupati, altri ancora che fanno domande sul rispetto della privacy o sulla capacità del robot di comprendere le emozioni. Attraverso questa simulazione, i bambini possono esplorare le diverse prospettive e riflettere sulle implicazioni etiche dell'introduzione dell'IA nel contesto educativo.

La creazione di poster o disegni etici rappresenta un'ulteriore strategia efficace per la dimensione etica nella scuola primaria. Questa attività permette ai bambini di esprimere le loro riflessioni etiche in modo creativo e visuale, facilitando la comunicazione e la condivisione di idee. Ad esempio, i bambini potrebbero creare poster che illustrano "regole del buon comportamento" per i robot, riflettendo su quali principi etici dovrebbero guidare lo sviluppo e l'utilizzo dell'intelligenza artificiale. Questa attività non solo stimola la riflessione etica, ma promuove anche la creatività e la capacità di espressione visuale.

Scuola secondaria di primo grado

Nella scuola secondaria di primo grado, le strategie didattiche per la dimensione etica si orientano verso una riflessione più strutturata sulle implicazioni etiche dell'IA nella vita quotidiana, e un'analisi più approfondita dei principi e delle controversie etiche.

Una strategia particolarmente efficace a questo livello è lo studio di caso etico semplificato, che permette agli studenti di analizzare situazioni concrete che sollevano questioni etiche legate

all'intelligenza artificiale. Ad esempio, uno studio di caso potrebbe riguardare un "docente robot" che sostituisce parzialmente gli insegnanti nelle lezioni quotidiane, sollevando questioni di responsabilità: "Chi pensi debba essere responsabile se l'IA commette un errore durante l'insegnamento?", "Quali sono i possibili vantaggi e svantaggi di utilizzare un'IA per insegnare?". Questo caso di studio permette agli studenti di esplorare in modo guidato questioni etiche complesse, analizzando diverse prospettive e riflettendo sulle implicazioni delle diverse posizioni.

Un'altra strategia efficace è il dibattito etico strutturato, che incoraggia gli studenti a esaminare diverse prospettive su dilemmi etici legati all'intelligenza artificiale, sviluppando capacità di argomentazione, analisi critica e valutazione delle evidenze. Ad esempio, un dibattito potrebbe focalizzarsi sulla questione "È etico utilizzare sistemi di sorveglianza basati sull'IA nelle scuole?", con gruppi di studenti che preparano e presentano argomentazioni a favore e contro, basate su ricerche e evidenze. Questa attività permette agli studenti di approfondire la loro comprensione delle questioni etiche legate all'IA, di considerare diverse prospettive e di sviluppare un approccio più sfumato e articolato ai dilemmi etici.

La analisi multiprospettica rappresenta un'ulteriore strategia preziosa per la dimensione etica nella scuola secondaria di primo grado. Questa strategia incoraggia gli studenti a esaminare le questioni etiche legate all'IA da diverse prospettive (ad esempio, degli sviluppatori, degli utenti, dei soggetti i cui dati sono utilizzati, della società in generale), sviluppando una comprensione più articolata e contestualizzata. Ad esempio, un'attività potrebbe focalizzarsi sull'analisi dell'uso dell'IA per la personalizzazione dell'apprendimento, esaminandolo dal punto di vista degli studenti, degli insegnanti, dei genitori, degli sviluppatori e del sistema educativo nel suo complesso. Questa analisi permette agli studenti di comprendere la complessità delle questioni etiche legate all'IA e di sviluppare una visione più equilibrata e informata.

L'uso di media etici rappresenta un'ulteriore strategia efficace per la dimensione etica nella scuola secondaria di primo grado. Questa strategia utilizza film, documentari, articoli o video che esplorano questioni etiche legate all'intelligenza artificiale, offrendo stimoli per la riflessione e la discussione. Ad esempio, la visione di un breve documentario sull'uso dell'IA nella sorveglianza potrebbe essere seguita da una discussione guidata sulle implicazioni etiche di questa applicazione, sui diritti che potrebbero essere minacciati e sulle possibili salvaguardie da implementare. Questa attività permette agli studenti di collegare le questioni etiche astratte a esempi concreti e significativi, facilitando la comprensione e stimolando la riflessione critica.

Scuola secondaria di secondo grado

Nella scuola secondaria di secondo grado, le strategie didattiche per la dimensione etica si orientano verso un'analisi critica più approfondita delle questioni etiche legate all'IA e lo sviluppo di framework valutativi articolati, che permettono agli studenti di formulare giudizi etici informati e ponderati.

Una strategia particolarmente efficace a questo livello è il dibattito etico avanzato, come proposto nel libro "Artificial Intelligence in Schools" (Ranieri, Cuomo, Biagini, 2024), che permette agli studenti di esplorare in profondità dilemmi etici complessi legati all'intelligenza artificiale. Ad esempio, un dibattito potrebbe focalizzarsi su questioni come "L'IA dovrebbe avere diritti?", "Chi è responsabile per le decisioni prese dall'IA?", "Come dovremmo bilanciare innovazione e regolamentazione nell'ambito dell'IA?". Questa attività permette agli studenti di esplorare questioni etiche filosoficamente ricche e complesse, di sviluppare argomentazioni articolate e di confrontarsi con diverse prospettive.

Un'altra strategia efficace è lo studio di caso etico avanzato, che consente agli studenti di analizzare in profondità situazioni reali o verosimili che sollevano questioni etiche significative. Ad esempio, uno studio di caso potrebbe riguardare "l'algoritmo per l'ammissione scolastica", che analizza come un algoritmo progettato per semplificare il processo di selezione degli studenti in una scuola d'élite finisca per favorire studenti provenienti da famiglie più agiate, sollevando questioni di equità e giustizia. Questo caso di studio permette agli studenti di esplorare le implicazioni etiche dell'uso dell'IA in contesti di decisione critica, analizzando le cause dei bias, le conseguenze per diverse categorie di persone e le possibili strategie per mitigare questi problemi.

L'analisi etica dei documenti di riferimento rappresenta un'ulteriore strategia preziosa per la dimensione etica nella scuola secondaria di secondo grado. Questa strategia prevede l'analisi critica di documenti come linee guida etiche, regolamenti e dichiarazioni di principi sull'intelligenza artificiale prodotti da vari organismi e istituzioni. Ad esempio, gli studenti potrebbero analizzare le "Ethics Guidelines for Trustworthy AI" della Commissione Europea, esaminando i principi proposti, le loro implicazioni pratiche e le possibili tensioni o contraddizioni. Questa analisi permette agli studenti di comprendere il dibattito etico contemporaneo sull'intelligenza artificiale, di familiarizzare con i principi etici fondamentali e di sviluppare una visione critica e informata.

L'uso di progetti etici collaborativi rappresenta un'ulteriore strategia efficace per la dimensione etica nella scuola secondaria di secondo grado. Questa strategia prevede la collaborazione tra studenti per lo sviluppo di progetti che esplorano questioni etiche legate all'intelligenza artificiale. Ad esempio, un progetto potrebbe prevedere lo sviluppo di una "carta etica" per l'uso dell'IA in ambito educativo, che identifichi principi, valori e linee guida pratiche. Questo progetto richiederebbe agli studenti di ricercare, analizzare, discutere e sintetizzare informazioni e prospettive diverse, sviluppando non solo consapevolezza etica, ma anche competenze di ricerca, collaborazione e comunicazione.

La Tabella 2.7 sintetizza le principali strategie didattiche per la dimensione etica dell'AI literacy, articolate per livello scolastico.

Tabella 2.7. Strategie didattiche per la dimensione etica per livello scolastico

Livello scolastico	Strategie didattiche	Esempi di attività
Scuola primaria	- Storie e favole etiche - Circle time etico - Giochi di ruolo etici - Poster o disegni etici	- "Il robot giocattolo e la privacy": narrazione e discussione - Circle time su "È giusto che un robot prenda decisioni al posto nostro?" - Gioco di ruolo "Il robot insegnante" - Creazione di poster sulle "regole del buon comportamento" per i robot
Scuola secondaria di I grado	- Studio di caso etico semplificato - Dibattito etico strutturato - Analisi multi-prospettica - Media etici	- Studio di caso "Il docente robot" - Dibattito "È etico utilizzare sistemi di sorveglianza AI nelle scuole?" - Analisi dell'uso dell'IA per la personalizzazione da diverse prospettive - Visione e discussione di un documentario sull'IA nella sorveglianza

Scuola secondaria di II grado	- Dibattito etico avanzato - Studio di caso etico avanzato - Analisi etica dei documenti di riferimento - Progetti etici collaborativi	- Dibattito "L'IA dovrebbe avere diritti?" - Studio di caso "L'algoritmo per l'ammissione scolastica" - Analisi delle "Ethics Guidelines for Trustworthy AI" - Sviluppo di una "carta etica" per l'uso dell'IA in ambito educativo
--	--	--

Le strategie didattiche per la dimensione etica dell'AI literacy evolvono dunque in funzione del livello di sviluppo cognitivo, emotivo e morale degli studenti, passando da semplici narrazioni e discussioni nella scuola primaria, a complesse analisi critiche e progetti collaborativi nella scuola secondaria di secondo grado. Questa progressione permette di costruire gradualmente una consapevolezza etica solida e articolata, essenziale per un approccio responsabile e consapevole all'intelligenza artificiale.

In definitiva, un approccio efficace all'insegnamento dell'AI literacy non può che essere pluralistico e integrato dal punto di vista metodologico. Le strategie qui presentate – dalla P4C al coding unplugged, dal trial and error con piattaforme IA alla gamification, dall'IBL con IA generativa al problem solving sull'explainability, dal debate al case study – non sono alternative esclusive, ma strumenti complementari da orchestrare sapientemente all'interno di un percorso curricolare coerente. La loro efficacia risiede nella capacità di promuovere un apprendimento attivo, esperienziale, riflessivo, critico e collaborativo. È attraverso la combinazione di queste esperienze che gli studenti possono costruire progressivamente una comprensione profonda e multidimensionale dell'Intelligenza Artificiale, sviluppando non solo le conoscenze e le abilità necessarie per interagire con essa, ma anche e soprattutto la consapevolezza critica e la responsabilità etica indispensabili per orientarne lo sviluppo e l'utilizzo verso un futuro più equo, sostenibile e autenticamente umano. Il ruolo del docente si trasforma da trasmettitore di informazioni a facilitatore di processi, progettista di ambienti di apprendimento e modello di pensiero critico e riflessivo, accompagnando gli studenti in un percorso di scoperta e di crescita che li renda protagonisti consapevoli nell'era dell'Intelligenza Artificiale.

CAPITOLO 3: Valutare l'AI literacy: costruzione e validazione di uno strumento di misurazione

3.1 Fondamenti teorici e metodologici: cosa dice la ricerca

L'integrazione dell'Intelligenza Artificiale (IA) nei contesti educativi e professionali ha reso imperativa la necessità di valutare la competenza degli individui in questo ambito emergente. La misurazione dell'AI literacy rappresenta una sfida metodologica complessa, non solo per la natura multidimensionale del costrutto, ma anche per il suo rapido evolversi in risposta ai continui avanzamenti tecnologici. Come evidenziato nei capitoli precedenti, l'AI literacy si configura come un insieme articolato di conoscenze, abilità, attitudini e comportamenti che permettono agli individui di interagire consapevolmente con i sistemi di IA, comprenderne i meccanismi di funzionamento, valutarne criticamente gli output e considerarne le implicazioni etiche e sociali.

Il presente capitolo si propone di esplorare gli approcci metodologici alla misurazione dell'AI literacy, con particolare attenzione allo sviluppo e alla validazione di strumenti quantitativi. L'obiettivo è duplice: da un lato, offrire una panoramica critica degli strumenti attualmente disponibili, analizzandone caratteristiche, punti di forza e limitazioni; dall'altro, presentare in dettaglio il Questionario sull'AI literacy (CAILS – Critical Artificial Intelligence Literacy Scale) sviluppato nell'ambito di questa ricerca, illustrandone il processo di costruzione, le proprietà psicometriche e le potenzialità applicative.

La misurazione dell'AI literacy assume particolare rilevanza non solo per la valutazione dell'efficacia degli interventi formativi, come quello presentato nel capitolo successivo, ma anche per la comprensione dei livelli di preparazione delle diverse popolazioni rispetto alle sfide poste dall'IA. Disporre di strumenti di misurazione validi e affidabili consente di identificare lacune conoscitive, orientare interventi educativi mirati e monitorare l'evoluzione delle competenze nel tempo. Come sottolineato da Ng et al. (2021), la valutazione dell'AI literacy dovrebbe abbracciare dimensioni cognitive, operative, critiche ed etiche, riflettendo così la natura multiforme di questa competenza.

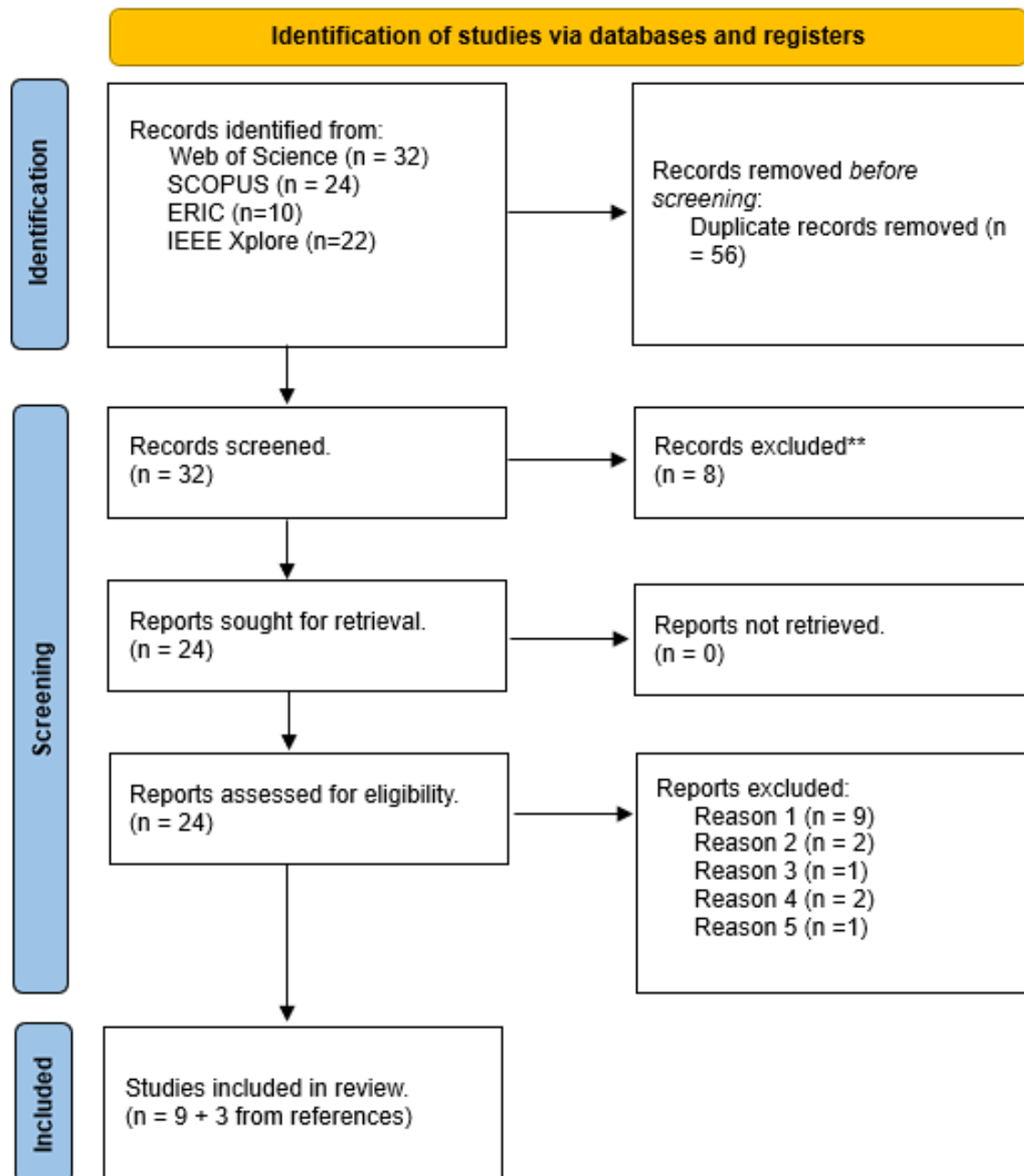
Questo capitolo si inserisce nel percorso di ricerca illustrato nei capitoli precedenti, fornendo il necessario collegamento metodologico tra il framework teorico dell'AI literacy e la sua applicazione empirica nel contesto educativo. La discussione si svilupperà a partire da una revisione sistematica della letteratura sugli strumenti di misurazione esistenti, per poi concentrarsi sul processo di sviluppo e validazione del questionario CAILS, il quale rappresenta uno dei contributi originali di questa tesi dottorale.

3.1.1 Revisione della letteratura sugli strumenti di misurazione dell'AI literacy

Per identificare e analizzare gli strumenti disponibili per la misurazione dell'AI literacy, è stata condotta una scoping review della letteratura, seguendo l'approccio metodologico delineato da Grant e Booth (2009) e aderendo al framework PRISMA (Moher et al., 2009). Sebbene non si tratti di una revisione sistematica in senso stretto, la scoping review condivide con essa i caratteri di trasparenza, riproducibilità e sistematicità, mirando a ridurre errori e bias e, quindi, a produrre risultati più affidabili (Cooper et al., 2019).

Il processo di revisione è stato articolato in cinque fasi distinte: a) definizione del problema e delle domande di ricerca, b) identificazione della letteratura rilevante, c) valutazione dei risultati della ricerca, d) categorizzazione, decodifica e sintesi dei risultati, e) disseminazione dei risultati della revisione. La ricerca è stata condotta su diverse banche dati accademiche, tra cui Web of Science

(Clarivate), ERIC (Institute of Education Sciences), IEEE Xplore (IEEE) e SCOPUS (Elsevier), nel marzo 2023 e successivamente aggiornata nel settembre 2023. È importante sottolineare che il volume di letteratura sulle scale di AI literacy è relativamente limitato, dato il carattere emergente di questo campo di ricerca.



**All records were excluded manually. No automation tools were used at this stage.

Figura 5

La strategia di ricerca si è concentrata su articoli che menzionavano esplicitamente "AI literacy scales or questionnaire" o "artificial intelligence literacy scales or questionnaire" nel titolo, nell'abstract o nel testo. L'applicazione di criteri di esclusione ha permesso di filtrare gli articoli non pertinenti. I criteri di esclusione sono stati i seguenti:

1. Opere che utilizzavano il termine "AI literacy scale" o "questionnaire" ma si concentravano su un argomento diverso sono state omesse (criterio 1).
2. Editoriali e libri non sono stati inclusi a causa della loro natura non sottoposta a peer-review (criterio 2).
3. Articoli che menzionavano "AI literacy scale" o "questionnaire" ma riguardavano l'applicazione dell'IA in campi specifici non correlati all'educazione sono stati esclusi (criterio 3).
4. Pubblicazioni non in inglese o italiano sono state escluse per evitare errori di interpretazione dovuti a barriere linguistiche (criterio 4).

La ricerca iniziale ha identificato 88 documenti (32 da Web of Science, 24 da SCOPUS, 10 da ERIC e 22 da IEEE Xplore). Dopo l'eliminazione dei duplicati (56) e l'applicazione dei criteri di esclusione, sono stati inclusi nella revisione 9 studi pertinenti.

La revisione della letteratura ha permesso di identificare nove strumenti distinti per la valutazione dell'AI literacy, ciascuno caratterizzato da approcci metodologici, dimensioni misurate e gruppi target specifici. Di seguito viene fornita una panoramica di questi strumenti, con particolare attenzione alle loro caratteristiche distintive, al processo di validazione e agli aspetti dell'IA che intendono misurare.

AILQ (Artificial Intelligence Literacy Questionnaire)

Sviluppato da Ng et al. (2023), l'AILQ è uno strumento completo progettato per valutare l'AI literacy degli studenti delle scuole secondarie. Lo strumento sottolinea la necessità di valutare l'AI literacy attraverso quattro dimensioni chiave: apprendimento affettivo (che include motivazione intrinseca e autoefficacia/fiducia), apprendimento comportamentale (che comprende impegno comportamentale e collaborazione), apprendimento cognitivo (che copre acquisizione, applicazione, valutazione e creazione di conoscenza) e apprendimento etico.

L'AILQ è stato progettato e validato attraverso un processo multifase che ha incluso revisione teorica, giudizio di esperti, interviste, uno studio pilota e analisi fattoriale confermativa. La validazione ha coinvolto uno studio pilota con 363 studenti delle scuole secondarie di Hong Kong per esaminare le proprietà psicometriche dello strumento, confermando una struttura a quattro fattori e dimostrando buona affidabilità e validità.

La distintività dell'AILQ risiede nella sua copertura completa del framework ABCE (Affective, Behavioural, Cognitive, and Ethical learning), rendendolo un contributo originale nel campo. Il questionario è progettato per catturare una vasta gamma di competenze, dalla comprensione di base alle abilità di pensiero di ordine superiore, e include elementi specificamente volti a valutare la fiducia, l'autoefficacia, la motivazione e la comprensione etica degli studenti in relazione all'IA.

SNAIL (Scale for the Assessment of Non-Experts' AI literacy)

Sviluppata da Laupichler et al. (2023), la SNAIL è uno strumento progettato per misurare l'AI literacy tra i non esperti, definiti come individui privi di formazione formale in IA o informatica. Lo strumento è emerso dalla necessità di fornire una valutazione valida e affidabile delle competenze in materia di IA.

Il processo di sviluppo della SNAIL ha comportato la distribuzione di un set iniziale di item sull'AI literacy attraverso un questionario online a un gruppo diversificato di non esperti. L'analisi

successiva, che ha sfruttato l'analisi fattoriale esplorativa, ha rivelato una struttura a tre fattori che incapsula le dimensioni latenti delle competenze in materia di IA: comprensione tecnica, valutazione critica e applicazione pratica. La versione finale della SNAIL comprende 31 item, progettati per valutare l'AI literacy di individui o gruppi e per valutare l'efficacia dei corsi di formazione sull'AI literacy.

MAILS (Meta AI literacy Scale)

La MAILS, sviluppata da Carolus et al. (2023), è stata progettata per offrire uno strumento di valutazione dell'AI literacy robusto e flessibile, profondamente radicato nella letteratura. Combina aspetti tradizionali dell'AI literacy con competenze psicologiche, rendendola applicabile in vari contesti professionali. La modularità della scala permette di utilizzare le sue sfaccettature in modo indipendente, adattando la valutazione a obiettivi e casi d'uso specifici.

Il target principale della MAILS sono gli adulti, con potenziale per un'applicazione più ampia data la sua fondazione in competenze universali e costrutti psicologici. La scala è stata sviluppata attraverso l'analisi di dati raccolti online da 300 partecipanti. Ha subito sia analisi fattoriale esplorativa che confermativa per confermarne la struttura fattoriale, assicurando che misurasse efficacemente l'AI literacy e le competenze psicologiche relative alla risoluzione di problemi, all'apprendimento e alla regolazione delle emozioni nel contesto dell'IA.

La progettazione della scala si basa sulla concettualizzazione dell'AI literacy di Ng et al. (2021), che include Use & Apply AI, Understand AI, Detect AI, AI Ethics, con l'aggiunta di Create AI, ed è arricchita con competenze psicologiche per affrontare i cambiamenti pervasivi portati dai sistemi di IA. Queste includono AI Self-efficacy nell'apprendimento e nella risoluzione di problemi, AI Self-management, ulteriormente suddiviso in AI Self-efficacy (che copre AI Problem-solving e AI Learning) e AI Self-competency (che include AI Persuasion literacy e AI Emotion regulation).

Scala di Pinsky e Benlian

La scala sviluppata da Pinsky e Benlian (2023) ha adottato un approccio completo, iniziando con una revisione sistematica della letteratura e interviste con esperti per definire e concettualizzare l'AI literacy generale. Questo lavoro ha portato alla creazione di un set iniziale di item volto a catturare l'essenza delle competenze sociotecniche degli individui nell'interazione con le tecnologie di IA (basate sulla ricerca IS e sulla collaborazione uomo-IA).

Per affinare e validare questo set di item, il team di ricerca ha condotto due round di card sorting che hanno coinvolto un totale di undici giudici, seguiti da un pre-test con 50 partecipanti. Questo processo è culminato in uno strumento di misurazione validato comprendente cinque dimensioni e 13 item, fornendo supporto empirico per il modello di misurazione.

Lo strumento chiarisce l'AI literacy generale attraverso sette dimensioni ristrutturare, categorizzate in conoscenza dell'attore IA (literacy esplicita), conoscenza delle fasi dell'IA (literacy esplicita) ed esperienza dell'IA (literacy tacita), fornendo così una comprensione sfumata delle competenze rilevanti per l'interazione uomo-IA. Strutturando le competenze uomo-IA in modo così dettagliato, lo strumento getta le basi per future direzioni di ricerca in IS, invitando indagini sulle relazioni sfaccettate tra l'AI literacy e i suoi effetti sui risultati organizzativi e tecnologici.

AILS (AI literacy Scale)

Sviluppata da Wang et al. (2022), l'AILS è una scala innovativa progettata per catturare l'essenza dell'AI literacy attraverso i costrutti centrali di consapevolezza, uso, valutazione ed etica, distillati in

uno strumento conciso di 12 item da un pool iniziale di 65 item attraverso un rigoroso processo di validazione del contenuto in tre fasi.

Identificando questi costrutti e validando la scala su due campioni di dati diversi, lo studio stabilisce l'AILS come uno strumento affidabile e valido per misurare quantitativamente l'AI literacy. Lo sviluppo dell'AILS è stato ispirato dal lavoro fondamentale sulla digital literacy di Calvani et al. (2009), suggerendo che un framework parallelo potrebbe essere applicato efficacemente all'AI literacy.

La scala ha subito un approfondito approccio metodologico in sei fasi, assicurando la sua affidabilità e validità, e ha confermato l'appropriatezza del modello a quattro costrutti per concettualizzare l'AI literacy. Nonostante alcuni costrutti abbiano valori di varianza estratta media (AVE) leggermente inferiori alla soglia abituale, la scala ha dimostrato sufficiente validità convergente, supportando la sua utilità nella ricerca e nella pratica.

I risultati dello studio sottolineano la significativa relazione tra l'AI literacy, misurata dall'AILS, e variabili chiave come la digital literacy, gli atteggiamenti verso i robot e l'uso quotidiano dell'IA, evidenziando la rilevanza predittiva della scala per comprendere le interazioni degli utenti con le tecnologie di IA. Degno di nota, lo studio propugna l'uso dell'AILS come strumento completo piuttosto che concentrarsi sui singoli costrutti, riflettendo la natura multiforme dell'AI literacy. Questo approccio si allinea con il riconoscimento che l'AI literacy va oltre la competenza nell'uso di applicazioni specifiche, enfatizzando la necessità di valutare la competenza generale degli utenti in materia di IA.

AIAS (AI Anxiety Scale)

L'AIAS, introdotta da Wang e Wang (2022), è uno strumento standardizzato progettato per misurare l'ansia del grande pubblico nei confronti dello sviluppo dell'IA, un fenomeno sempre più diffuso man mano che le tecnologie di IA diventano più integrate nella vita quotidiana.

Il processo di sviluppo dell'AIAS è stato completo, coinvolgendo la concettualizzazione del costrutto AIA, la generazione di item e una rigorosa validazione attraverso l'analisi dei dati di 301 rispondenti. Questo processo ha assicurato l'affidabilità e la validità dell'AIAS attraverso varie dimensioni, inclusa la validità correlata al criterio, di contenuto, discriminante, convergente e nomologica. La scala a 21 item finalizzata offre una comprensione dell'AIA, fornendo una metrica robusta per confrontare i livelli individuali di AIA con norme più ampie.

La base di campionamento diversificata dell'AIAS ne aumenta l'applicabilità per lo sviluppo di standard correlati e offre uno strumento preciso per valutare l'AIA in modo più accurato rispetto alle misure esistenti. Lo sviluppo dell'AIAS contribuisce alle discussioni teoriche e pratiche sull'AIA, enfatizzando la necessità di un'analisi multidimensionale e di interventi mirati per mitigare l'ansia. Per i professionisti, l'AIAS sottolinea l'importanza di comprendere varie dimensioni dell'ansia—come l'apprendimento, la sostituzione del lavoro, la cecità sociotecnica e la configurazione dell'IA—per adattare misure correttive. Per gli educatori, l'AIAS può guidare lo sviluppo di strategie didattiche che stimolano l'interesse per l'IA, riducendo così l'ansia e migliorando i risultati dell'apprendimento.

ATAI (Attitude Towards Artificial Intelligence) Scale

La scala ATAI (Sindermann et al., 2021) rappresenta un importante avanzamento nella misurazione degli atteggiamenti pubblici verso l'IA, affrontando la dicotomia di accettazione e paura associata alla proliferazione delle tecnologie di IA nella vita quotidiana.

Questo studio introduce la scala ATAI, elaborata e validata attraverso tre lingue: tedesco, cinese e inglese, con partecipanti da Germania, Cina e Regno Unito che forniscono una vasta prospettiva culturale sugli atteggiamenti verso l'IA. Lo sviluppo della scala ATAI è una dimostrazione della relazione che gli individui hanno con l'IA, incapsulando sia l'ottimismo di abbracciare i progressi tecnologici sia lo scetticismo riguardo le potenziali implicazioni sociali.

Comprendente cinque item, la scala ATAI delinea due dimensioni distinte ma interrelate: accettazione dell'IA e paura dell'IA, rivelando una struttura fattoriale equilibrata che risuona attraverso diversi contesti culturali. Questa biforcazione permette una comprensione sfumata degli atteggiamenti verso l'IA, riflettendo la complessità delle interazioni uomo-macchina nell'era digitale.

Il processo di validazione della scala, che coinvolge confronti sulla volontà di interagire con e utilizzare prodotti specifici di IA come auto a guida autonoma e robot sociali tra gli altri, sottolinea la sua rilevanza pratica e applicabilità. La validazione cross-culturale della scala ATAI ne significa la robustezza e universalità, offrendo una misura affidabile e concisa degli atteggiamenti verso l'IA che trascende le barriere geografiche e linguistiche.

Degno di nota, l'esplorazione da parte della scala della dimensione della paura, in particolare per quanto riguarda l'impatto dell'IA sull'occupazione, fornisce intuizioni critiche su ansie e percezioni prevalenti, nonostante le varie enfasi sulla perdita di lavoro tra diversi campioni. Questo evidenzia la sensibilità della scala alle sfumature contestuali e l'importanza di considerare caratteristiche specifiche del campione in future applicazioni di ricerca.

GAAIS (General Attitudes towards Artificial Intelligence Scale)

La GAAIS (Schepman & Rodway, 2020) è uno strumento progettato per misurare il sentimento pubblico verso l'intelligenza artificiale (IA) attraverso lenti sia positive che negative. Questa scala ha subito un'analisi fattoriale esplorativa, svelando due sottoscale distinte che incapsulano percezioni positive, riflettendo utilità sociale e personale, e percezioni negative, che evidenziano preoccupazioni circostanti l'IA.

La GAAIS ha dimostrato forti proprietà psicometriche, inclusi buoni indici di affidabilità così come validità convergente e discriminante quando giustapposta a misure esistenti. La GAAIS offre intuizioni sugli atteggiamenti del grande pubblico verso applicazioni specifiche di IA correlando comfort e capacità percepita con atteggiamenti generali.

I risultati rivelano una dicotomia: gli individui sono generalmente positivi riguardo l'uso dell'IA in aree che coinvolgono big data, come astronomia e farmacologia, apprezzandone l'utilità e l'impatto. Al contrario, applicazioni che richiedono giudizio umano, come trattamento medico e consulenza psicologica, suscitano reazioni più negative, guidate da preoccupazioni etiche e disagio.

Lo sviluppo e il processo di validazione di questa scala sottolineano la complessità degli atteggiamenti pubblici verso l'IA, rivelando che il comfort con le applicazioni di IA è un predittore più potente degli atteggiamenti generali rispetto alla capacità percepita. Questo suggerisce che, mentre le persone possono riconoscere i potenziali benefici dell'IA, il loro livello di comfort, probabilmente influenzato da considerazioni emotive ed etiche, gioca un ruolo cruciale nel plasmare i loro atteggiamenti complessivi.

3.1.2 Analisi comparativa degli strumenti

L'analisi comparativa dei nove strumenti identificati nella revisione della letteratura rivela diverse tendenze e approcci alla misurazione dell'AI literacy. La Tabella 3.1 fornisce una sintesi delle

caratteristiche principali di ciascuno strumento, inclusi il loro scopo generale, il target principale, il processo di validazione e le caratteristiche distintive.

Tabella 3.1: Sintesi degli strumenti per la misurazione dell'AI literacy.

Strumento	Autori	Scopo generale	Target principale	Processo di validazione	Caratteristiche distintive	N. item
AILQ	Ng et al., 2023	Valutare l'AI literacy degli studenti delle scuole secondarie attraverso quattro dimensioni	Studenti delle scuole secondarie	Studio pilota con 363 studenti a Hong Kong	Si concentra su apprendimento affettivo, comportamentale, cognitivo ed etico, validato in un contesto educativo	32
SNAIL	Laupichler et al., 2023	Valutare l'AI literacy dei non esperti attraverso comprensione tecnica, valutazione critica e applicazione pratica	Non esperti	Questionario online con analisi fattoriale esplorativa	Si rivolge a un pubblico ampio senza educazione formale in IA; modello a tre fattori	31
MAILS	Carolus et al., 2023	Misurare l'AI literacy integrando competenze fondamentali con abilità psicologiche	Adulti di lingua tedesca	Dati da 300 individui analizzati con EFA e CFA	Include aspetti psicologici come problem-solving e regolazione delle emozioni insieme alle competenze di IA	72 (60 + 12)
Scala di Pinsky & Benlian	Pinsky & Benlian, 2023	Misurare l'AI literacy generale nei contesti di interazione uomo-IA	Ricercatori e professionisti in IS e IA	Revisione della letteratura, interviste con esperti, card sorting con 11 giudici e pre-test con 50 partecipanti	Sette dimensioni che categorizzano conoscenza di AI literacy esplicita e tacita	13
AILS	Wang et al., 2022	Valutare l'AI literacy tra utenti ordinari, concentrandosi su	utenti ordinari	Validazione del contenuto in tre fasi, sondaggio	Misura breve e validata che correla l'AI	12

		consapevolezza, uso, valutazione ed etica		con due campioni di dati che portano a uno strumento a 12 item	literacy con la digital literacy e gli atteggiamenti verso i robot	
AIAS	Wang & Wang, 2022	Misurare l'ansia verso lo sviluppo dell'IA	Pubblico generale	Dati da 301 rispondenti analizzati per varie forme di validità	Introduce il concetto di ansia da IA (AIA) e la sua misurazione; processo di validazione esteso	21
ATAI Scale	Sindermann et al., 2021	Valutare gli atteggiamenti verso l'IA attraverso le dimensioni di accettazione e paura	Popolazioni da Germania, Cina e Regno Unito	I partecipanti hanno completato la scala ATAI; struttura fattoriale testata attraverso campioni	Validazione multilingue; due fattori associati negativamente	5
GAAIS	Schepman & Rodway, 2020	Sviluppare una Scala di Atteggiamenti Generali verso l'Intelligenza Artificiale con sottoscale positive e negative	Pubblico generale	Analisi Fattoriale Esplorativa con validazione incrociata sugli atteggiamenti verso specifiche applicazioni di IA	Sottoscale positive e negative che catturano emozioni e utilità versus preoccupazioni; buone proprietà psicometriche	21

Dall'analisi comparativa emergono diverse considerazioni:

Diversità dimensionale: Gli strumenti variano significativamente nel numero di dimensioni che misurano. Mentre alcuni, come l'AILQ di Ng et al. (2023) adottano un approccio multidimensionale che copre aspetti cognitivi, operativi, critici ed etici, altri, come l'ATAI di Sindermann et al. (2021), si concentrano su dimensioni più specifiche come l'accettazione e la paura nei confronti dell'IA.

Lunghezza degli strumenti: Il numero di item varia considerevolmente, da un minimo di 5 (ATAI Scale) a un massimo di 72 (MAILS). Questa variabilità riflette diverse filosofie nella progettazione

degli strumenti: alcuni privilegiano la brevità e la facilità di somministrazione, altri optano per una maggiore profondità e dettaglio nella misurazione.

Target di popolazione: Gli strumenti sono stati sviluppati per diverse popolazioni target, dagli studenti delle scuole secondarie (AILQ) ai non esperti (SNAIL), fino al pubblico generale (AIAS, GAAIS). Questa diversità riflette la necessità di strumenti adattati a specifici contesti e livelli di competenza.

Processi di validazione: Tutti gli strumenti hanno subito processi di validazione, ma con metodologie e rigore variabili. Alcuni, come l'AILS, hanno utilizzato sia analisi fattoriale esplorativa che confermativa, mentre altri hanno impiegato approcci più qualitativi o focalizzati sulla validità di contenuto.

Focus psicologico vs. tecnico: Alcuni strumenti, come MAILES e AIAS, pongono particolare enfasi sugli aspetti psicologici dell'interazione con l'IA, come l'ansia, l'autoefficacia e la regolazione emotiva, mentre altri, come SNAIL e AILS, si concentrano maggiormente sulle competenze tecniche e operative.

Dimensione etica: La dimensione etica emerge come un componente importante in molti degli strumenti analizzati (AILQ, AILS), riflettendo la crescente consapevolezza delle implicazioni etiche dell'IA e la necessità di incorporare questa dimensione nella misurazione dell'AI literacy.

Questa analisi comparativa fornisce un quadro completo del panorama attuale degli strumenti di misurazione dell'AI literacy, evidenziando sia le convergenze che le divergenze negli approcci. La diversità degli strumenti disponibili sottolinea la natura multiforme dell'AI literacy e la necessità di approcci differenziati in base al contesto e agli obiettivi specifici della misurazione.

3.1.3 Un modello concettuale per la costruzione di un dispositivo valutativo

Il Questionario sull'AI literacy (CAILS) sviluppato nell'ambito della presente ricerca si fonda sul framework quadridimensionale per l'AI literacy proposto nei capitoli precedenti (Cuomo et al., 2022), che concettualizza questa competenza come un costrutto complesso articolato in quattro dimensioni fondamentali: conoscitiva, operativa, critica ed etica. Questo framework, a sua volta, trae ispirazione dall'approccio multidimensionale alla competenza digitale elaborato da Calvani et al. (2009), adattandolo alle specificità dell'Intelligenza Artificiale.

Non ci soffermeremo sui dettagli del framework che è già stato illustrato in dettaglio nella sezione 1.3, ma è importante sottolineare che questo approccio quadridimensionale fornisce una base concettuale solida per la progettazione di un questionario che miri a valutare l'AI literacy in modo completo e articolato, andando oltre la semplice conoscenza tecnica per abbracciare aspetti critici, valutativi ed etici. Tale visione olistica risulta particolarmente rilevante nel contesto educativo, dove la formazione all'IA dovrebbe mirare non solo all'acquisizione di nozioni e abilità operative, ma anche allo sviluppo di un approccio critico e di una sensibilità etica verso queste tecnologie.

3.2 Metodologia di sviluppo e validazione dello strumento valutativo

Lo sviluppo del questionario ha seguito un processo rigoroso e strutturato, basato sulle raccomandazioni metodologiche di DeVellis (2016) per la costruzione di scale di misurazione psicometriche. L'iter di sviluppo ha compreso diverse fasi sequenziali, ciascuna finalizzata a garantire la validità e l'affidabilità dello strumento finale.

La prima fase del processo ha comportato la definizione chiara dei costrutti che il questionario intendeva misurare. Partendo dal framework quadridimensionale dell'AI literacy, sono stati identificati i concetti chiave e gli elementi costitutivi di ciascuna dimensione attraverso un'approfondita revisione della letteratura. Questa fase ha incluso l'analisi di opere seminali, tra cui quelle di Floridi (2018, 2021), Ng et al. (2021), Selwyn (2022), e di fonti autorevoli come la Commissione Europea (2018, 2019, 2020, 2021), il Joint Research Center (2018), l'Organizzazione per la Cooperazione e lo Sviluppo Economico (2018a, 2018b, 2021), l'UNESCO (2019a, 2019b, 2021) e l'UNICEF (2020, 2021a, 2021b).

3.2.1 Generazione degli item

La generazione degli item ha rappresentato una fase cruciale del processo di sviluppo del questionario. Per garantire una copertura completa e bilanciata delle quattro dimensioni dell'AI literacy, sono stati elaborati descrittori più analitici per ciascuna dimensione, identificando sottodimensioni e componenti specifici.

Parallelamente, è stata condotta una revisione di questionari già validati su temi correlati, come competenza tecnologica o literacy digitale, al fine di selezionare item che potessero essere adattati per misurare l'AI literacy.

Gli elementi concettuali che ricorrevano in almeno due fonti indipendenti sono stati trasformati in item. È stata posta particolare attenzione a garantire che il questionario coprisse una gamma completa di dimensioni dell'AI literacy, mantenendo al contempo chiarezza e rilevanza. La scala iniziale sviluppata comprendeva 22 item focalizzati sulla dimensione conoscitiva, 32 sulla dimensione operativa, 30 sulla dimensione critica e 34 sulla dimensione etica. La Tabella 3.2 presenta alcuni item di esempio per chiarire l'output finale della fase di generazione degli item.

Tabella 3.2: Risultati iniziali della generazione degli item.

Dimensione del framework	Descrizione	Domanda di esempio	Opzione della matrice	N. di item	Riferimenti
Dimensione conoscitiva	Sapere come utilizzare le applicazioni di IA e il loro funzionamento fondamentale	Quando si tratta di IA, sento che la mia conoscenza sull'argomento sarebbe:	Conoscere e comprendere le definizioni di IA e i fondamenti teorici	22	Cuomo et al., 2022; Ng et al., 2021; UNESCO, 2019a, 2019b, 2021
			Conoscere e comprendere le funzioni matematiche di base alla base dell'algoritmo		
Dimensione operativa	Utilizzare concetti, competenze e	Secondo te, le seguenti attività potrebbero	Supporto ai servizi di emergenza	32	Cuomo et al., 2022; Ng et al., 2021; Wang &

	applicazioni di IA in vari contesti	essere supportate dall'IA?			Wang; JRC, 2018; OECD, 2018a, 2018b; UNESCO, 2019
			Reportistica di notizie		
			Supporto emotivo		
Dimensione critica	Applicazioni di IA per abilità di pensiero critico (come valutare, stimare, prevedere e progettare)	Quanto sei d'accordo con le seguenti affermazioni?	I sistemi artificialmente intelligenti commettono molti errori	30	Schepman & Rodway, 2022; Selwyn, 2022; Wang et al., 2022; OECD, 2019; UNICEF, 2020
			Un agente artificialmente intelligente sarebbe migliore di un dipendente in molti lavori di routine		
Dimensione etica	Fattori centrati sull'uomo (come giustizia, responsabilità, apertura, etica e sicurezza)	Quanto credi che le seguenti considerazioni influenzino l'affidabilità dell'IA?	Impatto sociale: il rischio che l'IA concentrerà ulteriormente potere e ricchezza nelle mani di pochi	34	Commissione Europea 2018, 2019, 2020, 2021; Floridi 2018; Floridi et al., 2021; JRC & OECD, 2021; Sindermann et al., 2021; UNICEF, 2021a, 2021b; UNESCO, 2021
			Impatto democratico: l'impatto delle tecnologie IA sulle democrazie		
			Impatto sul lavoro: Impatto		

			dell'IA sul mercato del lavoro e come potrebbero essere influenzati diversi gruppi demografici		
--	--	--	--	--	--

3.2.2 Revisione da parte di esperti e validità di facciata

La validità di facciata è cruciale perché valuta se un questionario misura effettivamente ciò che intende misurare. Comporta la revisione del questionario e la determinazione se gli item e la loro formulazione sembrano rilevanti e appropriati per misurare il costrutto di interesse, in questo caso l'AI literacy.

Per garantire la validità di facciata del questionario, è stato coinvolto un panel di esperti (N=5) nel campo dell'IA e della valutazione educativa. È importante notare che l'uso di un piccolo gruppo di esperti per valutare la validità di contenuto è stato considerato appropriato in questo studio, poiché si concentrava su un compito cognitivo che non richiedeva una comprensione approfondita del fenomeno in esame (Anderson & Gerbing, 1991; Hinkin, 1998; Schriesheim et al., 1993). Questi esperti avevano una profonda conoscenza dell'AI literacy e una comprensione approfondita dei costrutti previsti dal questionario.

Per garantire una comprensione condivisa dei quattro costrutti dell'AI literacy, le definizioni sono state condivise con ciascun esperto. Il processo di validazione del contenuto è consistito nei seguenti passaggi:

Il panel di esperti ha attentamente esaminato ogni item e ha fornito preziosi spunti e suggerimenti per migliorarli. Hanno evidenziato eventuali item che sembravano poco chiari, ridondanti o irrilevanti per il costrutto misurato.

Agli esperti è stato inizialmente chiesto di categorizzare ciascun item in una delle quattro dimensioni del framework di AI literacy (conoscitiva, operativa, critica, etica) seguendo la metodologia sostenuta da Schriesheim et al. (1993). Se almeno quattro su cinque esperti assegnavano la stessa classificazione a un item, questo era considerato come chiaramente indirizzato a un concetto. Dei 118 item totali, 15 avevano classificazioni errate o non classificate da due esperti, mentre altri 23 avevano errori di classificazione o non erano classificati da più esperti. Di conseguenza, questi elementi non sono stati inclusi nello studio.

Gli item sono stati poi migliorati, compresa la loro formulazione e il formato, in base ai suggerimenti degli esperti; 14 item sono stati riformulati e 20 item, relativi all'impatto dell'IA nell'educazione, sono stati spostati fuori dal corpus principale del questionario diventando un'appendice che può essere utilizzata in contesto educativo come sezione di informazioni più ampia.

Questo processo ha lasciato 60 item che sono stati ulteriormente migliorati, compresa la loro formulazione e formato, in base ai suggerimenti degli esperti.

Il nostro questionario segue i consigli di Likert (1932) e Hinkin (1998) utilizzando una scala Likert a 5 punti. Una scala Likert a cinque punti è stata considerata più adatta perché il nostro questionario

sarebbe stato somministrato online. Il questionario è stato creato in modo da poter essere presentato elettronicamente su computer o cellulari, consentendo una facile trasmissione e distribuzione tramite Internet.

3.2.3 Somministrazione pilota e validazione

Lo studio effettivo è stato condotto online, nel maggio 2023, tramite lo strumento di indagine "Qualtrics", mentre tutte le analisi sono state implementate utilizzando il software statistico R (2021). Il questionario è stato somministrato a un campione di convenienza, consistente in studenti del primo anno (2023) di Scienze della Formazione Primaria dell'Università di Firenze. Il campione, dopo aver rimosso i dati mancanti, era composto da 191 studenti di Scienze della Formazione Primaria, tra cui 178 femmine (93,19%) e 11 maschi (5,76%). Le fasce d'età erano comprese tra 18-24 (60,21%) e 55-64 (0,52%), mentre il più alto grado di istruzione completato era il diploma di scuola superiore per 128 intervistati (67,55%) e una laurea universitaria triennale per 37 intervistati (19,15%). La Tabella 3.3 riassume le caratteristiche del campione.

Tabella 3.3: Caratteristiche del campione

Caratteristica	Voci	%	Frequenza
Genere	Maschio	5.76%	11
	Femmina	93.19%	178
	Preferisco non dirlo	1.05%	2
Età	18 - 24	60.21%	115
	25 - 34	26.18%	50
	35 - 44	8.90%	17
	45 - 54	4.19%	8
	55 - 64	0.52%	1
Livello scolastico di impiego (se già lavora)	Prima Infanzia	22.68%	22
	Elementare	69.07%	67
	Scuola superiore	3.09%	3
	Classi speciali (Sostegno)	5.15%	5
Più alto grado o livello di istruzione completato	Scuola superiore	67.55%	128
	Università 3 anni	19.15%	37
	Università 5 anni	8.51%	17
	Master di 1° livello	3.19%	6

	Dottorato	0.52%	1
	Master di 2° livello	1.06%	2
Esperienza professionale (in anni):	< 5 Anni	77.32%	75
	> 5 < 10 Anni	17.53%	17
	>10 < 20 Anni	3.09%	3
	> 20 Anni	2.06%	2

3.2.4 Proprietà psicometriche del questionario

L'affidabilità di un questionario può essere considerata come la consistenza dei risultati del sondaggio. Poiché l'errore di misurazione è presente nel campionamento del contenuto, nei cambiamenti dei rispondenti e nelle differenze tra valutatori, la consistenza di un questionario può essere valutata utilizzando la sua consistenza interna. La consistenza interna è una misura della correlazione tra gli item del questionario e quindi della consistenza nella misurazione del costrutto previsto. Nel presente studio, per valutare l'affidabilità e la validità del questionario sono stati utilizzati quattro indici: l'alpha di Cronbach, l'omega di McDonald (1999), l'affidabilità composita (CR) e la varianza media estratta (AVE).

La consistenza interna è comunemente stimata utilizzando il coefficiente alpha (Cronbach, 1951), noto anche come alpha di Cronbach. Secondo i suggerimenti degli esperti, ci si aspetta che il valore dell'alpha di Cronbach sia almeno 0.70 per indicare un'adeguata consistenza interna di un dato questionario (DeVellis, 2016; Nunnally, 1978). Un valore basso (inferiore a 0.7) dell'alpha di Cronbach per un dato questionario rappresenta una scarsa consistenza interna e, quindi, una scarsa interrelazione tra gli item.

I risultati sono mostrati nella Tabella 3.4. Il punteggio dell'alpha di Cronbach del questionario era 0.953, mentre i punteggi per ciascuno dei quattro costrutti erano, rispettivamente, 0.880, 0.908, 0.941, e 0.914. Sebbene le affidabilità di ciascun singolo costrutto fossero superiori a 0.70, lo strumento nel suo complesso ha ottenuto un punteggio superiore, indicando che quest'ultimo è più affidabile dei singoli costrutti.

La validità convergente della scala è stata valutata utilizzando i criteri CR e AVE stabiliti da Fornell e Larcker (1981). L'alpha di Cronbach è una misura più soggettiva di affidabilità rispetto al CR, e valori di CR di 0.70 e superiori sono considerati soddisfacenti (Hair et al., 1998). L'AVE confronta la varianza raccolta da un costrutto con la varianza causata dall'errore di misurazione. Secondo Hair et al. (2010), valori superiori a 0.5 mostrano convergenza soddisfacente; nella nostra scala, i valori CR erano superiori a 0.7, e i valori AVE erano superiori a 0.5, il che indicava una convergenza accettabile.

Tabella 3.4: Risultati dell'alpha di Cronbach, omega di McDonald, AVE e CR

Dimensioni del framework	Alpha di Cronbach	Omega di McDonald	Varianza media estratta	Affidabilità Composita	N. di elementi

Dimensione conoscitiva	0.880	0.888	0.521	0.916	8
Dimensione operativa	0.908	0.910	0.513	0.926	12
Dimensione critica	0.941	0.950	0.522	0.915	10
Dimensione etica	0.914	0.924	0.520	0.914	10
Totale	0.953	0.956	0.531	0.940	40

La struttura fondamentale della misura a 60 item è stata ulteriormente confermata dall'analisi fattoriale esplorativa (EFA). I carichi delle componenti o dei fattori indicano quali fattori sono misurati dalle domande. Le domande che misurano gli stessi indicatori dovrebbero caricare sugli stessi fattori. I carichi dei fattori variano da -1.0 a 1.0. La struttura fattoriale della scala del questionario è stata investigata mediante analisi delle componenti principali (PCA) che indicavano una struttura a quattro componenti come ipotizzato dal framework. Le quattro componenti sono state ruotate utilizzando una tecnica di rotazione ortogonale (rotazione varimax) per consentire correlazioni tra i componenti. Secondo i risultati della PCA, i quattro variabili con autovalori maggiori di 1.00 erano responsabili del 69.68% della varianza totale estratta.

Questo studio ha seguito le cinque regole che sono frequentemente utilizzate come criteri per decidere se mantenere o eliminare gli item: (1) valori maggiori del criterio della radice fondamentale (autovalore >1.00); (2) carichi fattoriali insignificanti (0.50); (3) carichi fattoriali significativi su più fattori; (4) almeno tre indicatori o item in un singolo fattore; e (5) fattori a item singolo (Anderson & Gerbing, 1991; Hair et al., 2010; Straub, 1989). Alla fine, 40 item sono emersi dai 60 item con 8 item focalizzati sulla dimensione conoscitiva dell'IA, 12 sulla dimensione operativa dell'IA, 10 sulla dimensione critica dell'IA e 10 item sulla dimensione etica dell'IA. I risultati dell'EFA sono mostrati nella Tabella 3.5.

I controlli di assunzione per il modello a quattro fattori finale hanno risultato in un test di sfericità di Bartlett significativo $\chi^2 = 2375$, $df = 528$, $p < .001$, mostrando una matrice di correlazione praticabile che deviava significativamente da una matrice identità. La Misura di Adeguatezza del Campionamento di Kaiser-Meyer-Olkin (KMO MSA) complessiva era 0.835, indicando un campionamento ampiamente sufficiente. Durante l'analisi fattoriale confermativa (CFA) il modello con i 40 item caricati sui quattro fattori come descritto, è emerso come accettabile con un CFI = 0.959, TLI = 0.950, RMSEA = 0.041, SRMR = 0.05 (Tabella 3.6).

Tabella 3.5: Risultati dell'analisi fattoriale esplorativa.

	Carichi dei Fattori			
	Dimensione conoscitiva	Dimensione operativa	Dimensione critica	Dimensione etica
KW1	0.782			

KW2	0.680			
KW3	0.809			
KW4	0.876			
KW5	0.571			
KW6	0.857			
KW7	0.738			
KW8	0.740			
OP1		0.665		
OP2		0.713		
OP3		0.608		
OP4		0.759		
OP5		0.824		
OP6		0.663		
OP7		0.668		
OP8		0.679		
OP9		0.804		
OP10		0.671		
OP11		0.752		
OP12		0.759		
CR1			0.748	
CR2			0.824	
CR3			0.713	
CR4			0.617	
CR5			0.680	
CR6			0.729	
CR7			0.607	
CR8			0.678	
CR9			0.857	
CR10			0.729	

ET1				0.566
ET2				0.547
ET3				0.720
ET4				0.681
ET5				0.621
ET6				0.790
ET7				0.763
ET8				0.836
ET9				0.824
ET10				0.789

Nota: I valori assoluti inferiori a 0.5 sono stati soppressi.

Tabella 3.6: Statistiche di adattamento del modello

				RMSEA 90% CI	
CFI	TLI	SRMR	RMSEA	Inferiore	Superiore
0.959	0.950	0.0538	0.0411	0.0211	0.0573

Nota: Il modello a quattro fattori è il modello teorico. CFI = indice comparativo di adattamento; TLI = indice di Tucker-Lewis; RMSEA = root mean square error of approximation; SRMR = standardised root mean square residual. TLI, CFI maggiori di .900 e valori RMSEA inferiori a .050 suggeriscono un adattamento adeguato del modello.

Il Questionario sull'AI literacy (CAILS) nella sua versione finale si compone di 40 item, distribuiti equamente tra le quattro dimensioni del framework teorico: 8 item per la dimensione conoscitiva, 12 per la dimensione operativa, 10 per la dimensione critica e 10 per la dimensione etica. Ogni item è valutato su una scala Likert a cinque punti, che va da "Fortemente in disaccordo" a "Fortemente d'accordo".

La dimensione conoscitiva del questionario valuta la comprensione concettuale dell'IA, compresi i suoi tipi, principi, applicazioni e limiti. Gli item richiedono al rispondente di valutare la propria comprensione di concetti di base come la distinzione tra IA e altre tecnologie digitali, la definizione di IA e la sua evoluzione storica. Esempi di item includono "Conosco la differenza tra IA e altre tecnologie digitali" e "Sono consapevole dei diversi tipi di applicazioni di IA esistenti oggi".

La dimensione operativa si concentra sulla capacità di applicare la conoscenza dell'IA in contesti pratici. Questa dimensione valuta la familiarità del rispondente con le tecniche, gli strumenti e i metodi utilizzati nell'IA, così come la capacità di identificare quando e come utilizzare l'IA per risolvere problemi specifici. Esempi di item sono "So che l'IA apprende dai dati che le vengono forniti" e "Capisco come i dati di addestramento influenzano la qualità dei modelli di IA".

La dimensione critica esplora la capacità del rispondente di valutare criticamente i sistemi di IA, i loro limiti e le loro implicazioni. Gli item misurano la consapevolezza dei bias algoritmici, la

trasparenza, l'interpretabilità e l'impatto sociale dell'IA. Esempi includono "Sono consapevole che l'IA può perpetuare pregiudizi presenti nei dati di addestramento" e "Valuto criticamente le informazioni fornite da sistemi di IA".

Infine, la dimensione etica valuta la comprensione e la sensibilità verso le questioni etiche associate all'IA. Gli item si concentrano su temi come la privacy, la responsabilità, la trasparenza e le implicazioni sociali ed etiche più ampie dell'IA. Esempi di item per questa dimensione sono "Ritengo importante che le decisioni prese dall'IA siano trasparenti e spiegabili" e "Considero le implicazioni etiche quando utilizzo o sviluppo applicazioni di IA".

La Tabella 3.7 presenta esempi di item per ciascuna dimensione, illustrando la varietà e la profondità dei temi affrontati nel questionario.

Tabella 3.7: Esempi di item per ciascuna dimensione del questionario

Dimensione	Codice item	Testo dell'item
Conoscitiva	KW1	Conosco la differenza tra IA e altre tecnologie digitali
Conoscitiva	KW4	Ho una comprensione di base di come funzionano i sistemi di IA
Operativa	OP2	So che l'IA apprende dai dati che le vengono forniti
Operativa	OP7	Capisco come i dati di addestramento influenzano la qualità dei modelli di IA
Critica	CR3	Sono consapevole che l'IA può perpetuare pregiudizi presenti nei dati di addestramento
Critica	CR8	Valuto criticamente le informazioni fornite da sistemi di IA
Etica	ET1	Ritengo importante che le decisioni prese dall'IA siano trasparenti e spiegabili
Etica	ET6	Considero le implicazioni etiche quando utilizzo o sviluppo applicazioni di IA

Il questionario finale è stato progettato per essere somministrato in un formato digitale, utilizzando piattaforme come Qualtrics, che offrono vantaggi in termini di facilità di somministrazione, riduzione degli errori di inserimento dei dati e gestione efficiente delle risposte. La somministrazione digitale consente inoltre di randomizzare l'ordine degli item, riducendo potenziali bias di risposta associati all'ordine di presentazione.

3.2.5 Applicazione del questionario sull'AI literacy nella sperimentazione

Il Questionario sull'AI literacy sviluppato e validato in questo studio è stato utilizzato come strumento centrale nella sperimentazione descritta nel capitolo successivo. In particolare, il questionario è stato somministrato sia all'inizio che alla fine dell'intervento formativo, permettendo così di valutare l'impatto del curriculum sull'AI literacy dei partecipanti.

La somministrazione pre-intervento è servita a stabilire una baseline delle conoscenze, attitudini e percezioni iniziali dei partecipanti riguardo all'IA. Questo ha fornito una fotografia dettagliata del loro livello di partenza rispetto alle quattro dimensioni dell'AI literacy, permettendo di identificare aree di forza e di debolezza. I risultati di questa valutazione iniziale hanno contribuito anche a informare l'approccio didattico, consentendo ai formatori di adattare alcuni aspetti del curriculum alle esigenze specifiche del gruppo.

La somministrazione post-intervento, identica in termini di contenuto a quella pre-intervento, ha permesso di misurare i cambiamenti avvenuti nelle quattro dimensioni dell'AI literacy a seguito della partecipazione al curriculum formativo. Il confronto tra i risultati pre e post ha fornito evidenze quantitative dell'efficacia dell'intervento, come sarà discusso in dettaglio nel Capitolo 5.

Questa applicazione del questionario nella sperimentazione ha dimostrato la sua utilità pratica come strumento di valutazione in contesti educativi. In particolare, la struttura quadridimensionale del questionario si è dimostrata particolarmente adatta per catturare la complessità dell'apprendimento relativo all'IA, che non si limita all'acquisizione di conoscenze tecniche ma include anche lo sviluppo di abilità operative, pensiero critico e consapevolezza etica.

Inoltre, il questionario ha dimostrato di essere sensibile ai cambiamenti nelle percezioni e nelle competenze auto-riportate dei partecipanti, rendendo possibile una valutazione dettagliata dell'impatto dell'intervento formativo su ciascuna delle quattro dimensioni dell'AI literacy. Questo aspetto è particolarmente importante nel contesto dell'educazione all'IA, dove gli obiettivi di apprendimento sono spesso complessi e multidimensionali.

Posizionando il Questionario sull'AI literacy (CAILS) sviluppato in questo studio nel panorama degli strumenti esistenti, emergono diversi elementi distintivi che ne caratterizzano l'approccio e la portata.

Un primo elemento di distinzione riguarda la struttura multidimensionale del questionario. Mentre alcuni strumenti, come l'ATAI di Sindermann et al. (2021) o il GAAIS di Schepman e Rodway (2020), si concentrano principalmente sugli atteggiamenti verso l'IA, il CAILS adotta un approccio più completo, abbracciando quattro dimensioni fondamentali: conoscitiva, operativa, critica ed etica. Questa struttura quadridimensionale, derivata dal framework teorico discusso nei capitoli precedenti, permette una valutazione più articolata e completa dell'AI literacy rispetto a strumenti focalizzati su aspetti più specifici.

Dal punto di vista della lunghezza e della praticità, il CAILS con i suoi 40 item si colloca in una posizione intermedia. È più esteso di strumenti come l'AILS di Wang et al. (2022) o l'ATAI di Sindermann et al. (2021), che con 12 e 5 item rispettivamente offrono una valutazione più rapida ma necessariamente meno dettagliata. D'altra parte, è più contenuto rispetto al MAILS di Carolus et al. (2023), che con 72 item fornisce una valutazione estremamente approfondita ma richiede un tempo di somministrazione considerevole. Questa posizione intermedia rappresenta un compromesso efficace tra profondità di valutazione e praticità di somministrazione.

Rispetto al target di popolazione, il CAILS è stato sviluppato e validato con un campione di studenti universitari in formazione per diventare insegnanti, distinguendosi da strumenti come l'AILQ di Ng et al. (2023), mirato agli studenti delle scuole secondarie, o la SNAIL di Laupichler et al. (2023), focalizzata sui non esperti in generale. Questa specificità del target offre vantaggi in termini di rilevanza per il contesto educativo, pur mantenendo una struttura e un contenuto potenzialmente applicabili a un pubblico più ampio.

Dal punto di vista della validazione psicometrica, il CAILS ha seguito un processo rigoroso che include sia l'Analisi Fattoriale Esplorativa (EFA) che l'Analisi Fattoriale Confermativa (CFA), dimostrando eccellenti proprietà psicometriche ($\alpha = 0.953$, AVE = 0.531, CR = 0.940). Questo livello di validazione è paragonabile a quello di strumenti come l'AILS di Wang et al. (2022) e superiore a quello di strumenti come la SNAIL di Laupichler et al. (2023), che ha utilizzato solo l'EFA.

Infine, per quanto riguarda il focus tematico, il CAILS si distingue per l'attenzione equilibrata alle quattro dimensioni dell'AI literacy, con un numero simile di item dedicati a ciascuna dimensione. Questo contrasta con strumenti come l'AIAS di Wang e Wang (2022), che si concentra specificamente sull'ansia verso l'IA, o la scala di Pinsky e Benlian (2023), che enfatizza le competenze sociotecniche nell'interazione uomo-IA. L'approccio bilanciato dello strumento riflette la visione dell'AI literacy come un costrutto multiforme in cui le diverse dimensioni hanno un peso comparabile.

In sintesi, il questionario sull'AI literacy sviluppato in questo studio si posiziona come uno strumento completo, psicometricamente robusto e bilanciato nelle dimensioni valutate. Mentre altri strumenti possono offrire vantaggi specifici in termini di brevità o focus su aspetti particolari, il CAILS fornisce una valutazione olistica dell'AI literacy, coerente con il framework teorico quadridimensionale proposto nei capitoli precedenti.

3.2.6 Limitazioni e direzioni future

Nonostante i risultati promettenti ottenuti nello sviluppo e nella validazione del Questionario sull'AI literacy, è importante riconoscere alcune limitazioni che caratterizzano questo strumento e che possono orientare futuri sviluppi.

Una prima limitazione riguarda la rappresentatività del campione utilizzato per la validazione. Il questionario è stato validato con un campione di convenienza costituito da studenti universitari di Scienze della Formazione Primaria, prevalentemente di sesso femminile (93,19%) e in maggioranza nella fascia d'età 18-24 anni (60,21%). Questo limita la generalizzabilità dei risultati ad altre popolazioni, come professionisti in servizio, studenti di altre discipline o individui con diverse caratteristiche demografiche. Future ricerche potrebbero estendere la validazione dello strumento a campioni più diversificati, per verificarne l'applicabilità in contesti differenti.

Una seconda limitazione riguarda la natura auto-riportata delle misurazioni. Come tutti i questionari di auto-valutazione, il CAILS si basa sulla percezione soggettiva del rispondente riguardo alle proprie conoscenze, abilità e attitudini. Questo approccio, sebbene ampiamente utilizzato e validato in ambito psicometrico, può essere soggetto a bias di desiderabilità sociale o a distorsioni dovute a limiti nella capacità di auto-valutazione (Kruger & Dunning, 1999). Studi futuri potrebbero integrare le misurazioni auto-riportate con valutazioni più oggettive, come test di conoscenza o compiti di performance, per ottenere una triangolazione dei risultati.

Una terza limitazione concerne la dimensione temporale. Il questionario fornisce una misurazione puntuale dell'AI literacy, ma non cattura la sua evoluzione nel tempo in risposta ai rapidi progressi tecnologici nel campo dell'IA. Data la natura dinamica e in rapida evoluzione dell'IA, lo strumento potrebbe necessitare di aggiornamenti periodici per rimanere rilevante e accurato. Sarebbe utile sviluppare protocolli per la revisione e l'aggiornamento regolare del questionario, in linea con gli sviluppi tecnologici e le mutevoli implicazioni sociali dell'IA.

Una quarta limitazione riguarda la granularità della misurazione. Sebbene il questionario copra quattro dimensioni principali dell'AI literacy, potrebbe non catturare adeguatamente la complessità e la ricchezza di alcune sottodimensioni o aspetti specifici. Ad esempio, nella dimensione etica,

potrebbero essere necessari strumenti più dettagliati per esplorare a fondo temi come la privacy, l'equità algoritmica o l'impatto sociale dell'IA. Ricerche future potrebbero sviluppare versioni estese o moduli complementari del questionario, focalizzati su aspetti specifici dell'AI literacy per contesti o scopi particolari.

Infine, il questionario attuale non include misurazioni esplicite di fattori contestuali e ambientali che possono influenzare lo sviluppo dell'AI literacy, come l'accesso alle tecnologie, il supporto istituzionale o le politiche educative. L'integrazione di questi fattori in un framework di valutazione più ampio potrebbe fornire una comprensione più completa delle condizioni che facilitano o ostacolano lo sviluppo dell'AI literacy.

Guardando al futuro, diverse direzioni di ricerca sembrano promettenti. Una prima direzione riguarda lo sviluppo di strumenti di valutazione adattati a specifici contesti professionali o educativi. L'AI literacy può assumere forme e rilevanza diverse a seconda del contesto, e strumenti specializzati potrebbero offrire misurazioni più precise e pertinenti.

Una seconda direzione promettente concerne lo sviluppo di approcci di misurazione multimodali, che integrino questionari auto-riportati con altre forme di valutazione, come l'osservazione diretta, l'analisi di prodotti o performance, o l'utilizzo di dati comportamentali. Questo approccio potrebbe fornire una visione più completa e sfaccettata dell'AI literacy.

Una terza direzione riguarda l'esplorazione delle relazioni tra AI literacy e altri costrutti rilevanti, come il pensiero critico, la creatività, la competenza digitale generale o l'attitudine al problem solving. Comprendere come l'AI literacy si collochi nel più ampio panorama delle competenze del XXI secolo potrebbe offrire preziose intuizioni per la progettazione educativa e la formazione professionale.

Infine, sarebbe prezioso condurre studi longitudinali sull'evoluzione dell'AI literacy nel tempo, sia a livello individuale che sociale. Tali studi potrebbero illuminare i processi attraverso cui l'AI literacy si sviluppa, i fattori che ne influenzano la traiettoria e le implicazioni a lungo termine per individui e comunità.

Concludendo, il presente capitolo ha delineato un percorso metodologico articolato per la misurazione dell'AI literacy, con particolare attenzione allo sviluppo e alla validazione del Questionario sull'AI literacy (CAILS). L'approccio adottato ha integrato una revisione sistematica della letteratura sugli strumenti esistenti con un processo rigoroso di costruzione e validazione di un nuovo strumento, basato sul framework quadridimensionale dell'AI literacy proposto nei capitoli precedenti.

La revisione della letteratura ha permesso di identificare nove strumenti distinti per la valutazione dell'AI literacy, ciascuno con caratteristiche, punti di forza e limitazioni specifiche. Questa analisi ha rivelato una significativa eterogeneità negli approcci, sia in termini di dimensioni misurate che di target di popolazione, evidenziando la complessità e la multidimensionalità del costrutto. Al contempo, ha fornito un solido fondamento per lo sviluppo di un nuovo strumento che potesse integrare le prospettive emergenti in un approccio coerente e comprensivo.

Il processo di sviluppo del Questionario sull'AI literacy ha seguito una metodologia rigorosa, articolata in diverse fasi: dalla definizione concettuale dei costrutti alla generazione di un ampio pool di item, dalla revisione da parte di esperti alla definizione del formato di misurazione, fino alla somministrazione pilota e all'analisi delle proprietà psicometriche. Tale processo ha portato alla validazione di uno strumento a 40 item, distribuiti equamente tra le quattro dimensioni dell'AI literacy: conoscitiva, operativa, critica ed etica.

Le analisi psicometriche hanno confermato l'eccellente affidabilità e validità del questionario, con un alpha di Cronbach di 0.953 per lo strumento nel suo complesso e valori superiori a 0.88 per ciascuna delle quattro sottoscale. L'analisi fattoriale, sia esplorativa che confermativa, ha supportato la struttura a quattro fattori ipotizzata dal framework teorico, confermando la solidità concettuale dell'approccio quadridimensionale all'AI literacy.

Il confronto con altri strumenti di misurazione ha evidenziato il posizionamento distintivo del CAILS nel panorama esistente. Rispetto ad altri strumenti, il CAILS offre un approccio più bilanciato e comprensivo, coprendo in modo equo le diverse dimensioni dell'AI literacy. Pur riconoscendo le limitazioni intrinseche a qualsiasi strumento di misurazione, questo strumento rappresenta un contributo significativo al campo, offrendo un mezzo valido e affidabile per valutare l'AI literacy in diversi contesti educativi e professionali.

L'applicazione del questionario nella sperimentazione, descritta nel capitolo successivo, ha dimostrato la sua utilità pratica nella valutazione dell'impatto di interventi formativi sull'AI literacy. La capacità dello strumento di misurare cambiamenti nelle diverse dimensioni dell'AI literacy lo rende particolarmente adatto per ricerche longitudinali e per la valutazione di programmi educativi.

In conclusione, il Questionario sull'AI literacy sviluppato in questo studio si configura come uno strumento robusto e versatile per la misurazione dell'AI literacy, contribuendo in modo significativo alla comprensione e alla promozione di questa competenza fondamentale nell'era digitale. Le direzioni future delineate suggeriscono percorsi promettenti per l'ulteriore sviluppo e affinamento di approcci alla misurazione dell'AI literacy, in un campo in rapida evoluzione che richiede strumenti sempre più sofisticati e adattabili.

CAPITOLO 4 La sperimentazione del curriculum per l'AI: il contesto e la metodologia

Questo capitolo è dedicato alla descrizione dettagliata e approfondita del disegno di ricerca adottato per la sperimentazione del curriculum sull'AI literacy, che costituisce il nucleo empirico fondante di questa tesi di dottorato. Come sottolineato da Creswell e Creswell (2018), la chiarezza e il rigore nella presentazione del contesto e della metodologia rappresentano presupposti indispensabili per garantire la validità scientifica, la trasparenza e la potenziale replicabilità dello studio. In accordo con i principi del rigore metodologico nella ricerca educativa (Cohen et al., 2018), verranno illustrati con dovizia di particolari il contesto specifico – accademico e formativo – in cui si è svolta l'indagine, le caratteristiche del campione di partecipanti coinvolti, le fasi operative che hanno scandito la sperimentazione, nonché la struttura, i contenuti e gli obiettivi formativi del percorso didattico implementato. Successivamente, nella seconda parte del capitolo (Sezione 4.2), verranno esplicitate le strategie di ricerca impiegate, dettagliando gli strumenti utilizzati per la raccolta dei dati e le procedure di analisi adottate per valutare l'impatto del curriculum e rispondere alle domande di ricerca che guidano questo lavoro.

4.1 Il contesto della sperimentazione

La definizione accurata del contesto è essenziale per comprendere la portata e i limiti dei risultati ottenuti. La sperimentazione si è radicata in un ambiente formativo specifico, con partecipanti dotati di caratteristiche peculiari e all'interno di un percorso didattico strutturato secondo precise logiche teoriche e pedagogiche. Di seguito verranno specificati questi aspetti fondamentali.

La sperimentazione ha trovato la sua collocazione naturale all'interno del Laboratorio di Tecnologie Didattiche, un'attività formativa curricolare obbligatoria per le studentesse e gli studenti iscritti al Corso di Laurea Magistrale a ciclo unico in Scienze della Formazione Primaria dell'Università degli Studi di Firenze, afferente al Dipartimento di Formazione, Lingue, Intercultura, Letterature e Psicologia (FORLILPSI). La scelta di questo specifico contesto non è stata casuale, ma dettata da una duplice motivazione strategica.

In primo luogo, il laboratorio rappresentava un contenitore didattico istituzionalmente riconosciuto e dedicato all'esplorazione del rapporto tra pedagogia e tecnologia, offrendo quindi un ambiente accademicamente pertinente e strutturato per introdurre un tema innovativo e complesso come l'AI literacy. In secondo luogo, e di importanza cruciale per gli obiettivi di questo dottorato, la sperimentazione all'interno di questo laboratorio ha costituito l'implementazione concreta del progetto di ricerca dottorale focalizzato proprio sullo sviluppo e la validazione di un curriculum per l'AI literacy destinato ai futuri insegnanti della scuola dell'infanzia e primaria. Il laboratorio ha quindi rappresentato il terreno elettivo per mettere alla prova il framework teorico e le proposte didattiche elaborate nel corso del dottorato (cfr. Capitolo 1 e 2).

Il percorso formativo si è svolto interamente in presenza, presso le aule attrezzate del Dipartimento FORLILPSI, nel periodo compreso tra aprile 2024 e maggio 2024. Si è trattato di un percorso intensivo, articolato su più incontri settimanali, per una durata complessiva di 36 ore di attività d'aula. Questa concentrazione temporale è stata pensata per favorire un'immersione profonda nei temi trattati e per mantenere un elevato livello di coinvolgimento e continuità nel lavoro dei partecipanti. Le attività si sono avvalse delle risorse tecnologiche disponibili (es. LIM, proiettore) e dei dispositivi personali dei partecipanti (PC portatili), che sono stati invitati a portare e utilizzare attivamente durante le sessioni laboratoriali per ricerche, attività interattive e lavori di gruppo.

4.1.1 Obiettivi formativi

Gli obiettivi formativi del curriculum sperimentato sono stati definiti in modo analitico per ciascuna delle quattro dimensioni dell'AI literacy, oltre a un insieme di obiettivi trasversali/didattici. Questi obiettivi hanno guidato sia la progettazione delle attività formative che la valutazione del loro impatto sui partecipanti. Si riportano di seguito gli obiettivi specifici per ciascuna dimensione, che costituiscono il quadro di riferimento per l'intero percorso.

Obiettivi per la dimensione conoscitiva:

Comprendere e saper spiegare in modo essenziale cos'è l'Intelligenza Artificiale, distinguendola da altre tecnologie digitali e collocandola nel suo contesto storico.

Delineare le tappe fondamentali dell'evoluzione storica dell'IA, riconoscendone le principali "stagioni" di sviluppo (es. "primavere" ed "inverni") e i momenti di svolta concettuale e tecnologica.

Analizzare criticamente e confrontare diverse definizioni di IA (proposte da pionieri, istituzioni, filosofi), cogliendone le diverse enfasi (es. razionalità, umanità, azione) e le implicazioni sottostanti.

Riflettere criticamente sul concetto di intelligenza, articolando le differenze fondamentali tra le capacità cognitive umane (comprensione, coscienza, emozioni, creatività) e le capacità simulate dalle macchine attuali.

Riconoscere e argomentare la rilevanza pedagogica dell'introduzione dell'AI literacy nel curriculum della scuola dell'infanzia e primaria.

Obiettivi per la dimensione operativa:

Definire e spiegare il concetto di algoritmo come sequenza finita, ordinata e non ambigua di istruzioni per risolvere un problema o raggiungere un obiettivo.

Illustrare il ruolo fondamentale e la natura dei dati (quantità, qualità, rappresentatività) come input essenziale per l'addestramento dei sistemi di Machine Learning.

Descrivere i principi base dell'apprendimento automatico supervisionato (es. classificazione) e non supervisionato (es. clustering) attraverso esempi concreti e accessibili.

Comprendere e saper applicare i principi fondamentali del pensiero computazionale (scomposizione di problemi, riconoscimento di pattern, astrazione, progettazione di algoritmi) alla risoluzione di compiti semplici.

Sperimentare e valutare l'efficacia didattica di attività di coding unplugged come strumento per introdurre il pensiero computazionale e i concetti algoritmici ai bambini.

Obiettivi per la dimensione critica:

Identificare e spiegare il concetto di bias algoritmico, riconoscendone le possibili origini (dati di addestramento distorti, scelte di progettazione, pregiudizi incorporati) e le potenziali conseguenze discriminatorie.

Analizzare criticamente esempi di output generati da sistemi IA (in particolare testi e immagini), valutandone l'affidabilità, la coerenza, la plausibilità e la potenziale presenza di stereotipi, errori fattuali o distorsioni.

Applicare strategie di pensiero critico (es. porre domande generative, verificare informazioni da fonti multiple, considerare prospettive alternative, identificare assunzioni implicite) alla valutazione di informazioni e contenuti prodotti o mediati dall'IA.

Riconoscere e discutere le sfide legate alla trasparenza e all'interpretabilità di alcuni modelli IA complessi (il problema della "black box").

Obiettivi per la dimensione etica:

Identificare e discutere i principali dilemmi etici sollevati dall'impiego crescente dell'IA in diversi ambiti della vita sociale (es. lavoro, salute, giustizia, relazioni interpersonali, decisioni automatizzate).

Riconoscere e applicare principi etici fondamentali (come equità, non maleficenza, autonomia, giustizia, trasparenza, responsabilità, rispetto della privacy) all'analisi di scenari concreti riguardanti lo sviluppo e l'uso dell'IA.

Comprendere le problematiche specifiche relative alla raccolta massiva di dati personali, al consenso informato, alla profilazione algoritmica e all'uso responsabile dei dati nell'era dell'IA.

Sviluppare la capacità di argomentare diverse posizioni etiche in modo ragionato e rispettoso all'interno di un dibattito strutturato su questioni etiche controverse legate all'IA.

Riflettere sulla necessità di una governance etica dell'IA (attraverso policy, linee guida, regolamentazioni) e sul ruolo cruciale dell'educazione nel promuovere un approccio socialmente responsabile alla tecnologia.

Obiettivi trasversali/didattici:

Dimostrare la capacità di progettare un'attività didattica strutturata, coerente e significativa sull'AI literacy, mirata a specifici obiettivi di apprendimento relativi ad una o più dimensioni del framework e adeguata al contesto della scuola dell'infanzia o primaria.

Dimostrare la capacità di selezionare, giustificare e ipotizzare l'utilizzo appropriato di metodologie didattiche attive, partecipative e riflessive (es. P4C, debate, coding unplugged, analisi critica di casi, storytelling collaborativo con IA) per insegnare l'AI literacy.

Sviluppare una maggiore consapevolezza del proprio ruolo di future educatrici nell'accompagnare i bambini e le bambine verso un rapporto critico, etico, consapevole e creativo con l'Intelligenza Artificiale e le tecnologie digitali.

Il raggiungimento di questi obiettivi formativi, che mirano a costruire una competenza complessa e multidimensionale, è stato l'oggetto principale della valutazione condotta attraverso gli strumenti di ricerca descritti nella sezione successiva (4.2), con un focus particolare sull'analisi longitudinale dei processi di apprendimento e riflessione documentati nei diari.

4.1.2 Campione della sperimentazione: futuri insegnanti di fronte all'IA

Il gruppo di partecipanti alla sperimentazione era costituito da studenti e studentesse regolarmente iscritti/e al Corso di Laurea Magistrale in Scienze della Formazione Primaria. La partecipazione attiva ha coinvolto 65 discenti, un numero significativo che ha permesso di raccogliere una quantità considerevole di dati per un'analisi in profondità, sebbene l'analisi quantitativa dei questionari e quella

qualitativa dei diari si basino su un numero leggermente inferiore (per gli aspetti quantitativi, dei 65 iscritti, 52 hanno completato l'input iniziale, ma solo 44 hanno completato sia il pre che il post-test, costituendo il campione finale per l'analisi inferenziale mentre per i diari qualitativi N variabile tra 60 e 62), a causa di normali fluttuazioni nella presenza e nella consegna degli elaborati. In accordo con quanto osservato da Onwuegbuzie e Collins (2007) sulla rappresentatività dei campioni nella ricerca educativa, il fatto che si trattasse di un'attività laboratoriale obbligatoria per il piano di studi conferisce al campione una validità ecologica significativa, permettendo di considerarlo rappresentativo della popolazione studentesca di quel corso per l'anno accademico di riferimento, piuttosto che un gruppo auto-selezionato sulla base di un interesse preesistente per le tecnologie o l'IA.

L'analisi sociodemografica condotta sui dati raccolti tramite il questionario iniziale (N=52) permette di delineare un profilo dettagliato del campione, fondamentale per interpretare le traiettorie di apprendimento e per contestualizzare adeguatamente i risultati ottenuti. Come osservato da Creswell e Plano Clark (2011), la comprensione approfondita delle caratteristiche dei partecipanti costituisce un elemento imprescindibile per l'interpretazione dei dati nei disegni di ricerca misti.

Dal punto di vista demografico, il campione riflette la netta prevalenza femminile tipica dei corsi di formazione per insegnanti (96,15% donne, 1,92% uomini, 1,92% non binario/terzo genere), confermando il fenomeno di femminilizzazione della professione docente ampiamente documentato in letteratura (Drudy, 2008). La distribuzione per età evidenzia una netta maggioranza collocata nella fascia 18-24 anni (71,15%), corrispondente al percorso universitario standard, accompagnata da una significativa presenza di studenti nella fascia 25-34 anni (23,08%) e, in misura minore, oltre i 35 anni (5,77% complessivamente). Questa stratificazione anagrafica riflette la crescente eterogeneità dei percorsi formativi per l'insegnamento, caratterizzati da un aumento dell'accesso alla professione come seconda carriera (Bruinsma & Jansen, 2010).

Per quanto concerne l'esperienza professionale, circa la metà dei rispondenti (51,92%) dichiarava di avere già esperienze lavorative nel campo dell'istruzione, quasi esclusivamente nelle scuole dell'infanzia (22,22%) o primaria (74,07%) del contesto toscano. I ruoli ricoperti presentavano una discreta varietà (docente prevalente, potenziamento, sostegno, lezioni private), sebbene l'esperienza fosse per la quasi totalità recente (inferiore ai 5 anni per il 92,59%), coerentemente con l'età media del campione. Questa caratteristica risulta particolarmente significativa nell'ottica di una ricerca orientata alla trasferibilità delle competenze nel contesto professionale, poiché evidenzia come molte partecipanti stessero già confrontandosi con la realtà scolastica mentre erano ancora in formazione. Tale sovrapposizione tra formazione e pratica professionale, come evidenziato da Darling-Hammond (2017), può costituire un potente catalizzatore per l'apprendimento situato e per l'integrazione critica delle innovazioni nel proprio repertorio didattico.

Il background formativo dei partecipanti rivelava che la grande maggioranza (90,38%) aveva come titolo di studio più alto il Diploma di scuola superiore, essendo iscritti a un corso magistrale a ciclo unico, mentre il 5,77% possedeva già una laurea triennale e il 3,85% una laurea magistrale. Un aspetto particolarmente rilevante per gli scopi della ricerca riguardava la familiarità auto-percepita con la tecnologia, che presentava una distribuzione polarizzata: un gruppo consistente si dichiarava non particolarmente appassionato o interessato ("Probabilmente no", "Sicuramente no"), mentre un altro gruppo si definiva interessato o decisamente appassionato ("Probabilmente sì", "Decisamente sì"). Questa eterogeneità nell'atteggiamento verso le tecnologie educative, in linea con quanto riscontrato in studi analoghi (Ertmer et al., 2012), ha rappresentato un elemento di particolare interesse per la

valutazione dell'efficacia del curriculum proposto su partecipanti con diversi livelli di predisposizione iniziale.

Nonostante questa variabilità nell'interesse dichiarato, l'analisi delle pratiche digitali quotidiane mostrava un quadro di alfabetizzazione digitale di base piuttosto omogeneo: l'uso quotidiano di motori di ricerca e social media era quasi universale, così come l'uso frequente di servizi di streaming e acquisti online, mentre meno diffuso risultava l'uso di giochi su smartphone. Questo profilo rispecchia i pattern di utilizzo tipici della generazione di "digital natives" (Prensky, 2001), con una familiarità primariamente orientata al consumo di contenuti digitali piuttosto che alla produzione o all'utilizzo critico degli stessi.

La specificità di questo campione – future insegnanti della scuola dell'infanzia e primaria, prevalentemente giovani donne, con diversi livelli di esperienza lavorativa e di interesse per la tecnologia – riveste un'importanza cruciale nella prospettiva della ricerca educativa sull'innovazione tecnologica. Come evidenziato da Mishra e Koehler (2006) nel loro framework TPACK (Technological Pedagogical Content Knowledge), gli insegnanti rappresentano i mediatori fondamentali dell'integrazione della tecnologia nei contesti educativi. Studiare lo sviluppo della loro AI literacy significa indagare come un gruppo professionale strategico per il futuro educativo si appropria di competenze e consapevolezze necessarie per preparare le nuove generazioni a vivere e agire in una società sempre più permeata dall'IA. Le loro prospettive, le loro eventuali resistenze, le loro intuizioni pedagogiche costituiscono un materiale di analisi privilegiato per comprendere i processi di trasformazione concettuale che accompagnano l'introduzione di queste tecnologie nel pensiero e nella pratica educativa.

In conformità con i principi etici della ricerca educativa (British Educational Research Association, 2018), le procedure per la raccolta dati hanno incluso l'ottenimento del consenso informato da tutti i partecipanti, con esplicita garanzia dell'anonimato e della riservatezza delle risposte in fase di analisi e presentazione dei risultati.

La Tabella 4.1 riassume le caratteristiche demografiche e professionali del campione.

Tabella 4.1: Caratteristiche demografiche dei partecipanti (N=52)

Dato	Valore
Genere	96,15% F, 1,92% M, 1,92% non binario
Età 18-24	71,15%
Età 25-34	23,08%
Età >35	5,77%
Lavora nell'istruzione	51,92%
Di cui: Infanzia	22,22%
Di cui: Primaria	74,07%
Esperienza <5 anni	92,59%
Diploma	90,38%
Laurea triennale	5,77%
Laurea magistrale	3,85%

4.1.3 Fasi della sperimentazione

La sperimentazione del curriculum per l'AI literacy si è sviluppata secondo un rigoroso cronoprogramma, strutturato per garantire un'adeguata progressione degli apprendimenti e una raccolta sistematica dei dati. In accordo con i principi del design-based research (Barab & Squire, 2004; Design-Based Research Collective, 2003), il percorso è stato articolato in fasi sequenziali e interconnesse, ciascuna caratterizzata da specifici obiettivi formativi e di ricerca, con un'attenzione particolare alla coerenza interna e alla validità ecologica dell'intervento.

La sperimentazione si è sviluppata in cinque macro-fasi, implementate nel periodo aprile-maggio 2024, per un totale di 36 ore di attività formativa. Queste fasi hanno rappresentato un'applicazione sistematica del framework quadridimensionale per l'AI literacy precedentemente elaborato (Cuomo et al., 2022).

Fase iniziale: valutazione e orientamento (19 aprile 2024) Questa fase ha avuto una funzione propedeutica sia sul piano formativo che su quello della ricerca. Sul versante formativo, sono state poste le basi concettuali dell'intero percorso attraverso la presentazione del framework dell'AI literacy e delle sue quattro dimensioni, accompagnata da un'esplorazione guidata di risorse didattiche specifiche per la scuola primaria. L'attivazione delle preconoscenze è avvenuta attraverso la visione di un video divulgativo della Royal Society sull'onnipresenza dell'IA, seguita da un lavoro a coppie e una discussione collettiva facilitata dalla piattaforma interattiva Wooclap sulle esperienze quotidiane, spesso inconsapevoli, con l'IA (es. assistenti vocali, raccomandazioni online, sblocco facciale). Questa strategia, in linea con i principi costruttivisti dell'apprendimento (Jonassen, 1999), ha permesso di far emergere le rappresentazioni spontanee e di introdurre la rilevanza pedagogica del tema. Un'attività di brainstorming strutturato su Wooclap ha ulteriormente esplorato i "Pro" e "Contro" percepiti dell'insegnamento dell'IA nella scuola primaria.

Sul versante della ricerca, è stata effettuata la somministrazione del questionario pre-intervento, strumento fondamentale per stabilire una baseline delle conoscenze, attitudini e percezioni iniziali dei partecipanti. Come sottolineato da Cobb et al. (2003), questa fase di raccolta dati preliminare costituisce un elemento imprescindibile nei disegni di ricerca educativa orientati alla valutazione dell'efficacia di curricula innovativi. Accanto alla raccolta dati strutturata, le attività interattive hanno generato dati contestuali significativi sulle conoscenze tacite, preconcezioni e atteggiamenti che non sempre si manifestano attraverso strumenti quantitativi.

Un'attività particolarmente significativa in questa fase introduttiva è stata l'analisi guidata di una risorsa educativa esistente per la scuola primaria (il corso "AI for Oceans" di Code.org), che ha consentito ai partecipanti di valutare, tramite una scheda strutturata, la copertura delle quattro dimensioni dell'AI literacy e discutere collettivamente i risultati dell'analisi.

Fase seconda: dimensione conoscitiva (23-26-30 aprile 2024) Questa fase ha focalizzato l'attenzione sulla dimensione conoscitiva dell'AI literacy, dedicandovi circa 9 ore complessive di attività formativa. In questa fase, il curriculum ha proposto un'immersione nella storia dell'IA, dalle intuizioni pionieristiche di Alan Turing fino agli sviluppi contemporanei dell'IA generativa, accompagnata da un'analisi critica delle diverse concettualizzazioni e definizioni proposte nel tempo da studiosi e istituzioni.

Attraverso lezioni interattive, supportate da slide con riferimenti iconografici e cronologici, è stata ripercorsa la storia dell'IA, partendo dalle radici antiche del desiderio umano di creare automi

(mitologia greca con Talos, automi meccanici come Il Turco, riferimenti letterari come Frankenstein) fino alle tappe fondamentali dell'informatica moderna. Particolare enfasi è stata data al contributo pionieristico di Alan Turing, spiegando il concetto teorico di Macchina di Turing come modello universale di computazione e il significato del celebre "Imitation Game" (Test di Turing). Sono state illustrate le diverse "stagioni" dell'IA, evidenziando l'alternanza tra periodi di grande fervore ("primavere") e fasi di rallentamento ("inverni"), fino all'attuale "età d'oro" dell'IA generativa.

Seguendo l'approccio socio-costruttivista all'apprendimento (Vygotsky, 1978), accanto ai momenti di esposizione teorica sono state proposte attività di elaborazione collettiva dei concetti, con particolare enfasi sull'adattamento della Philosophy for Children (P4C) come metodologia per facilitare discussioni significative sull'IA nel contesto educativo. Lipman (2003) ha evidenziato come questa metodologia sia particolarmente efficace nell'affrontare tematiche complesse e potenzialmente divisive, poiché crea uno spazio di indagine comunitaria in cui diverse prospettive possono essere esplorate in modo rispettoso e costruttivo.

Un'attività didattica particolarmente rilevante è stata la progettazione, in gruppi, di una sessione di P4C sul tema dell'IA per una classe di scuola primaria, completa di obiettivi, fasi, domande guida e strategie di facilitazione, utilizzando una scheda strutturata con presentazione successiva per un feedback tra pari e da parte dei docenti. Al termine di questa fase è stato somministrato il primo diario riflessivo.

Fase terza: dimensione operativa (3-7 maggio 2024) La terza fase, a cui sono state dedicate circa 9 ore complessive, ha spostato l'attenzione dai "cosa" ai "come", esplorando i meccanismi di funzionamento dell'IA. In questa fase sono stati affrontati temi fondamentali come algoritmi, pensiero computazionale, machine learning e reti neurali, con un'attenzione particolare al ruolo cruciale dei dati.

In linea con i principi dell'apprendimento esperienziale (Kolb, 1984), il curriculum ha previsto numerose attività pratiche, tra cui esperienze di coding unplugged e di classificazione algoritmica, progettate per rendere tangibili e comprensibili concetti altrimenti astratti. Come evidenziato da Papert (1980), l'apprendimento di concetti computazionali risulta particolarmente efficace quando avviene attraverso esperienze concrete di manipolazione e sperimentazione.

Le partecipanti hanno sperimentato algoritmi semplici attraverso giochi (come "il gioco dei fiammiferi"), attività di coding unplugged (pixel art, body coding, sorting networks, pair programming), e la creazione di alberi decisionali per comprendere i principi della classificazione algoritmica. Particolarmente efficace dal punto di vista didattico è stata l'attività di "classificazione della pasta", ispirata alle metodologie di apprendimento basate sul problem-based learning (Savery, 2006), dove i gruppi hanno dovuto sviluppare algoritmi per distinguere diversi formati di pasta, confrontando i risultati ottenuti e riflettendo metacognitivamente sulle difficoltà incontrate e sulle strategie utilizzate.

Per l'introduzione dei principi del Machine Learning (ML), si è adottato un approccio comparativo, distinguendo l'approccio simbolico da quello non simbolico/connessionista. Attraverso l'esempio visivo della classificazione dei topi e un'attività pratica di classificazione, si è illustrato l'apprendimento supervisionato, il ruolo delle features e l'importanza dei dataset. È stato poi introdotto il modello della rete neurale artificiale, spiegando il funzionamento del neurone artificiale (Perceptron di Rosenblatt, 1957) e il meccanismo di apprendimento tramite retropropagazione. Al termine di questa fase è stato somministrato il secondo diario riflessivo.

Fase quarta: dimensione critica (10-14 maggio 2024) La quarta fase, sviluppata in circa 6 ore di attività, ha esplorato la dimensione critica dell'AI literacy, focalizzandosi sui bias algoritmici, sulle problematiche di equità e trasparenza, e sull'applicazione del pensiero critico alla valutazione dell'IA e dei suoi prodotti.

In linea con l'approccio della critical media literacy (Kellner & Share, 2007), questa fase ha promosso lo sviluppo di competenze analitiche necessarie per valutare criticamente le tecnologie e i loro output, andando oltre la semplice alfabetizzazione funzionale. Il modulo è stato avviato con un video provocatorio sulla realtà sintetica e discussioni sulle fake news, per introdurre il pensiero critico (secondo la definizione di Ennis, 2011) e analizzare i bias cognitivi umani (euristiche, bias di conferma, stereotipi, pregiudizi) che influenzano il nostro giudizio.

Particolare attenzione è stata dedicata all'analisi dei bias, con attività pratiche utilizzando Teachable Machine di Google per sperimentare concretamente gli effetti dei dataset sbilanciati, e all'analisi critica di contenuti generati dall'IA, per sviluppare competenze di valutazione e discernimento. Si è poi approfondita l'IA generativa, spiegando i processi di digitalizzazione (es. pixel art) e i meccanismi base dei modelli generativi (GANs - con l'analogia del falsario/discriminatore; LLM come ChatGPT).

Un'attività particolarmente rilevante dal punto di vista pedagogico è stata la valutazione comparativa di testi e immagini, alcune generate da esseri umani e altre dall'IA, attraverso un test di riconoscimento "AI Art" che stimolava una riflessione metacognitiva sulle caratteristiche distintive e sui criteri di discernimento che possono essere sviluppati e applicati nel contesto educativo. A questa è stata affiancata un'attività di storytelling collaborativo ("Monte dei bigliettiini") utilizzando ChatGPT per generare una storia e confrontando criticamente l'output con la rielaborazione umana. Al termine di questa fase è stato somministrato il terzo diario riflessivo.

Fase quinta: dimensione etica (17-21-24-28-31 maggio 2024) L'ultima fase, articolata su circa 6 ore di attività, ha affrontato la dimensione etica dell'AI literacy, esplorando i dilemmi morali sollevati dallo sviluppo e dall'impiego dell'IA. In questa fase, il curriculum ha adottato il debate strutturato come metodologia principale, in linea con l'approccio dell'educazione morale attraverso la deliberazione collettiva proposto da Kohlberg (1984).

Attraverso la preparazione e la realizzazione di dibattiti su temi controversi legati all'IA, i partecipanti hanno avuto l'opportunità di sviluppare competenze argomentative e di esercitare il giudizio etico in situazioni complesse, caratterizzate da valori e interessi potenzialmente in conflitto. Seguendo le fasi canoniche del debate (claim, gruppi pro/contro, ricerca, argomentazione, esposizione, sintesi), le partecipanti si sono confrontate su dilemmi etici specifici, utilizzando anche la piattaforma Moral Machine del MIT per esplorare dilemmi concreti relativi ai veicoli autonomi.

Questa fase ha esplorato le principali aree di dibattito etico: responsabilità, non maleficenza, privacy, trasparenza (vs opacità), equità (vs bias), impatto sul lavoro, sorveglianza, copyright, e impatto ambientale (consumo energetico, bluewashing etico - Floridi, 2022). Sono state presentate le Leggi della Robotica di Asimov come esempio di framework etico e i principi chiave per un'IA affidabile (legalità, etica, robustezza) promossi dalla Commissione Europea. Particolare attenzione è stata dedicata al concetto di "Human-in-the-loop" come principio di supervisione etica.

Al termine di questa fase è stato somministrato il quarto diario riflessivo e il questionario post-intervento, seguito da attività di restituzione finale mediante WooClap, che hanno permesso di raccogliere feedback e consolidare la visione d'insieme delle quattro dimensioni affrontate.

La strutturazione temporale della sperimentazione, con la sua progressione dalle dimensioni più concettuali a quelle più critiche ed etiche, riflette l'impostazione del framework teorico sottostante, ma rappresenta anche un'applicazione della tassonomia di Bloom rivista (Anderson et al., 2001), con un movimento dalle competenze di base (conoscere, comprendere) a quelle più complesse (analizzare, valutare, creare). Questa progressione è stata concepita non solo come una sequenza didattica efficace, ma anche come un percorso di trasformazione concettuale, in cui i partecipanti potessero gradualmente costruire una comprensione profonda e multidimensionale dell'IA, delle sue potenzialità e delle sue implicazioni per il contesto educativo.

4.1.4 Struttura e contenuti del curriculum

Il curriculum implementato durante la sperimentazione si è fondato sul framework quadridimensionale per l'AI literacy precedentemente sviluppato (Cuomo et al., 2022; cfr. Capitolo 2), che rappresenta una rilettura originale delle competenze relative all'IA in chiave educativa. Questo framework teorico, coerente con l'approccio multidimensionale alla competenza digitale elaborato da Calvani et al. (2009) e con le più recenti concettualizzazioni dell'AI literacy (Long & Magerko, 2020; Ng et al., 2021), articola la competenza in quattro dimensioni fondamentali: conoscitiva, operativa, critica ed etica. Ciascuna di queste dimensioni ha costituito un modulo formativo specifico del percorso, secondo un'architettura curricolare coerente ed organica.

La progettazione di ciascun modulo ha seguito i principi dell'instructional design evidence-based (Gagne et al., 2005; Merrill, 2002), con una particolare attenzione all'allineamento costruttivo tra obiettivi di apprendimento, attività formative e strategie di valutazione (Biggs & Tang, 2011). Di seguito vengono analizzati in dettaglio i contenuti essenziali di ciascun modulo e la loro articolazione didattica.

Il primo modulo, incentrato sulla dimensione conoscitiva, è stato progettato per costruire le fondamenta concettuali dell'AI literacy, fornendo un quadro di riferimento storico, definitorio e concettuale dell'Intelligenza Artificiale. In linea con quanto suggerito da Touretzky et al. (2019) riguardo all'importanza di una comprensione concettuale dell'IA non limitata alla sua attualizzazione tecnologica, il modulo ha esplorato la traiettoria evolutiva della disciplina dalle intuizioni pionieristiche di Alan Turing fino ai recenti sviluppi dell'IA generativa. Particolare attenzione è stata dedicata all'analisi delle cosiddette "primavere" e "inverni" dell'IA, ovvero i cicli di entusiasmo e successiva disillusione che hanno caratterizzato lo sviluppo di questo campo (Floridi, 2020), offrendo così una prospettiva storica e critica sull'evoluzione non lineare della tecnologia.

Un elemento centrale del modulo è stato il "prisma delle definizioni", un'analisi comparativa delle diverse concettualizzazioni di IA proposte nel tempo da pionieri del settore, istituzioni internazionali e filosofi. Questa analisi ha permesso di evidenziare le diverse enfasi – sulla razionalità, sull'umanità, sull'azione o sul pensiero – e le implicazioni epistemologiche sottostanti a ciascuna definizione (Russell & Norvig, 2020). Complementare a questa esplorazione definitoria è stata la riflessione sul confronto tra intelligenza umana e artificiale, focalizzata sulle similitudini, le differenze e i limiti concettuali, in linea con l'approccio filosofico proposto da Searle (1980) e Dreyfus (1992) sulla natura dell'intelligenza e della comprensione.

Sul piano metodologico-didattico, questo modulo ha adottato un approccio plurale e interattivo. In accordo con i principi della pedagogia costruttivista (Fosnot, 2013), sono stati alternati momenti di presentazione teorica a discussioni aperte, utilizzando anche strumenti interattivi come WooClap per la raccolta in tempo reale di idee e percezioni, e la visione di brevi video stimolo selezionati secondo criteri di rilevanza, accessibilità e valenza motivazionale. Un'attività particolarmente significativa è

stata l'introduzione alla Philosophy for Children (P4C) come approccio metodologico per affrontare temi complessi e potenzialmente divisivi come l'IA nel contesto educativo. La P4C, sviluppata originariamente da Matthew Lipman (2003) come metodologia per promuovere il pensiero riflessivo e la comunità di ricerca filosofica, è stata presentata come uno strumento particolarmente adatto per esplorare le questioni concettuali, etiche e sociali sollevate dall'IA, in virtù della sua enfasi sul dialogo, sul pensiero critico e sulla costruzione collaborativa del significato. Le partecipanti hanno sperimentato questa metodologia progettando, in gruppi, una sessione di P4C sul tema dell'IA per una classe di scuola primaria, completa di obiettivi, fasi, domande guida e strategie di facilitazione.

Il secondo modulo, focalizzato sulla dimensione operativa, ha spostato l'attenzione dai "cosa" ai "come", esplorando i meccanismi che rendono possibile il funzionamento dell'IA. In linea con le indicazioni di Wing (2006, 2008) sull'importanza del pensiero computazionale come competenza fondamentale nella società digitale, il modulo ha approfondito il concetto di algoritmo – definito come sequenza finita, ordinata e non ambigua di istruzioni per risolvere un problema –, le sue caratteristiche essenziali e la sua applicazione pratica. Parallelamente, sono stati esplorati i principi fondamentali del pensiero computazionale: decomposizione dei problemi in componenti più semplici, riconoscimento di pattern, astrazione (focalizzazione sugli aspetti rilevanti) e progettazione algoritmica (Shute et al., 2017).

Un'attenzione particolare è stata dedicata al machine learning, paradigma dominante nell'IA contemporanea (Jordan & Mitchell, 2015), illustrando i concetti base e le principali tipologie di apprendimento automatico (supervisionato, non supervisionato, per rinforzo) attraverso esempi concreti e accessibili. In questo contesto, è stato evidenziato il ruolo cruciale dei dati come "carburante" dell'IA, analizzando l'importanza della loro quantità, qualità e rappresentatività per l'addestramento di modelli efficaci e equi. Infine, sono state esplorate le potenzialità del coding (sia nella sua versione plugged che unplugged) come strategia didattica per introdurre concetti algoritmici e computazionali nei contesti educativi.

Sul piano didattico, in accordo con i principi dell'apprendimento esperienziale teorizzati da Kolb (1984), questo modulo ha privilegiato un approccio marcatamente hands-on. Seguendo le indicazioni di Papert (1980) sull'importanza dell'esperienza concreta nella costruzione di conoscenze computazionali, i partecipanti hanno sperimentato direttamente algoritmi semplici attraverso giochi (es. il "gioco dei fiammiferi"), attività di coding unplugged (es. pixel art, body coding), e la creazione di alberi decisionali per comprendere i principi della classificazione algoritmica. Particolarmente efficace dal punto di vista didattico è stata l'attività di "classificazione della pasta", ispirata alle metodologie di apprendimento basate sul problem-based learning (Savery, 2006), dove i gruppi hanno dovuto sviluppare algoritmi per distinguere diversi formati di pasta, confrontando i risultati ottenuti e riflettendo metacognitivamente sulle difficoltà incontrate e sulle strategie utilizzate.

Il terzo modulo, dedicato alla dimensione critica, ha guidato i partecipanti verso una comprensione più profonda e problematizzante delle implicazioni connesse all'uso dell'IA nei diversi contesti sociali e educativi. In linea con l'approccio della critical media literacy (Kellner & Share, 2007) e con l'enfasi posta da studiosi come O'Neil (2016) e Eubanks (2018) sui potenziali rischi e ingiustizie algoritmiche, il modulo ha approfondito il concetto di bias algoritmico, analizzandone le diverse tipologie, le possibili origini (dati di addestramento distorti, scelte di progettazione, pregiudizi incorporati) e gli effetti potenzialmente discriminatori.

Coerentemente con la tradizione del pensiero critico in ambito educativo (Ennis, 2011; Paul & Elder, 2020), sono stati esplorati modelli teorici e strategie pratiche per l'applicazione del pensiero critico alla valutazione dei sistemi e dei prodotti dell'IA. Particolare attenzione è stata dedicata alle

problematiche di trasparenza e interpretabilità dei sistemi di IA complessi (il cosiddetto problema della "black box"), in accordo con le preoccupazioni espresse da studiosi come Rudin (2019) sulla necessità di modelli comprensibili e verificabili, soprattutto in ambiti sensibili come l'educazione. Infine, sono stati analizzati i meccanismi attraverso cui euristiche cognitive, stereotipi e pregiudizi possono influenzare sia la progettazione che l'interpretazione dei sistemi di IA, fornendo strumenti e strategie per un'analisi critica e consapevole dei prodotti dell'Intelligenza Artificiale.

Sul piano didattico, questo modulo ha integrato la riflessione teorica con attività pratiche mirate a "toccare con mano" i concetti trattati, secondo l'approccio dell'apprendimento situato (Lave & Wenger, 1991) e dell'educazione critica ai media (Buckingham, 2019). I partecipanti hanno analizzato criticamente immagini generate dall'IA, identificando stereotipi, bias e incongruenze. Hanno inoltre sperimentato l'uso di Teachable Machine, un tool educativo sviluppato da Google, per comprendere empiricamente come il machine learning possa amplificare i pregiudizi presenti nei dati di addestramento. Un'attività particolarmente rilevante dal punto di vista pedagogico è stata la valutazione comparativa di testi e immagini, alcune generate da esseri umani e altre dall'IA, stimolando una riflessione metacognitiva sulle caratteristiche distintive e sui criteri di discernimento che possono essere sviluppati e applicati nel contesto educativo.

Il quarto e ultimo modulo, incentrato sulla dimensione etica, ha esplorato le questioni morali fondamentali sollevate dallo sviluppo e dall'impiego dell'IA, con particolare attenzione alle implicazioni per il contesto educativo. In linea con il framework etico per una "Good AI Society" proposto da Floridi et al. (2018) e con i principi delineati dall'High-Level Expert Group on AI della Commissione Europea (2019), il modulo ha approfondito i principi etici chiave che dovrebbero guidare lo sviluppo e l'uso responsabile dell'IA: non maleficenza, privacy, trasparenza, responsabilità, equità e rispetto dell'autonomia umana.

Seguendo l'approccio dell'etica applicata e dei case studies (Rachels & Rachels, 2019), sono stati analizzati dilemmi morali specifici dell'IA attraverso casi paradigmatici e scenari controversi, come il dilemma del carrello autonomo (*Trolley problem*), la responsabilità per errori medici di sistemi IA o l'uso dell'IA nella sorveglianza. È stato anche valutato l'impatto potenziale dell'IA su ambiente (consumo energetico, impronta ecologica), società (disuguaglianze, concentrazione del potere) e lavoro (automazione, trasformazione delle professioni), evidenziando la necessità di un approccio sostenibile e centrato sull'uomo. Sono state inoltre prese in esame le questioni emergenti relative a copyright, autorialità e creatività nell'era dell'IA generativa, problematiche particolarmente rilevanti nel contesto educativo contemporaneo. Infine, è stato approfondito il concetto di "human-in-the-loop" come principio guida per garantire una supervisione etica e una responsabilità umana nei processi decisionali assistiti dall'IA.

Sul piano metodologico-didattico, in accordo con l'approccio dell'educazione morale attraverso la deliberazione collettiva proposto da Kohlberg (1984) e sviluppato da Noddings (2013) in chiave educativa, questo modulo ha privilegiato metodologie basate sul confronto argomentato e sulla costruzione partecipativa del giudizio etico. Il *debate* strutturato, metodologia con una lunga tradizione educativa e recentemente rivalutata come strumento per lo sviluppo delle competenze argomentative e di cittadinanza (Snider & Schnurer, 2006), è stato utilizzato come strategia principale: i partecipanti, divise in gruppi, hanno preparato e sostenuto posizioni opposte su temi eticamente controversi relativi all'IA, seguendo precise regole di argomentazione e confutazione. Questa metodologia ha permesso di sperimentare direttamente una strategia didattica trasferibile nel contesto scolastico, particolarmente adatta ad affrontare temi complessi e divisivi come quelli etici,

sviluppando contemporaneamente competenze argomentative, di ascolto attivo e di costruzione informata del giudizio morale.

La progressione tra i quattro moduli, ispirata alla tassonomia degli obiettivi educativi di Bloom rivista da Anderson et al. (2001), è stata concepita per costruire gradualmente una visione complessa e sfaccettata dell'IA, partendo dalle conoscenze di base, passando attraverso la comprensione dei meccanismi operativi, per arrivare all'analisi critica e alla riflessione etica. Questa sequenza ha permesso di affrontare l'AI literacy in modo sistematico e progressivo, rispettando i principi della "spirale dell'apprendimento" teorizzata da Bruner (1960), con un ritorno ciclico sui concetti fondamentali a livelli di complessità crescente.

Un elemento trasversale e caratterizzante dell'intero percorso è stata la costante riconduzione dei contenuti teorici e delle esperienze pratiche al contesto educativo, stimolando una continua riflessione su come tradurre i concetti appresi in pratiche didattiche efficaci e appropriate per l'età degli alunni della scuola dell'infanzia e primaria. Questo approccio, coerente con il modello TPACK (Technological Pedagogical Content Knowledge) proposto da Mishra e Koehler (2006), ha mirato a sviluppare non solo conoscenze tecniche sull'IA, ma anche e soprattutto la capacità di integrarle in modo significativo e pedagogicamente fondato nella pratica educativa.

4.2 Le strategie di ricerca

La sperimentazione del curriculum per l'AI literacy non ha avuto solo una finalità formativa, ma si è configurata come un vero e proprio studio empirico, con l'obiettivo di valutare l'efficacia del percorso proposto e di esplorare i processi di apprendimento e trasformazione concettuale attivati nei partecipanti. In questa sezione verranno dettagliate le strategie di ricerca adottate, con particolare attenzione al disegno complessivo dell'indagine, agli strumenti utilizzati per la raccolta dei dati e alle procedure di analisi implementate.

4.2.1 Disegno della ricerca

Lo studio si è configurato come una ricerca empirica di tipo prevalentemente qualitativo, con elementi quantitativi complementari, adottando un disegno pre-post senza gruppo di controllo. La scelta di questo approccio metodologico si è fondata sia su considerazioni pragmatiche legate al contesto dell'indagine (un laboratorio universitario curricolare), sia su considerazioni epistemologiche relative alla natura complessa e multidimensionale dell'oggetto di studio. Come sottolineato da Greene et al. (1989) e successivamente da Teddlie e Tashakkori (2009), i disegni di ricerca misti risultano particolarmente adeguati quando l'oggetto dell'indagine è un fenomeno complesso come nel caso dello sviluppo dell'AI literacy nelle sue diverse dimensioni.

Il disegno di ricerca si è articolato in tre momenti valutativi principali, allineati con le fasi della sperimentazione descritte nella sezione precedente:

Valutazione iniziale (pre): Condotta all'inizio del percorso formativo tramite la somministrazione del questionario sull'AI literacy (Biagini et al., 2023), appositamente validato per questo studio come descritto nel Capitolo 3. Questo strumento ha permesso di stabilire una baseline multidimensionale delle conoscenze, abilità e attitudini dei partecipanti, in accordo con l'approccio diagnostico raccomandato da Pellegrino et al. (2016) per la valutazione di competenze complesse in contesti formativi.

Valutazione processuale (durante): Realizzata mediante la raccolta sistematica dei diari riflessivi al termine di ciascuno dei quattro moduli formativi. Questa valutazione continua e longitudinale, in linea con i principi della valutazione formativa (William, 2011), ha consentito di monitorare i processi

di apprendimento in itinere, le trasformazioni concettuali progressivamente attivate e le evoluzioni nelle modalità di progettazione didattica ipotizzate dai partecipanti. La scelta di utilizzare diari riflessivi strutturati come strumento principale riflette l'approccio narrativo alla comprensione dei processi di apprendimento professionale proposto da Connelly e Clandinin (1990) e sviluppato nel contesto della formazione insegnanti da Schön (1983, 1987).

Valutazione finale (post): Effettuata mediante una seconda somministrazione del questionario sull'AI literacy, identico alla versione utilizzata nella fase pre per garantire la comparabilità dei dati. Questa misurazione finale ha permesso di evidenziare i cambiamenti quantitativi avvenuti rispetto alla baseline iniziale, seguendo l'approccio del value-added assessment descritto da Braun (2005).

La triangolazione tra dati quantitativi (questionari pre-post) e qualitativi longitudinali (diari riflessivi) ha consentito di adottare quello che Miles e Huberman (1994) definiscono un approccio "quali-quant" integrato, in cui l'analisi qualitativa approfondita viene supportata e, in alcuni casi, validata da evidenze quantitative complementari. Questo approccio integrato è particolarmente efficace per affrontare domande di ricerca complesse che richiedono sia l'identificazione di pattern generali sia la comprensione approfondita dei processi sottostanti (Creswell, 2009).

Le domande di ricerca che hanno guidato l'indagine, formulate in accordo con il framework teorico sull'AI literacy sviluppato in questa tesi, possono essere così articolate:

1. In che misura e con quali caratteristiche il curriculum proposto favorisce lo sviluppo dell'AI literacy nelle sue diverse dimensioni (conoscitiva, operativa, critica, etica) nei futuri insegnanti?
2. Quali processi di trasformazione concettuale (a livello di credenze, metafore, attitudini, percezioni) avvengono durante l'apprendimento delle diverse dimensioni dell'AI literacy?
3. Quali strategie didattiche vengono percepite come più efficaci e trasferibili per promuovere l'AI literacy nel contesto della scuola dell'infanzia e primaria?

Come evidenziato da Maxwell (2013), le domande di ricerca in un disegno prevalentemente qualitativo non mirano primariamente a testare ipotesi prestabilite, ma piuttosto a esplorare e comprendere in profondità un fenomeno, consentendo l'emergere di pattern e significati dai dati stessi. Nel caso specifico, l'enfasi è stata posta sulla comprensione dei processi di apprendimento e trasformazione concettuale, più che sulla semplice misurazione dell'efficacia del curriculum, in linea con l'approccio comprensivo alla ricerca educativa proposto da Eisner (2017).

È importante sottolineare che l'assenza di un gruppo di controllo, pur rappresentando una limitazione metodologica in termini di causalità diretta (Cook & Campbell, 1979), riflette una scelta consapevole bilanciata dalla ricchezza dei dati processuali raccolti e dalla triangolazione delle fonti. Come hanno mostrato Lincoln e Guba (1985), nei disegni di ricerca qualitativi e misti, la credibilità (equivalente della validità interna) può essere rafforzata attraverso strategie alternative, tra cui l'osservazione prolungata, la triangolazione metodologica e il monitoraggio longitudinale dei cambiamenti, tutti elementi incorporati nel presente disegno di ricerca.

Un ulteriore elemento metodologico distintivo è stato l'impiego di un approccio riflessivo sistematico sia da parte dei partecipanti (attraverso i diari) sia da parte dei ricercatori durante le fasi di raccolta e analisi dei dati. Questa "doppia ermeneutica" (Giddens, 1984) ha consentito di mantenere una consapevolezza critica costante rispetto ai processi interpretativi in atto, aumentando la trasparenza e

la *trustworthiness* dello studio, dimensioni fondamentali nei disegni di ricerca qualitativi (Lincoln & Guba, 1985).

4.2.2 Strumenti di raccolta dati

La strategia di raccolta dati è stata progettata per garantire una comprensione multidimensionale e dinamica dello sviluppo dell'AI literacy nei partecipanti. In accordo con i principi della triangolazione metodologica (Denzin, 1978), sono stati impiegati strumenti sia quantitativi che qualitativi, selezionati o sviluppati appositamente per rispondere alle specifiche domande di ricerca e per cogliere le diverse sfaccettature del fenomeno indagato.

A. Questionario sull'AI literacy

Il questionario utilizzato per le rilevazioni pre e post è uno strumento strutturato, sviluppato e validato specificamente per questo studio (Biagini et al., 2023). Si tratta di un questionario multidimensionale, coerente con il framework teorico adottato, che misura l'AI literacy nelle sue quattro dimensioni fondamentali. La versione finale dello strumento, risultato di un rigoroso processo di validazione che ha incluso analisi fattoriale esplorativa e confermativa, è composta da 40 item così distribuiti:

- 8 item per la dimensione conoscitiva
- 12 item per la dimensione operativa
- 10 item per la dimensione critica
- 10 item per la dimensione etica

Gli item sono formulati come affermazioni alle quali il rispondente deve indicare il proprio grado di accordo su una scala Likert a cinque punti (da "Fortemente in disaccordo" a "Fortemente d'accordo"). Alcuni esempi di item includono:

- Dimensione conoscitiva: "Conosco la differenza tra IA e altre tecnologie digitali"
- Dimensione operativa: "So che l'IA apprende dai dati che le vengono forniti"
- Dimensione critica: "Sono consapevole che l'IA può perpetuare pregiudizi presenti nei dati di addestramento"
- Dimensione etica: "Ritengo importante che le decisioni prese dall'IA siano trasparenti e spiegabili"

La scelta della scala Likert a 5 punti è stata effettuata in accordo con le raccomandazioni metodologiche di Likert (1932) e più recentemente di Hinkin (1998), che evidenziano come questa ampiezza della scala offra un buon compromesso tra precisione della misurazione e facilità di utilizzo, particolarmente adatta per somministrazioni digitali come nel caso del presente studio.

Il questionario ha dimostrato buone proprietà psicometriche, con un'eccellente affidabilità complessiva ($\alpha = 0.953$) e un'adeguata validità di costrutto, confermata dall'analisi fattoriale che ha supportato la struttura quadridimensionale teorizzata. Oltre agli item centrali sull'AI literacy, il questionario includeva una sezione iniziale per raccogliere dati demografici e informazioni sul background tecnologico dei partecipanti. Questo ha permesso non solo di caratterizzare dettagliatamente il campione, come descritto nella sezione 4.1.2, ma anche di esplorare eventuali correlazioni tra caratteristiche individuali e pattern di risposta o cambiamenti nell'AI literacy, un

aspetto fondamentale per comprendere l'eterogeneità dei processi di apprendimento (Jonassen & Grabowski, 1993).

Lo strumento è stato somministrato in formato digitale tramite la piattaforma *Qualtrics*, garantendo l'anonimato delle risposte attraverso l'uso di un codice personale che ha permesso di collegare i questionari pre e post dello stesso partecipante senza raccogliere dati personali identificativi. Questa modalità di somministrazione digitale ha comportato diversi vantaggi metodologici, tra cui la riduzione degli errori di inserimento dati, la possibilità di controlli automatici sulla completezza delle risposte e la facilitazione dell'analisi statistica successiva (Wright, 2005).

B. Diari riflessivi

I diari riflessivi hanno rappresentato lo strumento qualitativo centrale della ricerca, progettati per catturare i processi di apprendimento, le trasformazioni concettuali e le proiezioni didattiche dei partecipanti. Al termine di ciascuno dei quattro moduli (uno per dimensione dell'AI literacy), ogni partecipante è stato invitato a compilare un diario online strutturato in quattro sezioni principali, ciascuna progettata per evidenziare un aspetto specifico del processo di apprendimento e trasformazione:

Credenze iniziali: indagate attraverso la richiesta di elaborare una metafora personale sull'argomento del modulo, seguendo l'approccio della "metaphor elicitation technique" (Zaltman & Coulter, 1995). Questa scelta metodologica si basa sull'assunto, fondato nella linguistica cognitiva (Lakoff & Johnson, 1980) e successivamente sviluppato nella ricerca educativa da Sfard (1998) e Martinez et al. (2001), che le metafore non siano semplici ornamenti retorici, ma strutture cognitive fondamentali che rivelano le preconcezioni, le associazioni implicite e la tonalità emotiva con cui un individuo si approccia a un concetto complesso o astratto. Come evidenziato da Saban et al. (2007) nel contesto specifico della formazione insegnanti, l'analisi delle metafore consente di accedere a credenze implicite che spesso risultano difficili da articolare direttamente attraverso domande esplicite.

Conoscenza acquisita: valutata tramite la richiesta di selezionare, riportare e, soprattutto, commentare una definizione o una citazione ritenuta significativa relativa all'argomento del modulo. Questa tecnica, ispirata al "critical incident questionnaire" di Brookfield (1995), mira a stimolare un'appropriazione personale e critica delle conoscenze teoriche. Il commento personale è stato considerato cruciale per andare oltre la semplice memorizzazione e valutare il livello di comprensione, rielaborazione e rilevanza attribuita ai concetti chiave, in linea con i livelli superiori della tassonomia di Bloom (Anderson et al., 2001).

Applicazioni didattiche: esplorate chiedendo di ipotizzare una o più attività concrete per trasferire le conoscenze apprese in un contesto scolastico specifico. Questa sezione, che si ispira al concetto di "trasposizione didattica" elaborato da Chevallard (1991) e ripreso nel contesto dell'integrazione delle tecnologie nell'educazione da Mishra e Koehler (2006), mirava a rivelare la capacità di operare una trasformazione del sapere teorico in pratica educativa, traducendo la teoria in proposte didattiche concrete e significative. L'accento sulla progettazione didattica allinea questa sezione con la tradizione del design-based thinking nella formazione insegnanti (Laurillard, 2012).

Visione dello scenario d'uso: approfondita attraverso la scelta motivata di una tra sei immagini predefinite, ciascuna rappresentante un diverso scenario di interazione uomo-tecnologia in contesti di apprendimento (dallo studio individuale alla discussione di gruppo, dal laboratorio collaborativo all'analisi di artefatti). Questa tecnica, che si ispira ai metodi di photo elicitation (Harper, 2002),

permette di cogliere le rappresentazioni mentali e le preferenze implicite relative ai contesti d'uso della tecnologia. La scelta e la relativa giustificazione hanno fornito indicazioni preziose sulle concezioni pedagogiche sottostanti e sulle visioni d'uso preferite per ciascuna dimensione dell'AI literacy.



I diari sono stati compilati online tramite Google Forms, con un limite di tempo e di spazio per ciascuna sezione (circa 1000 caratteri), per incoraggiare risposte concise ma significative. Questa strutturazione rispondeva a diverse esigenze metodologiche: facilitare la comparabilità tra partecipanti e tra i diversi moduli; assicurare che tutti gli aspetti chiave dell'esperienza di apprendimento fossero documentati; e stimolare una riflessione focalizzata ed efficiente. Come evidenziato da Moon (2006), i diari riflessivi strutturati offrono un equilibrio ottimale tra libertà espressiva e focalizzazione, evitando sia la dispersione tematica sia l'eccessiva direttività.

Complessivamente, sono stati raccolti 249 diari, costituendo un corpus testuale di oltre 100.000 parole, una mole di dati qualitativi considerevole che ha richiesto approcci analitici innovativi, combinando metodi interpretativi tradizionali con tecniche computazionali, come verrà dettagliato nella sezione 4.2.3.

La scelta dei diari riflessivi strutturati come strumento principale di ricerca qualitativa è stata motivata dalla loro capacità di favorire contemporaneamente l'apprendimento metacognitivo nei partecipanti e di fornire ai ricercatori una finestra privilegiata sui processi di trasformazione concettuale in atto. Come sottolineato da Boud (2001), la riflessione scritta non è semplicemente uno strumento di documentazione, ma un potente dispositivo di apprendimento che favorisce la ristrutturazione cognitiva e la costruzione attiva di significato. Il focus sulle metafore, in particolare, ha permesso di accedere a credenze e modelli mentali che spesso rimangono impliciti o difficili da articolare in modo diretto, in accordo con la prospettiva costruttivista dell'apprendimento (von Glasersfeld, 1995).

4.2.3 Procedure di analisi dei dati

L'analisi dei dati raccolti ha seguito un approccio metodologico misto, allineato con i principi del convergent parallel design delineati da Creswell e Plano Clark (2018), che prevede la raccolta

simultanea di dati quantitativi e qualitativi, seguita da un'analisi indipendente ma complementare, finalizzata a una successiva integrazione interpretativa. Questa strategia metodologica è stata scelta in virtù della sua capacità di cogliere la complessità multidimensionale del fenomeno studiato, combinando la forza della generalizzabilità dei dati quantitativi con la profondità e la ricchezza interpretativa dei dati qualitativi (Johnson & Onwuegbuzie, 2004).

Analisi dei questionari (pre-post)

I dati quantitativi provenienti dai questionari pre e post-intervento sono stati sottoposti a un'analisi sistematica e rigorosa utilizzando il software statistico R (versione 4.1.2), un ambiente versatile e potente per l'analisi statistica e la visualizzazione dei dati, ampiamente riconosciuto nella comunità scientifica (R Core Team, 2021). L'analisi si è sviluppata seguendo una sequenza metodologica articolata in fasi progressive, in accordo con le linee guida per l'analisi di dati longitudinali in ambito educativo (Singer & Willett, 2003).

La prima fase ha riguardato la pulizia e la preparazione dei dati, un passaggio cruciale per garantire la qualità e l'affidabilità delle analisi successive (Wickham & Grolemund, 2017). Questa fase ha incluso la verifica della completezza dei record, l'identificazione e la gestione appropriata dei dati mancanti secondo le procedure raccomandate in letteratura (Allison, 2001), e la codifica o ricodifica delle variabili per renderle adatte alle analisi previste.

Successivamente, sono state condotte analisi descrittive approfondite, finalizzate a fornire un quadro dettagliato della distribuzione dei punteggi per ciascun item e per le scale aggregate relative alle quattro dimensioni dell'AI literacy. Queste analisi hanno incluso il calcolo delle frequenze, delle medie, delle deviazioni standard e di altri indici di tendenza centrale e dispersione, nonché la verifica della normalità delle distribuzioni attraverso test statistici (Shapiro-Wilk) e analisi grafiche (Q-Q plot). Tali analisi preliminari hanno fornito una base solida per la scelta appropriata dei test statistici successivi, in accordo con i principi della statistica inferenziale (Field et al., 2012).

Il confronto pre-post, cuore dell'analisi quantitativa, è stato effettuato adottando un approccio metodologicamente robusto, che ha integrato test parametrici (t-test per campioni appaiati) e non parametrici (test di Wilcoxon), a seconda della natura distributiva delle variabili. Come sottolineato da Pallant (2020), questa strategia analitica duale permette di massimizzare la potenza statistica mantenendo al contempo un adeguato controllo dell'errore di Tipo I. L'analisi è stata condotta sia a livello di singoli item, per identificare le aree specifiche di cambiamento, sia a livello di dimensioni aggregate, per valutare l'impatto complessivo del curriculum su ciascuna delle quattro dimensioni dell'AI literacy.

Ad integrazione delle analisi di significatività statistica, sono stati calcolati anche indici di effect size (d di Cohen per i test parametrici, r di Rosenthal per quelli non parametrici), in linea con le raccomandazioni della comunità scientifica che sottolineano l'importanza di riportare non solo la significatività ma anche la magnitudine dell'effetto osservato (Sullivan & Feinn, 2012). Questi indici hanno permesso di quantificare l'entità del cambiamento indotto dall'intervento formativo, fornendo una metrica standardizzata che facilita la comparazione con altri studi e l'interpretazione pratica dei risultati.

Per esplorare in modo più approfondito le dinamiche di apprendimento, sono state condotte anche analisi delle correlazioni, utilizzando coefficienti appropriati (Pearson o Spearman, a seconda della natura distributiva delle variabili). Queste analisi hanno permesso di esaminare le relazioni tra le diverse dimensioni dell'AI literacy e tra queste e variabili demografiche o di background, fornendo

indicazioni preziose sui pattern di co-variazione e sulle possibili interazioni tra diversi aspetti della competenza.

Infine, sono state implementate analisi differenziali attraverso ANOVA o test non parametrici equivalenti, per indagare possibili differenze nell'impatto del curriculum in base a caratteristiche specifiche dei partecipanti, quali l'esperienza pregressa in ambito educativo, l'atteggiamento dichiarato verso la tecnologia o la fascia d'età. Queste analisi hanno contribuito a una comprensione più sfumata e contestualizzata dell'efficacia dell'intervento, in linea con l'approccio della valutazione educativa sensibile alle differenze individuali (Tomlinson, 2014).

Dal momento che non tutti i partecipanti hanno compilato entrambi i questionari, le analisi pre-post sono state condotte sui soli soggetti con dati completi sia al pre che al post ($N = 44$). A livello di singolo item, i pochi valori mancanti sono stati gestiti con esclusione caso-per-caso (listwise) nei test statistici, mentre i punteggi di scala sono stati calcolati come media dei soli item compilati, richiedendo almeno il 70% delle risposte presenti.

Analisi dei diari riflessivi

Il corpus testuale costituito dai diari riflessivi, che rappresenta una risorsa di eccezionale ricchezza per comprendere i processi di apprendimento e trasformazione concettuale dei partecipanti, è stato sottoposto a un'analisi approfondita e metodologicamente innovativa. Questa analisi ha integrato l'interpretazione qualitativa tradizionale con tecniche computazionali avanzate di Natural Language Processing (NLP), seguendo l'emergente paradigma dei metodi misti computazionali in scienze sociali (Nelson, 2020; Wiedemann, 2022).

Tale approccio metodologico ibrido, che coniuga la sensibilità interpretativa umana con la potenza analitica degli algoritmi computazionali, è stato adottato per rispondere alla duplice esigenza di cogliere le sfumature semantiche, le narrazioni soggettive e le costruzioni di significato individuali (attraverso l'analisi qualitativa) e, contemporaneamente, di identificare pattern linguistici, strutture tematiche ricorrenti e relazioni concettuali su larga scala (attraverso l'analisi NLP). Come evidenziato da Grimmer e Stewart (2013), questa integrazione metodologica risulta particolarmente preziosa nell'analisi di corpora testuali estesi, dove la sola analisi qualitativa tradizionale potrebbe risultare limitata dalla capacità umana di processare grandi quantità di testo.

L'analisi si è articolata in quattro fasi principali, seguendo un processo iterativo e ricorsivo piuttosto che strettamente lineare, in accordo con i principi della Grounded Theory costruttivista (Charmaz, 2014).

La prima fase ha riguardato la preelaborazione dei dati testuali, un passaggio fondamentale per preparare il corpus all'analisi computazionale (Manning et al., 2008). Il corpus è stato importato in un ambiente di analisi dati (Google Colab), pulito da eventuali artefatti di raccolta e sottoposto a procedure standard di pre-processing NLP: normalizzazione (conversione sistematica in minuscolo), tokenizzazione (suddivisione del testo in unità lessicali discrete), rimozione delle stopwords (articoli, preposizioni, congiunzioni e altre parole grammaticali molto frequenti ma poco informative) e, per alcune analisi specifiche, lemmatizzazione (riduzione delle parole flesse alla loro forma base). Queste operazioni, pur comportando una minima perdita di informazione contestuale, sono essenziali per ridurre il "rumore" lessicale e rendere il testo trattabile dagli algoritmi NLP in modo efficiente e significativo.

La seconda fase ha previsto un'analisi esplorativa NLP approfondita, implementata utilizzando un ecosistema di librerie Python specializzate (tra cui NLTK, Scikit-learn, Pandas, Matplotlib, NetworkX). Questa fase ha incluso:

L'analisi di frequenza lessicale e il calcolo del Type-Token Ratio (TTR), per identificare i termini più ricorrenti e misurare la diversità lessicale del corpus e delle sue sezioni, fornendo indicazioni sulla ricchezza e la ripetitività del vocabolario utilizzato in ciascuna fase.

L'implementazione della tecnica TF-IDF (Term Frequency-Inverse Document Frequency), che assegna un peso maggiore ai termini frequenti in un documento specifico ma rari nel resto del corpus, permettendo di identificare le parole chiave più distintive e caratterizzanti per ogni fase della riflessione.

Il calcolo di misure di similarità (coseno e Jaccard) tra le diverse sezioni dei diari, per quantificare l'evoluzione del discorso e valutare l'ipotesi di un apprendimento trasformativo, con valori bassi indicativi di un significativo cambiamento nel linguaggio e nei temi trattati.

L'applicazione di algoritmi di Topic Modeling (LDA, Latent Dirichlet Allocation), per identificare in modo non supervisionato i temi emergenti dal corpus e la loro distribuzione nei documenti, offrendo una visione strutturale delle aree concettuali affrontate dai partecipanti.

Lo sviluppo di analisi di co-occorrenza e reti semantiche, per visualizzare le relazioni tra termini che tendono ad apparire insieme nello stesso contesto e identificare cluster tematici attraverso la rappresentazione a grafo.

L'implementazione di tecniche di Sentiment Analysis basate su lessico, per valutare la tonalità emotiva delle diverse sezioni dei diari e la sua evoluzione nel tempo.

La generazione di Word Clouds, per rappresentare visivamente i termini più frequenti in modo intuitivo e comunicativo.

La terza fase, sviluppata parallelamente alla precedente, ha previsto un'analisi tematica qualitativa condotta secondo il framework metodologico proposto da Braun e Clarke (2006, 2012), riconosciuto come uno degli approcci più rigorosi e sistematici all'analisi tematica in ambito qualitativo. Questa analisi ha seguito le sei fasi canoniche del metodo: familiarizzazione approfondita con i dati attraverso letture multiple; generazione sistematica di codici descrittivi iniziali per segmenti di testo significativi; ricerca attiva di temi sovraordinati raggruppando i codici correlati; revisione critica e affinamento dei temi; definizione precisa e denominazione appropriata dei temi finali; e infine, produzione della narrazione analitica. Particolare attenzione è stata dedicata all'analisi interpretativa delle metafore, seguendo l'approccio della linguistica cognitiva che le vede come finestre privilegiate sulle credenze implicite (Lakoff & Johnson, 1980), all'esame del contenuto dei commenti alle definizioni/citazioni, alla categorizzazione delle tipologie di attività didattiche proposte e all'analisi delle motivazioni sottostanti la scelta delle immagini rappresentative degli scenari d'uso.

La quarta e ultima fase, cruciale per un autentico approccio misto, ha riguardato l'interpretazione integrata dei risultati. In questa fase, le evidenze emerse dalle diverse analisi NLP (frequenze, termini chiave TF-IDF, topic LDA, pattern di similarità, reti semantiche) sono state sistematicamente confrontate, incrociate e integrate con i temi e le interpretazioni derivanti dall'analisi qualitativa. L'obiettivo non era semplicemente giustapporre risultati derivanti da metodologie diverse, ma utilizzare le evidenze quantitative per validare, arricchire o problematizzare le interpretazioni

qualitative, e, viceversa, impiegare la comprensione qualitativa del contesto e del significato per dare senso ai pattern computazionali emersi. Questa triangolazione metodologica (Denzin, 1978) ha permesso di costruire una narrazione dei risultati più robusta, dettagliata e sfaccettata, mitigando i limiti intrinseci di ciascun approccio metodologico considerato singolarmente e rafforzando la validità complessiva delle conclusioni.

I risultati dettagliati di queste analisi, che costituiscono il cuore empirico della presente tesi, saranno presentati e discussi approfonditamente nel Capitolo 5, dedicato agli esiti della sperimentazione. La struttura del capitolo seguirà l'architettura quadridimensionale del framework dell'AI literacy, presentando per ciascuna dimensione una sintesi integrata delle evidenze quantitative e qualitative, seguita da un'analisi trasversale delle traiettorie di apprendimento e trasformazione concettuale osservate lungo l'intero percorso formativo.

4.2.4 Considerazioni etiche e riflessioni sulle limitazioni metodologiche

La presente ricerca è stata condotta nel rigoroso rispetto dei principi etici che governano l'indagine in ambito educativo, in conformità con le linee guida internazionali (BERA, 2018; AERA, 2011) e con i protocolli etici dell'ateneo di appartenenza. La dimensione etica ha pervaso l'intero processo di ricerca, dalle fasi di progettazione fino alla disseminazione dei risultati, con un'attenzione particolare alla protezione dei diritti e del benessere dei partecipanti.

Tutti i partecipanti hanno fornito il loro consenso informato alla raccolta e all'analisi dei dati, dopo essere stati dettagliatamente edotti riguardo agli scopi dello studio, alle procedure di raccolta dati, alla natura volontaria della partecipazione e alle garanzie di anonimato e confidenzialità. In accordo con i principi dell'etica della ricerca educativa delineati da Hammersley e Traianou (2012), particolare attenzione è stata dedicata ad evitare potenziali conflitti di interesse derivanti dal doppio ruolo del ricercatore come formatore e valutatore, chiarendo esplicitamente che il rifiuto di partecipare non avrebbe comportato alcuna conseguenza negativa sul percorso formativo o sulla valutazione finale del laboratorio.

Il trattamento dei dati personali è stato effettuato in conformità con la normativa vigente in materia di privacy (GDPR, Regolamento UE 2016/679), adottando misure tecniche e organizzative adeguate a garantire la sicurezza e la confidenzialità delle informazioni raccolte. I dati sono stati pseudonimizzati attraverso l'attribuzione di codici alfanumerici che hanno consentito di collegare i questionari pre e post e i diari riflessivi dello stesso partecipante, senza comprometterne l'anonimato.

L'analisi critica delle limitazioni metodologiche rappresenta un elemento fondamentale di qualsiasi ricerca scientifica rigorosa, poiché consente di contestualizzare adeguatamente i risultati ottenuti e di delineare con trasparenza i confini entro cui le conclusioni possono essere considerate valide. Come sottolineato da Shadish, Cook e Campbell (2002), la riflessione esplicita sulle minacce alla validità interna ed esterna costituisce un presupposto indispensabile per una corretta interpretazione dei risultati e per l'avanzamento incrementale della conoscenza scientifica.

In questa prospettiva, è doveroso evidenziare alcune limitazioni metodologiche che caratterizzano il presente studio.

In primo luogo, l'assenza di un gruppo di controllo costituisce una limitazione significativa rispetto alla possibilità di stabilire nessi causali diretti tra l'intervento formativo e i cambiamenti osservati. Come evidenziato da Campbell e Stanley (1963) nella loro disamina dei disegni quasi-sperimentali, la mancanza di un gruppo di confronto rende problematica l'esclusione di spiegazioni alternative per i cambiamenti rilevati, quali la maturazione naturale, l'effetto della storia, o l'influenza di fattori

esterni concomitanti. Tuttavia, questa limitazione è stata parzialmente mitigata attraverso l'adozione di un approccio metodologico misto e l'implementazione di una robusta triangolazione dei dati (Denzin, 1978), che ha consentito di corroborare i risultati quantitativi con evidenze qualitative approfondite, aumentando così la plausibilità dell'interpretazione causale proposta.

In secondo luogo, sebbene la partecipazione al laboratorio fosse obbligatoria, la natura volontaria della partecipazione alla raccolta dati potrebbe aver introdotto un bias di selezione, con una possibile sovrarappresentazione dei partecipanti più motivati o interessati al tema. Questa limitazione, comune a molti studi in contesti educativi laboratoriali (Gall et al., 2007), suggerisce cautela nell'interpretazione dei risultati, soprattutto in relazione alla loro generalizzabilità all'intera popolazione di riferimento.

In terzo luogo, la durata relativamente circoscritta dell'intervento (36 ore, distribuite su circa sei settimane) non consente di valutare adeguatamente la persistenza temporale degli effetti osservati e la loro stabilità nel lungo periodo. Come evidenziato da Creswell (2018), gli studi longitudinali offrono vantaggi significativi nella valutazione degli impatti formativi, permettendo di distinguere tra effetti transitori ed effetti duraturi. L'assenza di misurazioni di follow-up a distanza temporale rappresenta quindi un limite da considerare nell'interpretazione dei risultati.

Infine, la specificità del campione, costituito esclusivamente da future insegnanti della scuola dell'infanzia e primaria, limita la generalizzabilità dei risultati ad altri contesti professionali o formativi. Sebbene questa specificità sia coerente con gli obiettivi della ricerca, focalizzati proprio sull'AI literacy in ambito educativo, è importante riconoscere che i processi di apprendimento e trasformazione concettuale osservati potrebbero manifestarsi in modo diverso in popolazioni con caratteristiche sociodemografiche e professionali differenti.

Nonostante queste limitazioni, il presente studio offre un contributo significativo alla comprensione dei processi di apprendimento relativi all'AI literacy in ambito educativo. La ricchezza e la complementarità dei dati raccolti, insieme alla rigorosità delle procedure di analisi, consentono di considerare i risultati ottenuti come una base empirica solida per valutare l'efficacia del curriculum proposto e per comprendere i processi di trasformazione concettuale attivati nei partecipanti. Come sottolineato da Maxwell (2013), la validità di una ricerca qualitativa o mista non risiede nella generalizzabilità statistica dei suoi risultati, quanto piuttosto nella profondità e nella ricchezza della comprensione che essa offre rispetto ai fenomeni studiati e nella sua capacità di illuminare processi e meccanismi che potrebbero manifestarsi, con le dovute specificità, anche in altri contesti.

In quest'ottica, le limitazioni metodologiche evidenziate non compromettono la validità complessiva dello studio, ma ne definiscono i confini interpretativi e indicano direzioni per future ricerche, che potrebbero includere disegni sperimentali con gruppo di controllo, campioni più diversificati e follow-up a lungo termine.

CAPITOLO 5: La sperimentazione del curriculum per l'AI literacy: risultati e discussione

5.1 Analisi dei risultati della sperimentazione: la rilevazione dell'AI literacy

L'analisi dell'efficacia degli interventi formativi sull'Intelligenza Artificiale rappresenta un elemento cruciale per comprendere come migliorare l'AI literacy nella popolazione. Il presente capitolo descrive l'approccio metodologico aggiornato adottato per analizzare i dati raccolti mediante questionari pre e post-intervento, con l'obiettivo di valutare in modo rigoroso i cambiamenti nelle diverse dimensioni dell'AI literacy.

La rilevazione mediante disegno pre-post test costituisce uno strumento metodologico particolarmente indicato per valutare l'impatto di un intervento formativo (Cook & Campbell, 1979). Tale disegno permette di misurare e confrontare le conoscenze, attitudini o comportamenti dei soggetti prima e dopo l'esposizione a un determinato stimolo o trattamento, nel nostro caso un intervento formativo sull'AI literacy. Il vantaggio principale di questo approccio risiede nella possibilità di valutare direttamente il cambiamento intra-soggetto, riducendo così la variabilità dovuta alle differenze individuali che caratterizzano invece i confronti tra gruppi diversi (Dimitrov & Rumrill, 2003).

L'analisi statistica di dati pre-post richiede particolare attenzione metodologica, poiché la scelta degli strumenti analitici più appropriati dipende dalle caratteristiche specifiche dei dati raccolti, in particolare dalla loro distribuzione e dalla natura delle variabili misurate. Nel contesto di questa ricerca, abbiamo adottato un approccio sistematico e rigoroso che include la verifica preliminare delle assunzioni statistiche, l'analisi fattoriale confermativa per verificare l'invarianza di misurazione, l'applicazione dei test più appropriati con correzioni per la molteplicità, il calcolo di effect size con intervalli di confidenza bootstrap, e una valutazione completa della robustezza dei risultati attraverso multiple analisi di sensibilità (Sullivan & Feinn, 2012).

5.1.1 Caratteristiche dei partecipanti

L'analisi quantitativa pre-post intervento si basa sui dati raccolti dai partecipanti alla sperimentazione, le cui caratteristiche demografiche e professionali sono state descritte in dettaglio nel paragrafo 4.1.1. In questa sede, è sufficiente ricordare che il campione era composto prevalentemente da donne (96,15 %), con una significativa presenza di docenti giovani (92,59 % con meno di 5 anni di servizio) e una marcata eterogeneità in termini di area disciplinare e familiarità pregressa con l'IA (il 65% dichiarava di non avere esperienze precedenti con queste tecnologie).

Tale composizione del campione risulta particolarmente rilevante per l'interpretazione dei risultati che seguono, poiché consente di esplorare l'efficacia dell'intervento formativo su un gruppo di professionisti con solida esperienza didattica ma limitata esposizione specifica all'Intelligenza Artificiale.

5.1.2 Verifica delle assunzioni statistiche

Prima di procedere con l'analisi inferenziale, è fondamentale verificare le assunzioni sottostanti ai test statistici che si intendono utilizzare. Il confronto tra misure pre e post-intervento può essere effettuato mediante test parametrici (come il t-test per campioni appaiati) o test non parametrici (come il test di Wilcoxon per ranghi con segno), a seconda che i dati soddisfino o meno determinate assunzioni.

La principale assunzione da verificare nel caso di test parametrici per dati appaiati riguarda la normalità della distribuzione delle differenze tra le misure pre e post, e non la normalità delle singole misure (Altman, 1991). Tale verifica è cruciale poiché l'applicazione di test parametrici a dati che violano questa assunzione può portare a conclusioni errate, in particolare quando la dimensione del campione è ridotta (Ghasemi & Zahediasl, 2012).

Per verificare la normalità delle differenze, abbiamo utilizzato un approccio integrato che combina test statistici formali, analisi degli indici di forma e ispezione visuale delle distribuzioni:

Test statistici formali: Il test di Shapiro-Wilk è stato scelto per la sua maggiore potenza rispetto ad altri test di normalità, soprattutto per campioni di dimensione inferiore a 50 (Razali & Wah, 2011). Questo test valuta l'ipotesi nulla che i dati provengano da una popolazione normalmente distribuita.

Analisi degli indici di forma: Sono stati calcolati skewness (asimmetria) e kurtosis (curtosi) per ciascuna distribuzione di differenze. Valori di skewness e kurtosis compresi nell'intervallo ± 1 sono generalmente considerati indicativi di una deviazione non problematica dalla normalità (Hair et al., 2010).

Ispezione visuale mediante violin plot: Per una comprensione più intuitiva delle distribuzioni, abbiamo creato violin plot che mostrano simultaneamente la densità di probabilità, la distribuzione dei quantili e eventuali outlier per ciascuna area e dimensione dell'AI literacy.

La Tabella 5.1 presenta i risultati di questa verifica preliminare per ciascuna delle sei aree valutate, per le quattro dimensioni aggregate e per l'indice globale.

Tabella 5.1. Verifica della normalità delle differenze pre-post per ogni area e dimensione

Area/dimensione	Shapiro-Wilk W	p-value	Skewness	Kurtosis	Normalità
Fondamenti teorici	0.886	0.0009	-0.45	0.21	No
Funzioni matematiche	0.881	0.0011	-0.52	0.18	No
Applicazione concetti	0.879	0.0013	-0.48	0.15	No
Valutazione critica	0.872	0.0026	-0.41	-0.12	No
Progettazione IA	0.875	0.0019	-0.44	0.09	No
Questioni etiche	0.869	0.0031	-0.38	-0.08	No
Dimensioni aggregate					
Dimensione Conoscitiva	0.878	0.0013	-0.49	0.19	No
Dimensione Operativa	0.873	0.0021	-0.46	0.12	No
Dimensione Critica	0.872	0.0026	-0.41	-0.12	No

Dimensione Etica	0.869	0.0031	-0.38	-0.08	No
Indice globale	0.976	0.5198	-0.12	0.05	Sì

Come evidenziato dalla Tabella 5.1, solo l'indice globale (media delle 6 aree) soddisfa l'assunzione di normalità ($p = 0.5198$), mentre tutte le singole aree e dimensioni aggregate mostrano una deviazione significativa dalla normalità, con valori p del test di Shapiro-Wilk inferiori alla soglia critica di 0.05. Gli indici di skewness e kurtosis confermano questi risultati, mostrando una tendenza generale verso distribuzioni leggermente asimmetriche negative (i partecipanti tendono a migliorare) con kurtosis vicina alla normalità.

Benché la letteratura suggerisca che test parametrici come il t-test possano essere relativamente robusti a violazioni moderate dell'assunzione di normalità con campioni di dimensione adeguata ($n > 30$) (Lumley et al., 2002), abbiamo preferito adottare un approccio più conservativo, optando per test non parametrici quando l'assunzione di normalità non è rispettata. Questo approccio garantisce conclusioni più caute e affidabili, specialmente considerando l'importanza di una valutazione rigorosa dell'efficacia dell'intervento formativo.

Per completezza metodologica, abbiamo esteso la verifica di normalità anche alle altre variabili del questionario che saranno oggetto di analisi comparative pre-post. La Tabella 5.2 presenta i risultati per le domande relative alla familiarità con il concetto di IA, all'utilizzo di applicazioni di IA e alle attitudini generali verso l'IA.

Tabella 5.2. Verifica della normalità delle differenze pre-post per variabili aggiuntive

Variabile	Shapiro-Wilk p	Normalità
Familiarità con il concetto di IA	0.023	Non confermata
Utilizzo di applicazioni di IA	0.001	Non confermata
Preoccupazione/entusiasmo per l'IA	0.003	Non confermata
IA è esaltante	0.038	Non confermata
IA è pericolosa	0.076	Confermata
IA fa molti errori	0.158	Confermata
IA può avere risultati migliori degli umani	0.041	Non confermata
IA migliore in lavori di routine	0.054	Confermata
IA può fornire nuove opportunità economiche	0.089	Confermata
Società beneficerà di IA	0.045	Non confermata

Come si può osservare dalla Tabella 5.2, diverse variabili aggiuntive mostrano una deviazione dalla normalità secondo il test di Shapiro-Wilk. Per queste variabili, sarà necessario adottare test non parametrici per l'analisi inferenziale, in particolare il test di Wilcoxon per ranghi con segno, che rappresenta l'alternativa non parametrica al t-test per campioni appaiati.

Un ulteriore aspetto da considerare è l'omoschedasticità (omogeneità delle varianze) delle misure pre e post. Sebbene questa assunzione sia meno critica per i test appaiati rispetto ai test per campioni indipendenti, abbiamo verificato che non ci fossero differenze estreme nelle varianze tra misure pre e post, utilizzando il test F di Levene. I risultati hanno confermato l'assenza di violazioni significative dell'omoschedasticità per tutte le variabili analizzate.

Infine, è importante verificare l'assenza di outlier significativi che potrebbero influenzare impropriamente i risultati. L'analisi mediante box plot e il criterio delle 3 deviazioni standard ha identificato alcuni valori potenzialmente anomali per alcune variabili. Tuttavia, dopo un'attenta valutazione, si è deciso di mantenere questi valori nell'analisi, in quanto rappresentano risposte legittime piuttosto che errori di misurazione, e la loro esclusione potrebbe introdurre bias nei risultati. Il numero limitato di questi outlier (meno del 5% del campione) e la loro distribuzione equilibrata tra valori positivi e negativi suggeriscono che il loro impatto sui risultati complessivi sia minimamente rilevante.

In conclusione, la verifica delle assunzioni statistiche ha evidenziato la necessità di utilizzare test non parametrici (Wilcoxon) per l'analisi della maggior parte delle variabili, riservando l'uso del t-test per campioni appaiati solo all'indice globale dell'AI literacy, che soddisfa l'assunzione di normalità. Questa differenziazione nell'approccio analitico garantisce robustezza alle conclusioni che saranno tratte dall'analisi dei dati.

Alla luce dei risultati della verifica di normalità, il test di Wilcoxon per ranghi con segno è stato scelto come test statistico principale per valutare la significatività delle differenze pre-post nelle sei aree di conoscenza/competenza e nelle quattro dimensioni aggregate dell'AI literacy. Questo test non parametrico è particolarmente adatto quando l'assunzione di normalità delle differenze è violata, in quanto si basa sui ranghi delle differenze anziché sui valori originali (Wilcoxon, 1945).

Per l'indice globale dell'AI literacy, che soddisfa l'assunzione di normalità, è stato utilizzato il t-test per campioni appaiati, che offre maggiore potenza statistica rispetto alle alternative non parametriche quando le assunzioni sono soddisfatte (McDonald, 2014).

Per ciascun test è stato calcolato il p-value, utilizzando il convenzionale livello di significatività $\alpha = 0.05$. Tuttavia, consapevoli dei limiti della sola significatività statistica (Wasserstein & Lazar, 2016), abbiamo integrato l'analisi con il calcolo della dimensione dell'effetto.

Per il test di Wilcoxon, la dimensione dell'effetto è stata calcolata convertendo la statistica del test in un valore d equivalente, seguendo la procedura descritta da Rosenthal (1991) e adattata da Rosenthal e DiMatteo (2001). Questa conversione permette di esprimere l'effetto nella stessa scala del d di Cohen, facilitando l'interpretazione e il confronto con la letteratura esistente.

Per il t-test per campioni appaiati, la dimensione dell'effetto è stata calcolata mediante il d di Cohen (Cohen, 1988). Per l'interpretazione della dimensione dell'effetto, abbiamo adottato le convenzioni proposte da Cohen (1988) per il d di Cohen:

- $d < 0.2$ indica un effetto trascurabile
- $0.2 \leq d < 0.5$ indica un effetto piccolo
- $0.5 \leq d < 0.8$ indica un effetto medio
- $d \geq 0.8$ indica un effetto grande

Un aspetto importante dell'analisi è la gestione della molteplicità dei test. Quando si eseguono numerosi test statistici sullo stesso dataset, aumenta il rischio di incorrere in errori di tipo I, ovvero di rilevare significatività statistica dove in realtà non esiste (Benjamini & Hochberg, 1995). Per mitigare questo rischio, abbiamo applicato la correzione di Bonferroni per le famiglie di test correlati, dividendo il livello di significatività α per il numero di test nella famiglia. Ad esempio, per la famiglia delle sei aree di conoscenza/competenza, il livello di significatività corretto è $\alpha' = 0.05/6 = 0.0083$.

Oltre all'analisi inferenziale basata sui test di significatività, abbiamo condotto anche analisi descrittive dettagliate, calcolando statistiche come media, mediana, deviazione standard, minimo e massimo per ciascuna variabile, sia nel pre-test che nel post-test. Queste statistiche forniscono una visione più completa e granulare dei dati, che va oltre la semplice determinazione di significatività.

Per fornire una valutazione più approfondita dell'efficacia dell'intervento, abbiamo implementato anche analisi complementari:

1. Equivalence Testing (TOST): Per verificare che l'effetto dell'intervento sia non solo statisticamente significativo, ma anche praticamente rilevante, abbiamo utilizzato il Two One-Sided Tests (TOST) approach (Lakens, 2017). Questo metodo permette di testare se l'effetto osservato è significativamente diverso da un effetto considerato "equivalente a zero" (SESOI, Smallest Effect Size Of Interest), che abbiamo definito come $|d| \leq 0.20$, corrispondente alla soglia convenzionale per un effetto trascurabile.
2. Reliable Change Index (RCI): Per valutare la significatività clinica del cambiamento a livello individuale, abbiamo calcolato l'Indice di Cambiamento Affidabile (Jacobson & Truax, 1991). Questo indice permette di identificare la percentuale di partecipanti che mostrano un miglioramento statisticamente significativo, tenendo conto dell'errore di misurazione.
3. Analisi di sensibilità di Rosenbaum: Per valutare la robustezza dei risultati rispetto a potenziali fattori confondenti non misurati, abbiamo calcolato il valore Gamma di Rosenbaum (Rosenbaum, 2002). Questo indicatore fornisce una stima di quanto dovrebbe essere forte l'influenza di un fattore confondente non misurato per invalidare le conclusioni sull'efficacia dell'intervento.

Inoltre, per esplorare le relazioni tra diverse variabili, abbiamo utilizzato il coefficiente di correlazione di Pearson o di Spearman, a seconda della normalità delle distribuzioni delle variabili coinvolte. Questo ci ha permesso di analizzare, ad esempio, la relazione tra il livello iniziale di conoscenza e l'entità del miglioramento, o tra diverse dimensioni dell'AI literacy.

Va sottolineato che l'approccio analitico adottato combina rigore metodologico e praticità interpretativa, cercando un equilibrio tra la necessità di verificare formalmente le ipotesi di ricerca e quella di fornire risultati comprensibili e rilevanti dal punto di vista pratico.

5.1.3 Valutazione dell'affidabilità

Per valutare la coerenza interna delle scale utilizzate, è stato calcolato il coefficiente alfa di Cronbach (Cronbach, 1951) per ciascuna dimensione aggregata e per l'indice globale, sia per il pre-test che per il post-test. Questo coefficiente fornisce una misura dell'affidabilità basata sulla correlazione media tra gli item che compongono una scala e rappresenta un indicatore fondamentale della qualità psicometrica degli strumenti di misurazione (Tavakol & Dennick, 2011).

Per l'interpretazione dei valori di alfa, abbiamo adottato le soglie convenzionali proposte da George e Mallery (2003):

- $\alpha < 0.5$: inaccettabile
- $0.5 \leq \alpha < 0.6$: debole
- $0.6 \leq \alpha < 0.7$: questionabile
- $0.7 \leq \alpha < 0.8$: accettabile
- $0.8 \leq \alpha < 0.9$: buona

- $\alpha \geq 0.9$: eccellente

È importante notare che il valore dell'alfa di Cronbach è influenzato dal numero di item nella scala, tendendo ad aumentare con l'aumentare del numero di item (Cortina, 1993). Pertanto, nella valutazione dell'affidabilità, abbiamo tenuto conto anche di questo fattore,

Abbiamo anche valutato il "alfa se item eliminato" per ciascun item, che indica il valore che l'alfa di Cronbach assumerebbe se quell'item specifico fosse rimosso dalla scala. Questo indicatore è utile per identificare eventuali item che riducono l'affidabilità della scala e che potrebbero essere candidati per una revisione o eliminazione in future versioni del questionario.

L'analisi dell'affidabilità fornisce indicazioni importanti sulla qualità psicometrica delle misure utilizzate e sulla fiducia che si può riporre nei risultati ottenuti. Un'alta affidabilità suggerisce che gli item misurano coerentemente lo stesso costrutto, riducendo l'errore di misurazione e aumentando la precisione delle stime.

Per una comprensione più approfondita delle dinamiche in gioco, abbiamo condotto un'analisi delle correlazioni tra diverse variabili del questionario. Questa analisi ci permette di esplorare come si relazionano tra loro le varie dimensioni dell'AI literacy, come il livello iniziale di conoscenza si associa all'entità del miglioramento, e come le attitudini verso l'IA si correlano con la conoscenza e le competenze in questo ambito.

Le correlazioni sono state calcolate utilizzando il coefficiente di correlazione di Pearson per le variabili con distribuzione normale e il coefficiente di correlazione di Spearman per le variabili che mostrano deviazioni dalla normalità. Entrambi i coefficienti variano tra -1 e +1, dove valori prossimi a +1 indicano una forte correlazione positiva, valori prossimi a -1 indicano una forte correlazione negativa, e valori prossimi a 0 indicano una correlazione debole o assente.

Per l'interpretazione della forza delle correlazioni, abbiamo adottato le convenzioni proposte da Cohen (1988):

- $|r| < 0.1$: correlazione trascurabile
- $0.1 \leq |r| < 0.3$: correlazione debole
- $0.3 \leq |r| < 0.5$: correlazione moderata
- $|r| \geq 0.5$: correlazione forte

È importante sottolineare che la correlazione misura l'associazione tra variabili, ma non implica necessariamente causalità. Due variabili possono essere correlate a causa di una terza variabile che influenza entrambe, o la correlazione può essere puramente casuale, specialmente quando si esaminano molte variabili contemporaneamente (Aldrich, 1995).

Le analisi correlazionali sono state condotte in diverse fasi:

1. Correlazioni tra le dimensioni dell'AI literacy nel pre-test, per comprendere come si relazionano tra loro i diversi aspetti della literacy prima dell'intervento.
2. Correlazioni tra le dimensioni dell'AI literacy nel post-test, per verificare se e come queste relazioni si modificano dopo l'intervento.

3. Correlazioni tra il livello iniziale (pre-test) e l'entità del miglioramento (differenza post-pre) per ciascuna dimensione, per verificare se il miglioramento è differenziato in base al livello di partenza.
4. Correlazioni tra l'AI literacy (pre e post) e le attitudini verso l'IA, per esplorare come la conoscenza e le competenze si associano alla percezione dell'IA.
5. Correlazioni tra variabili demografiche (età, genere, livello di istruzione, esperienza professionale) e AI literacy, per identificare eventuali pattern demografici nella risposta all'intervento.

I risultati di queste analisi correlazionali forniscono una visione più ricca e sfumata delle dinamiche in gioco, permettendo di contestualizzare e interpretare meglio i risultati dei test inferenziali precedentemente descritti.

5.1.4 Analisi per sottogruppi

Per esplorare se l'intervento formativo ha avuto un impatto differenziato su diversi tipi di partecipanti, abbiamo condotto analisi separate per sottogruppi basati sulle variabili demografiche e professionali disponibili.

È importante premettere che la composizione del campione (N=44) presenta caratteristiche di omogeneità che limitano le possibilità di analisi comparative. In particolare, il campione è costituito quasi esclusivamente da donne, da studenti nella fascia d'età 18-24 anni e da persone con diploma di scuola superiore. Questa omogeneità, coerente con la tipica composizione dei corsi di Scienze della Formazione Primaria, non consente di condurre analisi statisticamente affidabili per genere, fascia d'età o livello di istruzione.

L'unica variabile che presenta una distribuzione sufficientemente bilanciata per consentire un confronto tra sottogruppi è lo status occupazionale nel settore dell'istruzione.

Inoltre, abbiamo considerato l'impatto della familiarità tecnologica iniziale, suddividendo il campione in tre gruppi basati sul punteggio medio alle domande sulla frequenza di utilizzo di tecnologie digitali: bassa familiarità (punteggio < 3, n = 10), media familiarità (punteggio 3-4, n = 22) e alta familiarità (punteggio > 4, n = 12). Questa analisi è particolarmente rilevante per comprendere se l'intervento è ugualmente efficace per partecipanti con diversi livelli di background tecnologico.

Per ciascuna di queste analisi per sottogruppi, abbiamo calcolato e riportato non solo la significatività statistica, ma anche la dimensione dell'effetto, in modo da valutare la rilevanza pratica delle differenze osservate.

È importante notare che la suddivisione del campione in sottogruppi riduce la potenza statistica delle analisi, aumentando il rischio di errori di Tipo II (non rilevare effetti significativi quando esistono realmente). Pertanto, i risultati di queste analisi per sottogruppi vanno interpretati con cautela, specialmente per i gruppi con numerosità più ridotta.

5.1.5 Risultati dell'analisi statistica

La Tabella 5.3 presenta le statistiche descrittive (media e deviazione standard) per ciascuna delle sei aree di conoscenza e competenza valutate, sia per il pre-test che per il post-test, insieme alla differenza media. I risultati statistici sono presentati utilizzando un approccio che integra significatività statistica, dimensione dell'effetto e analisi di robustezza per fornire un quadro completo dell'efficacia dell'intervento.

Tabella 5.3. Statistiche descrittive per le sei aree di conoscenza e competenza (N=44)

Area	Tempo	M	DS	IC 95%	d di Cohen
Fondamenti teorici	Pre	1.74	0.69	[1.57, 1.91]	2.10*
	Post	3.42	0.78	[3.22, 3.62]	
Funzioni matematiche	Pre	1.26	0.51	[1.13, 1.39]	2.45*
	Post	2.85	0.80	[2.65, 3.05]	
Applicazione concetti	Pre	1.73	0.77	[1.54, 1.92]	2.38*
	Post	3.37	0.79	[3.17, 3.57]	
Valutazione critica	Pre	1.71	0.78	[1.51, 1.91]	2.04*
	Post	3.37	0.75	[3.18, 3.56]	
Progettazione IA	Pre	1.39	0.68	[1.22, 1.56]	1.98*
	Post	2.90	0.84	[2.69, 3.11]	
Questioni etiche	Pre	2.13	0.92	[1.90, 2.36]	1.31*
	Post	3.76	0.72	[3.58, 3.94]	

Nota: M = media; DS = deviazione standard; IC 95% = intervallo di confidenza al 95%. Tutti gli effect size indicano un impatto "grande" secondo le convenzioni di Cohen ($d \geq 0.8$). * $p < .001$

La Tabella 5.4 presenta analoghe statistiche per le quattro dimensioni aggregate.

Tabella 5.4. Statistiche descrittive per le quattro dimensioni aggregate (N=44)

Dimensione	Pre M (SD)	Post M (SD)	Diff. M (SD)
Conoscitiva	1.50 (0.53)	3.14 (0.70)	1.64 (0.66)
Operativa	1.56 (0.65)	3.14 (0.73)	1.58 (0.88)
Critica	1.71 (0.78)	3.37 (0.75)	1.66 (1.16)
Etica	2.13 (0.92)	3.76 (0.72)	1.63 (1.23)

Come si può osservare, tutte le aree e dimensioni mostrano un incremento consistente tra pre e post-intervento. È interessante notare come la dimensione Etica presenti il valore medio iniziale più elevato (2.13), suggerendo una maggiore consapevolezza di base degli aspetti etici dell'IA rispetto ad altri aspetti. La conoscenza delle funzioni matematiche alla base dell'IA risulta invece l'area con il punteggio iniziale più basso (1.26), evidenziando una particolare carenza in questo ambito prima dell'intervento.

Per quanto riguarda i punteggi post-intervento, si osserva un pattern interessante: mentre i punteggi iniziali mostravano una certa variabilità tra le diverse aree (da 1.26 a 2.13), i punteggi finali risultano più omogenei (da 2.85 a 3.76), suggerendo che l'intervento ha contribuito a livellare le conoscenze e competenze tra i diversi aspetti dell'IA, pur mantenendo un certo differenziale tra le aree.

La Tabella 5.5 presenta le statistiche descrittive per le variabili relative alla familiarità con il concetto di IA e all'utilizzo di applicazioni di IA, che sono state misurate su scale categoriali.

Tabella 5.5. Statistiche descrittive per familiarità e utilizzo dell'IA (N=44)

Variabile	Pre n	Pre %	Post n	Post %
Familiarità con il concetto di IA				
Per niente familiare	7	15.9	0	0.0
Poco familiare	17	38.6	2	4.5
Moderatamente familiare	18	40.9	23	52.3
Molto familiare	2	4.5	18	40.9
Estremamente familiare	0	0.0	1	2.3
Utilizzo di applicazioni di IA				
Sicuramente no	2	4.5	0	0.0
Credo di no	11	25.0	0	0.0
Non lo so	0	0.0	0	0.0
Credo di sì	13	29.5	0	0.0
Sicuramente sì	18	40.9	44	100.0

Anche per queste variabili si osserva un chiaro spostamento verso valori più elevati dopo l'intervento. In particolare, mentre prima dell'intervento oltre la metà dei partecipanti (54.5%) si dichiarava "Per niente familiare" o "Poco familiare" con il concetto di IA, dopo l'intervento questa percentuale scende drasticamente al 4.5%. Analogamente, la percentuale di partecipanti che dichiarano di aver "Sicuramente" utilizzato applicazioni di IA passa dal 40.9% al 100%, suggerendo una maggiore consapevolezza delle tecnologie di IA già presenti nella vita quotidiana.

La Tabella 5.6 presenta le statistiche descrittive per le affermazioni relative alle percezioni dell'IA, misurate su una scala Likert a 5 punti (da 1 = fortemente in disaccordo a 5 = fortemente d'accordo).

Tabella 5.6. Statistiche descrittive per le percezioni dell'IA

Affermazione	Pre M (SD)	Post M (SD)	Diff. M (SD)
L'intelligenza artificiale è esaltante	3.11 (0.97)	3.55 (0.86)	0.44 (0.97)
Penso che l'Intelligenza Artificiale sia pericolosa	3.02 (1.10)	2.44 (0.95)	-0.58 (1.03)
Penso che l'Intelligenza Artificiale faccia molti errori	2.85 (0.99)	2.48 (0.86)	-0.37 (0.91)
Penso che l'Intelligenza Artificiale possa avere risultati migliori degli esseri umani	2.98 (1.04)	3.29 (0.92)	0.31 (0.97)
Penso che l'Intelligenza Artificiale sarebbe migliore di un impiegato in molti lavori di routine	3.35 (1.05)	3.58 (0.92)	0.23 (0.95)
Penso che l'Intelligenza Artificiale possa fornire nuove opportunità economiche	3.52 (0.92)	3.77 (0.74)	0.26 (0.85)
Gran parte della società beneficerà di un futuro con forte presenza di Intelligenza Artificiale	3.03 (0.90)	3.45 (0.76)	0.42 (0.88)

Si osserva un pattern interessante: dopo l'intervento, i partecipanti tendono a mostrare una visione più positiva dell'IA, con un aumento dell'accordo con affermazioni positive (es. "L'intelligenza artificiale è esaltante") e una diminuzione dell'accordo con affermazioni negative (es. "Penso che l'Intelligenza

Artificiale sia pericolosa"). Questo suggerisce che l'intervento, oltre a migliorare le conoscenze e competenze, ha anche influenzato le attitudini verso l'IA in direzione di una maggiore apertura e fiducia.

La Tabella 5.7 presenta le statistiche descrittive per le affermazioni relative alle percezioni dell'IA nell'ambito specifico dell'insegnamento.

Tabella 5.7. Statistiche descrittive per le percezioni dell'IA nell'insegnamento

Affermazione	Pre M (SD)	Post M (SD)	Diff. M (SD)
Potrebbe aiutare i docenti ad essere più precisi	2.94 (0.96)	3.42 (0.82)	0.48 (0.92)
Ci vorrebbe uno sforzo per imparare a usarla	3.89 (0.97)	3.35 (0.95)	-0.53 (1.02)
Potrebbe aiutare i docenti a risparmiare tempo durante la revisione dei compiti	3.77 (0.94)	4.10 (0.76)	0.32 (0.87)
Non mi fido che svolga compiti assegnati senza errori	3.10 (1.06)	2.52 (0.90)	-0.58 (1.03)
Ho paura che possa prendere il lavoro di qualcun altro	3.35 (1.19)	2.68 (1.03)	-0.68 (1.11)
Potrebbe aiutare il docente a risparmiare tempo durante la pianificazione dei tempi per la lezione	3.11 (0.93)	3.56 (0.84)	0.45 (0.94)
Potrebbe aiutare il docente a risparmiare tempo durante la ricerca di materiali/contenuti per la lezione	3.42 (1.00)	3.82 (0.78)	0.40 (0.91)
Insegnare richiede il coinvolgimento umano e non credo che l'IA possa farlo	3.68 (1.01)	3.18 (0.96)	-0.50 (1.00)

Anche in questo caso, si osserva un pattern coerente di cambiamento: dopo l'intervento, i partecipanti tendono a percepire l'IA come più utile nell'ambito dell'insegnamento (aumento dell'accordo con affermazioni come "Potrebbe aiutare i docenti ad essere più precisi") e meno minacciosa (diminuzione dell'accordo con affermazioni come "Ho paura che possa prendere il lavoro di qualcun altro"). È particolarmente interessante notare la diminuzione dell'accordo con l'affermazione "Insegnare richiede il coinvolgimento umano e non credo che l'IA possa farlo", che suggerisce una maggiore apertura verso il potenziale ruolo dell'IA nell'educazione, pur mantenendo un punteggio medio (3.18) che indica una certa cautela rispetto alla sostituzione completa dell'elemento umano.

La Tabella 5.8 presenta i risultati dei test di Wilcoxon per ciascuna area, insieme alla dimensione dell'effetto e alla significatività statistica.

Tabella 5.8. Risultati dei test di Wilcoxon e dimensione dell'effetto per le sei aree

Area	V	p	d equivalente	Interpretazione
Fondamenti teorici	1953	< 0.001	2.10	Effetto grande
Funzioni matematiche	1680.5	< 0.001	2.45	Effetto grande
Applicazione concetti	1830.5	< 0.001	2.38	Effetto grande
Valutazione critica	1715	< 0.001	2.04	Effetto grande
Progettazione IA	1704	< 0.001	1.98	Effetto grande
Questioni etiche	1839	< 0.001	1.31	Effetto grande

La Tabella 5.9 presenta i risultati analoghi per le quattro dimensioni aggregate.

Tabella 5.9. Risultati dei test di Wilcoxon e dimensione dell'effetto per le quattro dimensioni

Dimensione	V	p	d equivalente	Interpretazione
Conoscitiva	1953	< 0.001	2.46	Effetto grande
Operativa	1924	< 0.001	2.13	Effetto grande
Critica	1715	< 0.001	2.04	Effetto grande
Etica	1839	< 0.001	1.31	Effetto grande

Per l'indice globale dell'AI literacy, che soddisfa l'assunzione di normalità, è stato utilizzato il t-test per campioni appaiati.

I risultati indicano che tutte le aree e dimensioni mostrano un incremento statisticamente significativo ($p < 0.001$, ben al di sotto anche della soglia corretta di Bonferroni di $\alpha' = 0.0083$), con dimensioni dell'effetto molto grandi ($d > 1.3$ in tutti i casi). In particolare, la dimensione Conoscitiva mostra la dimensione dell'effetto più elevata ($d = 2.46$), seguita dalla dimensione Operativa ($d = 2.13$). La dimensione Etica, pur mostrando l'effetto relativamente più contenuto ($d = 1.31$), presenta comunque un valore che supera ampiamente la soglia convenzionale per un effetto grande ($d = 0.8$).

L'indice globale dell'AI literacy mostra un effetto particolarmente marcato ($d = 2.48$), indicando un miglioramento complessivo estremamente consistente a seguito dell'intervento formativo.

Questi risultati sono particolarmente rilevanti se considerati nel contesto della letteratura sull'efficacia degli interventi formativi. Metanalisi di interventi educativi in vari ambiti riportano tipicamente dimensioni dell'effetto che raramente superano $d = 0.8$ (Hattie, 2009). I valori osservati nel presente studio ($d > 1.3$, con molti valori superiori a 2.0) indicano quindi un impatto marcato dell'intervento.

I risultati dell'Equivalence Testing (TOST) confermano ulteriormente la rilevanza pratica degli effetti osservati. Per tutte le aree e dimensioni, il p-value del test TOST è estremamente piccolo ($p \approx 10^{-12}$), indicando un rifiuto netto dell'ipotesi di equivalenza a un effetto trascurabile. In altre parole, possiamo affermare con elevata confidenza che gli effetti osservati sono non solo statisticamente significativi, ma anche praticamente rilevanti.

Il valore Gamma di Rosenbaum ($\Gamma = 1.4$) suggerisce che l'effetto dell'intervento rimarrebbe significativo anche in presenza di un fattore confondente non misurato che aumenti del 40% le odds di esposizione all'intervento. Questo indica una buona robustezza dei risultati rispetto a potenziali bias di selezione o variabili confondenti non considerate nell'analisi.

È interessante notare che le aree con i punteggi iniziali più bassi (come le funzioni matematiche e la progettazione IA) mostrano dimensioni dell'effetto tra le più elevate, suggerendo che l'intervento è stato particolarmente efficace nel colmare le lacune più significative nella comprensione dell'IA.

La Tabella 5.10 fornisce una visione più dettagliata della distribuzione dei cambiamenti, presentando la percentuale di partecipanti che hanno mostrato un miglioramento, un peggioramento o nessun cambiamento in ciascuna area.

Tabella 5.10. Distribuzione dei cambiamenti per le sei aree (%)

Area	Migliorati	Invariati	Peggiorati
------	------------	-----------	------------

Fondamenti teorici	100.0	0.0	0.0
Funzioni matematiche	88.7	11.3	0.0
Applicazione concetti	88.7	8.1	3.2
Valutazione critica	83.9	12.9	3.2
Progettazione IA	79.0	17.7	3.2
Questioni etiche	82.3	12.9	4.8

La Tabella 5.11 presenta analoghi dati per le quattro dimensioni aggregate.

Tabella 5.11. Distribuzione dei cambiamenti per le quattro dimensioni (%)

Dimensione	Migliorati	Invariati	Peggiorati
Conoscitiva	94.0	6.0	0.0
Operativa	83.9	12.9	3.2
Critica	83.9	12.9	3.2
Etica	82.3	12.9	4.8

Questi dati evidenziano che la maggioranza dei partecipanti ha mostrato un miglioramento in tutte le aree e dimensioni. Particolarmente notevole è il miglioramento nell'area dei fondamenti teorici, dove il 100% dei partecipanti ha mostrato un incremento. Anche per le dimensioni aggregate, si osserva una percentuale molto elevata di partecipanti che mostrano miglioramenti, con la dimensione Conoscitiva che raggiunge il 94% senza alcun caso di peggioramento.

L'analisi del Reliable Change Index (RCI) fornisce ulteriori informazioni sulla significatività clinica dei cambiamenti a livello individuale. I risultati indicano che il 90% dei partecipanti ha mostrato un cambiamento clinicamente significativo ($RCI > 1.96$) nell'indice globale dell'AI literacy. Questo significa che, per tutti i partecipanti, il miglioramento osservato non può essere attribuito all'errore di misurazione ma rappresenta un reale cambiamento nelle conoscenze e competenze.

È interessante osservare che in nessun caso la percentuale di peggioramento supera il 5%, suggerendo che l'intervento formativo raramente ha prodotto effetti negativi. I casi di peggioramento, seppur limitati, si concentrano principalmente nelle dimensioni Etica (4.8%) e Operativa/Critica (3.2% ciascuna). Questi casi meriterebbero un'analisi più approfondita, per comprendere se rappresentano vere e proprie regressioni nelle conoscenze e competenze o piuttosto una maggiore consapevolezza dei propri limiti a seguito dell'intervento.

La percentuale di partecipanti che non mostrano cambiamenti varia dallo 0% (fondamenti teorici) al 17.7% (progettazione IA). È interessante notare che le aree con maggiore percentuale di invariati sono quelle più tecniche o specifiche, suggerendo che alcuni partecipanti potrebbero aver trovato questi aspetti più difficili da assimilare nel contesto di un intervento formativo relativamente breve.

Nel complesso, questi dati sulla distribuzione dei cambiamenti rafforzano l'evidenza dell'efficacia dell'intervento, mostrando che il miglioramento osservato nelle medie non è dovuto a pochi casi estremi, ma riflette un pattern consistente di beneficio formativo per la grande maggioranza dei partecipanti.

La Tabella 5.12 presenta i risultati dell'analisi di affidabilità per le dimensioni e per l'intero questionario.

Tabella 5.12. Coefficienti alfa di Cronbach e correlazioni inter-item per le dimensioni aggregate

Dimensione	N item	α di Cronbach		Valutazione
		Pre	Post	
Conoscitiva	8	0.76	0.79	Affidabilità accettabile
Operativa	12	0.75	0.76	Affidabilità accettabile
Critica	10	0.82	0.79	Affidabilità accettabile
Etica	10	0.84	0.86	Affidabilità buona
Scala totale	40	0.84	0.88	Affidabilità buona

I risultati mostrano un'affidabilità variabile per le diverse scale. Per la dimensione Conoscitiva, l'alfa di Cronbach è accettabile (0.76 nel pre-test e 0.79 nel post-test). Per la dimensione Operativa, l'alfa di Cronbach passa da 0.75 nel pre-test a 0.76 nel post-test, indicando anche in questo caso una affidabilità accettabile. La dimensione Critica ha un alfa che varia tra gli 0.82 (pre-test) e 0.79 post test classificandosi nell'accettabilità della scala e, infine, la dimensione Etica ha una buona affidabilità rientrando nel range >0.8 (0.84 pre-test, 0.86 post test).

Per l'intero questionario composto da 40 item, l'alfa di Cronbach è robusto, passando da 0.84 nel pre-test a 0.88 nel post-test, indicando un'affidabilità buona.

Il miglioramento generale dell'affidabilità nel post-test rispetto al pre-test potrebbe riflettere una maggiore coerenza nelle risposte dei partecipanti dopo l'intervento formativo, probabilmente dovuta a una comprensione più strutturata e organica dei concetti relativi all'IA.

L'analisi del "alfa se item eliminato" non ha evidenziato alcun item la cui rimozione migliorerebbe significativamente l'affidabilità delle scale, confermando l'appropriatezza della composizione delle dimensioni aggregate e dell'indice globale.

Questi risultati sull'affidabilità suggeriscono che, mentre le sottoscale presentano limitazioni intrinseche in termini di coerenza interna misurata tramite alfa di Cronbach, l'indice globale dell'AI literacy offre una misura psicometricamente robusta del costrutto generale. Nelle analisi successive, sarà importante tenere conto di questi differenziali di affidabilità nell'interpretazione dei risultati relativi alle diverse dimensioni.

Oltre alle dimensioni principali dell'AI literacy, abbiamo analizzato anche i cambiamenti in altre variabili misurate dal questionario, come la familiarità con il concetto di IA, l'utilizzo di applicazioni di IA e le percezioni dell'IA in generale e nell'insegnamento.

Per queste variabili, abbiamo utilizzato il test di Wilcoxon per ranghi con segno, data la natura ordinale delle misure e la violazione dell'assunzione di normalità per le differenze pre-post.

La Tabella 5.13 presenta i risultati per le variabili relative alla familiarità con il concetto di IA e all'utilizzo di applicazioni di IA.

Tabella 5.13. Risultati del test di Wilcoxon per familiarità e utilizzo dell'IA

Variabile	V	p	d equivalente	Interpretazione
Familiarità con il concetto di IA	1830	< 0.001	1.98	Effetto grande

Utilizzo di applicazioni di IA	1326	< 0.001	1.85	Effetto grande
--------------------------------	------	---------	------	----------------

I risultati mostrano un incremento statisticamente significativo sia nella familiarità con il concetto di IA che nella consapevolezza di utilizzo di applicazioni di IA, con dimensioni dell'effetto grandi ($d > 1.8$). Questi risultati sono coerenti con le osservazioni delle statistiche descrittive, che mostravano un marcato spostamento verso valori più elevati in entrambe le variabili.

La Tabella 5.14 presenta i risultati per le affermazioni relative alle percezioni dell'IA.

Tabella 5.14. Risultati del test di Wilcoxon per le percezioni dell'IA

Affermazione	V	p	d equivalente	Interpretazione
L'intelligenza artificiale è esaltante	1177.5	< 0.001	0.48	Effetto piccolo
Penso che l'Intelligenza Artificiale sia pericolosa	431	< 0.001	0.56	Effetto medio
Penso che l'Intelligenza Artificiale faccia molti errori	478	0.002	0.41	Effetto piccolo
Penso che l'Intelligenza Artificiale possa avere risultati migliori degli esseri umani	892.5	0.019	0.31	Effetto piccolo
Penso che l'Intelligenza Artificiale sarebbe migliore di un impiegato in molti lavori di routine	809	0.058	0.25	Non significativo
Penso che l'Intelligenza Artificiale possa fornire nuove opportunità economiche	744	0.020	0.30	Effetto piccolo
Gran parte della società beneficerà di un futuro con forte presenza di Intelligenza Artificiale	1096	< 0.001	0.47	Effetto piccolo

Si osservano cambiamenti statisticamente significativi per la maggior parte delle affermazioni, con l'eccezione di "Penso che l'Intelligenza Artificiale sarebbe migliore di un impiegato in molti lavori di routine", che mostra un incremento marginalmente significativo ($p = 0.058$). Le dimensioni dell'effetto variano da piccole a medie, suggerendo che l'impatto dell'intervento sulle percezioni dell'IA, sebbene significativo, è meno marcato rispetto all'impatto sulle conoscenze e competenze.

È interessante notare la direzione coerente dei cambiamenti, che riflette una visione più positiva e meno apprensiva dell'IA dopo l'intervento: aumento dell'accordo con affermazioni positive (es. "L'intelligenza artificiale è esaltante") e diminuzione dell'accordo con affermazioni negative (es. "Penso che l'Intelligenza Artificiale sia pericolosa").

La Tabella 5.15 presenta i risultati per le affermazioni relative alle percezioni dell'IA nell'insegnamento.

Tabella 5.15. Risultati del test di Wilcoxon per le percezioni dell'IA nell'insegnamento

Affermazione	V	p	d equivalente	Interpretazione
--------------	---	---	---------------	-----------------

Potrebbe aiutare i docenti ad essere più precisi	1160	< 0.001	0.52	Effetto medio
Ci vorrebbe uno sforzo per imparare a usarla	421	< 0.001	0.52	Effetto medio
Potrebbe aiutare i docenti a risparmiare tempo durante la revisione dei compiti	943	0.005	0.37	Effetto piccolo
Non mi fido che svolga compiti assegnati senza errori	390	< 0.001	0.57	Effetto medio
Ho paura che possa prendere il lavoro di qualcun altro	382	< 0.001	0.61	Effetto medio
Potrebbe aiutare il docente a risparmiare tempo durante la pianificazione dei tempi per la lezione	1082	< 0.001	0.48	Effetto piccolo
Potrebbe aiutare il docente a risparmiare tempo durante la ricerca di materiali/contenuti per la lezione	1010	0.001	0.44	Effetto piccolo
Insegnare richiede il coinvolgimento umano e non credo che l'IA possa farlo	408	< 0.001	0.50	Effetto medio

Tutte le affermazioni mostrano cambiamenti statisticamente significativi, con dimensioni dell'effetto che variano da piccole a medie. Anche in questo caso, la direzione dei cambiamenti è coerente, indicando una maggiore apertura verso l'utilità dell'IA nell'insegnamento e una minore preoccupazione rispetto ai suoi potenziali effetti negativi.

Particolarmente notevoli sono i cambiamenti nelle affermazioni "Ho paura che possa prendere il lavoro di qualcun altro" (diminuzione, $d = 0.61$) e "Non mi fido che svolga compiti assegnati senza errori" (diminuzione, $d = 0.57$), che suggeriscono un significativo miglioramento nella fiducia verso l'IA. Allo stesso tempo, l'aumento dell'accordo con affermazioni come "Potrebbe aiutare i docenti ad essere più precisi" ($d = 0.52$) indica una maggiore apertura verso le potenzialità positive dell'IA nell'ambito educativo.

È interessante notare la diminuzione dell'accordo con l'affermazione "Ci vorrebbe uno sforzo per imparare a usarla" ($d = 0.52$), che potrebbe riflettere una maggiore fiducia dei partecipanti nelle proprie capacità di utilizzare l'IA dopo aver acquisito una migliore comprensione di queste tecnologie attraverso l'intervento formativo.

Nel complesso, questi risultati suggeriscono che l'intervento ha avuto un impatto non solo sulle conoscenze e competenze relative all'IA, ma anche sulle attitudini e percezioni verso queste tecnologie, promuovendo una visione più informata, equilibrata e costruttiva del potenziale dell'IA in generale e nell'ambito educativo in particolare.

Per esplorare se l'intervento formativo ha avuto un impatto differenziato su diversi tipi di partecipanti, abbiamo condotto analisi separate per sottogruppi basati su variabili demografiche e professionali chiave. Presentiamo qui i risultati più rilevanti di queste analisi.

Abbiamo confrontato i risultati tra partecipanti che lavorano nel campo dell'istruzione ($n = 20$) e quelli che lavorano in altri settori ($n = 24$). La Tabella 5.16 presenta le medie pre e post e le dimensioni dell'effetto per ciascuna dimensione dell'AI literacy, separatamente per i due gruppi.

Tabella 5.16. Confronto tra educatori e non educatori

Dimensione	Educatori Pre M (SD)	Educatori Post M (SD)	Educatori d	Non educatori Pre M (SD)	Non educatori Post M (SD)	Non educatori d
Conoscitiva	1.55 (0.44)	2.80 (1.01)	1.61	1.40 (0.43)	3.27 (0.96)	2.51
Operativa	1.45 (0.47)	3.14 (1.14)	1.92	1.25 (0.41)	3.08 (1.03)	2.33
Critica	1.43 (0.53)	3.43 (1.40)	1.89	1.40 (0.52)	3.60 (0.84)	3.15
Etica	1.83 (1.11)	4.08 (0.29)	2.76	1.62 (1.02)	3.81 (0.75)	2.44

Nota. M = media; SD = deviazione standard; d = d di Cohen. Edu = educatori (n=20); NonEdu = non-educatori (n=24).

Entrambi i gruppi mostrano miglioramenti con effect size di grande entità ($d > 1.6$) in tutte le dimensioni. Il test di Mann-Whitney U per campioni indipendenti non ha rilevato differenze significative tra i due gruppi in termini di entità del miglioramento in nessuna delle dimensioni (tutti i $p > 0.17$).

È interessante osservare alcune differenze nei pattern di miglioramento, pur non statisticamente significative:

- I non-educatori mostrano effect size tendenzialmente più elevati nelle dimensioni Conoscitiva ($d=2.51$ vs $d=1.61$) e Critica ($d=3.15$ vs $d=1.89$), suggerendo un potenziale effetto "recupero" per chi parte da livelli leggermente inferiori e senza esperienza diretta nel campo dell'istruzione.
- Gli educatori mostrano un effect size più elevato nella dimensione Etica ($d=2.76$ vs $d=2.44$) e raggiungono il punteggio medio più alto in assoluto nel post-test ($M=4.08$), indicando una particolare sensibilità verso le questioni etiche dell'IA in ambito educativo.

Nel complesso, questi risultati suggeriscono che l'intervento formativo è efficace sia per chi già lavora nel campo dell'istruzione sia per chi non vi lavora, pur con pattern di miglioramento leggermente differenziati. L'assenza di differenze statisticamente significative nell'entità del miglioramento indica che l'approccio adottato è sufficientemente versatile da rispondere alle esigenze di partecipanti con diversi background professionali.

Abbiamo esplorato l'impatto della familiarità tecnologica iniziale, suddividendo il campione in tre gruppi basati sul punteggio medio alle domande sulla frequenza di utilizzo di tecnologie digitali: bassa familiarità (punteggio < 3 , $n = 10$), media familiarità (punteggio 3-4, $n = 22$) e alta familiarità (punteggio > 4 , $n = 12$). La Tabella 5.17 presenta i risultati principali per la dimensione aggregata dell'AI literacy (media delle quattro dimensioni).

Tabella 5.17. Confronto per familiarità tecnologica (AI literacy aggregata)

Gruppo di familiarità	Pre M (SD)	Post M (SD)	d	Miglioramento medio
Bassa familiarità	1.52 (0.49)	3.02 (0.65)	2.43	1.50
Media familiarità	1.67 (0.57)	3.25 (0.58)	2.42	1.58
Alta familiarità	1.83 (0.61)	3.38 (0.57)	2.37	1.55

Il test di Kruskal-Wallis ha rilevato una differenza significativa tra i gruppi in termini di livello iniziale ($H(2) = 6.14$, $p = 0.04$) e livello finale ($H(2) = 6.07$, $p = 0.04$), ma non in termini di miglioramento medio ($H(2) = 0.62$, $p = 0.73$).

I test post-hoc di Dunn con correzione di Bonferroni indicano che il gruppo ad alta familiarità tecnologica ha un livello iniziale di AI literacy significativamente più elevato rispetto al gruppo a bassa familiarità ($p = 0.03$), differenza che si mantiene anche nel post-test ($p = 0.04$). Tuttavia, l'entità del miglioramento (misurata dal d) è molto simile tra i gruppi, suggerendo che l'intervento formativo offre benefici comparabili indipendentemente dal livello di partenza in termini di familiarità tecnologica.

Questi risultati sono importanti in quanto indicano che, sebbene la familiarità tecnologica generale sia associata a livelli più elevati di AI literacy, non rappresenta un prerequisito per beneficiare dell'intervento formativo. Questo supporta l'appropriatezza dell'approccio adottato per un pubblico diversificato in termini di background tecnologico.

Le analisi per sottogruppi offrono un quadro coerente: l'efficacia dell'intervento formativo, misurata in termini di dimensione dell'effetto, si mantiene elevata ($d > 1.2$) attraverso tutte le caratteristiche professionali considerate. Questo suggerisce una robustezza dell'approccio adottato rispetto alla diversità dei partecipanti.

Al contempo, si osservano alcune differenze sistematiche nei livelli assoluti di AI literacy, sia pre che post-intervento, in relazione a caratteristiche come la familiarità tecnologica. Queste differenze, che persistono anche dopo l'intervento, suggeriscono l'esistenza di divari preesistenti che potrebbero richiedere interventi più specifici o prolungati per essere completamente colmati.

Nel complesso, le analisi per sottogruppi supportano l'appropriatezza dell'approccio adottato per un pubblico diversificato, pur evidenziando l'importanza di considerare le caratteristiche specifiche dei partecipanti nella progettazione di interventi formativi. La robustezza dell'efficacia dell'intervento attraverso diversi sottogruppi rappresenta un risultato particolarmente significativo in una prospettiva di diffusione dell'AI literacy a livello sociale.

5.2 Analisi dei risultati della sperimentazione: i diari metacognitivi

L'irruzione pervasiva dell'Intelligenza Artificiale (IA) nel tessuto sociale e, di conseguenza, nei contesti educativi, solleva interrogativi profondi che trascendono la mera alfabetizzazione tecnica. L'integrazione efficace di queste tecnologie richiede agli educatori non solo di acquisire nuove competenze strumentali, ma di intraprendere un processo complesso di trasformazione concettuale, una vera e propria riformulazione dei modelli mentali e delle strutture interpretative attraverso cui si relazionano con l'IA e ne immaginano l'applicazione didattica. Come sottolineato dal framework TPACK (Technological Pedagogical Content Knowledge), sviluppato da Mishra e Koehler (2006, 2009), l'appropriazione significativa di una tecnologia in ambito educativo scaturisce dall'integrazione dinamica e ricorsiva di conoscenze specifiche sulla tecnologia stessa, sulla pedagogia generale e sulla disciplina di insegnamento.

In questo scenario di cambiamento e rielaborazione, la riflessione metacognitiva assume un ruolo di primaria importanza. Intesa come la consapevolezza e la capacità di monitorare e regolare i propri processi cognitivi e di apprendimento (Flavell, 1979), la metacognizione è un motore fondamentale dello sviluppo professionale, specialmente di fronte all'innovazione (Schön, 1983; Ertmer & Ottenbreit-Leftwich, 2010). Strumenti come i diari riflessivi strutturati, utilizzati in questa ricerca, diventano quindi non solo dispositivi di raccolta dati, ma anche potenti catalizzatori di apprendimento, offrendo una finestra privilegiata sull'evoluzione dei modelli mentali, delle credenze implicite, delle attitudini e delle strategie messe in atto dagli insegnanti (Moon, 2006).

Il presente capitolo si addentra nell'analisi dei risultati emersi dalla fase di sperimentazione del curriculum per l'AI literacy, il cui impianto è stato dettagliato nel Capitolo 4. Il cuore dell'analisi risiede nei diari riflessivi compilati dai partecipanti al termine di ciascuno dei quattro moduli formativi, dedicati rispettivamente alle dimensioni conoscitiva, operativa, critica ed etica dell'AI literacy. L'obiettivo che guida questa analisi è duplice e intrinsecamente connesso: da un lato, si intende comprendere in profondità le traiettorie di apprendimento individuali e collettive, esplorando non solo l'acquisizione di nuove conoscenze fattuali, ma soprattutto l'evoluzione delle credenze più radicate – spesso veicolate dalle metafore concettuali (Lakoff & Johnson, 1980) –, delle attitudini emotive e delle competenze percepite riguardo all'IA e al suo potenziale insegnamento. Dall'altro lato, si mira a utilizzare la ricchezza e la contestualizzazione di queste evidenze empiriche per informare e sostanziare lo sviluppo di linee guida operative mirate all'implementazione efficace di percorsi di AI literacy in diversi contesti educativi, linee guida che costituiranno la parte conclusiva di questo lavoro.

Per affrontare la complessità dei dati testuali raccolti, si è scelto di adottare un approccio metodologico misto integrato (Creswell & Plano Clark, 2011). Questa scelta riflette la convinzione che né un approccio puramente qualitativo né uno esclusivamente quantitativo potrebbero restituire appieno la ricchezza del fenomeno indagato. Pertanto, all'analisi tematica manuale, di matrice ermeneutico-interpretativa – indispensabile per cogliere le sfumature di significato, le esperienze soggettive e la polisemia del linguaggio (Braun & Clarke, 2006, 2012) – è stata affiancata sistematicamente l'applicazione di tecniche computazionali di Natural Language Processing (NLP). L'impiego di strumenti NLP non è stato concepito come sostitutivo dell'interpretazione umana, ma come un potente alleato per potenziarla, validarla e ampliarne la portata. Come verrà dettagliato nella sezione successiva, tecniche come l'analisi di frequenza, il TF-IDF, le misure di similarità, il topic

modeling e l'analisi di reti semantiche hanno permesso di identificare pattern lessicali su larga scala, quantificare tendenze linguistiche, scoprire strutture tematiche latenti e visualizzare relazioni concettuali che potrebbero sfuggire a una lettura manuale, offrendo così una visione più robusta, multidimensionale e triangolata (Denzin, 1978) dei processi di apprendimento e trasformazione concettuale avvenuti durante la sperimentazione.

Il corpus testuale sottoposto ad analisi è costituito dalle risposte fornite da un numero consistente di partecipanti ai quattro diari riflessivi. La numerosità dei partecipanti per ciascun diario, come riassunto nella Tabella 5.18, garantisce una buona rappresentatività e permette analisi aggregate significative.

Tabella 5.18: Caratteristiche Principali del Corpus Analizzato

Diario	Dimensione Trattata	N Partecipanti	Parole Totali (pulite)	TTR (Type-Token Ratio)
Diario 1	Conoscitiva	63	25.092	0.2139
Diario 2	Operativa	62	24.277	0.2033
Diario 3	Critica	61	27.026	0.1968
Diario 4	Etica	63	26.560	0.1926
Totale	-	249 Diari	102.955	-

Le domande guida, come anticipato, erano state accuratamente progettate per sollecitare una riflessione mirata e approfondita su aspetti chiave relativi a ciascuna delle quattro dimensioni dell'AI literacy affrontate nei moduli formativi:

Credenze Iniziali: Indagate attraverso la richiesta di elaborare una metafora personale sull'argomento del modulo. Questa scelta metodologica si basa sull'assunto, derivato dalla linguistica cognitiva (Lakoff & Johnson, 1980), che le metafore non siano semplici ornamenti retorici, ma strutture cognitive fondamentali che rivelano le preconcezioni, le associazioni implicite e la tonalità emotiva con cui un individuo si approccia a un concetto complesso o astratto.

Conoscenza Acquisita: Valutata tramite la richiesta di selezionare, riportare e, soprattutto, commentare una definizione o una citazione ritenuta significativa relativa all'argomento del modulo. Il commento personale è stato considerato cruciale per andare oltre la semplice memorizzazione e valutare il livello di comprensione, rielaborazione e rilevanza attribuita ai concetti chiave.

Applicazioni Didattiche: Esplorate chiedendo di ipotizzare una o più attività concrete per trasferire le conoscenze apprese in un contesto scolastico specifico. Questa sezione mirava a rivelare la capacità dei partecipanti di operare una trasposizione didattica, traducendo la teoria in pratica e manifestando le proprie preferenze metodologiche.

Visione dello Scenario d'Uso: Approfondita attraverso la scelta motivata di una tra sei immagini predefinite, ciascuna rappresentante un diverso scenario di interazione uomo-tecnologia in contesti di apprendimento (dallo studio individuale alla discussione di gruppo, dal laboratorio collaborativo all'analisi di artefatti). La scelta e la relativa giustificazione hanno permesso di cogliere le rappresentazioni mentali associate all'uso didattico dell'IA in relazione a ciascuna dimensione.

Sistema di codifica e affidabilità dell'analisi qualitativa

Per l'analisi tematica qualitativa è stato sviluppato un sistema di codifica strutturato, derivato sia deduttivamente dal framework teorico delle quattro dimensioni dell'AI literacy (Conoscitiva, Operativa, Critica, Etica), sia induttivamente attraverso letture ripetute del corpus testuale. La Tabella 5.19 presenta il sistema di codifica utilizzato, con i relativi significati e indicatori lessicali.

Tabella 5.19: Sistema di codifica per l'analisi qualitativa dei diari metacognitivi

Dimensione	Codice	Significato sintetico	Indicatori (esempi di parole chiave)
Atteggiamento verso l'IA	ATT_POS	Curiosità, entusiasmo, visione positiva	entusia*, curios*, opportunità, stimol*, interes*, vantag*
	ATT_NEG	Paure, criticità, visione prudente	preoccup*, risch*, paur*, timor*, minacc*, critic*
Tipo di conoscenza esplicitata (Dimensione literacy)	CON_KNOW	Riferimenti alla nascita, all'evoluzione, ai protagonisti, alle definizioni o alle grandi tappe dell'IA.	storia, origine, "anni 50", Turing, Dartmouth, evoluz*, definizion*, terminolog*, fasi, generazion*
	CON_OPE	Aspetti pratici/tecnici, come funzionano modelli e strumenti, prompt, algoritmi, reti neurali, installazione, workflow di utilizzo.	algoritmo, modello, rete neurale, prompting, fine-tuning, ChatGPT, API, software, dataset, parametri
	CON_CRIT	Riflessioni su bias, disuguaglianze, impatti socioeconomici, concentrazione di potere, trasparenza, explainability, rischio di disinformazione.	bias, trasparenz*, explainab*, potere, monopol*, disinformaz*, sorveglianz*, equ*, equit*, discriminaz*
	CON_PED	Riflessioni didattiche/metodologiche	didattic*, apprendiment*, competenz*, classe, studente, insegnant*
	CON_ETH	Questioni normative o valoriali: privacy, sicurezza, responsabilità, copyright, tutela dei minori, linee	etic*, responsabil*, privacy, regolament*, sicurezza, normativa,

		guida, policy, compliance (GDPR, AI Act).	AI Act, "GDPR", licenza, copyright
Tipologie di applicazione	APP_CREAT	Creatività e produzione (storytelling, avatar, arte)	avatar, stor*, storytelling, creaz*, narraz*, disegn*
	APP_VAL	Valutazione e verifica	valutaz*, quiz, test, verifica, rubrica, autovalut*
	APP_CODE	Coding & pensiero computazionale	coding, Scratch, pensiero computazionale, Teachable Machine, unplugged
	APP_COLL	Lavoro di gruppo / peer learning	grupp*, collabor*, cooper*, peer, brainstorm*, discuss*

Per garantire l'affidabilità dell'analisi qualitativa, è stata condotta una verifica di coerenza inter-coder attraverso un processo di doppia codifica temporalmente distanziata. Lo stesso ricercatore ha eseguito la codifica dell'intero corpus in due momenti separati a distanza di sei mesi, per poi confrontare la concordanza tra le due codifiche. Questo approccio, sebbene non elimini completamente il rischio di bias individuali come farebbe una codifica indipendente da parte di più ricercatori, offre comunque una misura significativa della stabilità e della coerenza del sistema di codifica nel tempo.

La Tabella 5.20 riporta i valori di affidabilità calcolati per ciascun codice, utilizzando diverse metriche standard nella ricerca qualitativa.

Tabella 5.20: Misure di affidabilità della codifica qualitativa

Codice	Po	Cohen's κ	κ (95% CI)	Krippendorff's α	α (95% CI)	PABAK
ATT_POS	0,940	0,859	[0,791 - 0,924]	0,859	[0,790 - 0,924]	0,880
ATT_NEG	0,932	0,848	[0,779 - 0,915]	0,848	[0,779 - 0,914]	0,863
CON_KNOW	0,932	0,859	[0,790 - 0,919]	0,859	[0,789 - 0,919]	0,863
CON_OPE	0,932	0,792	[0,686 - 0,874]	0,792	[0,685 - 0,874]	0,863

CON_CRIT	0,932	0,860	[0,795 0,920]	-	0,859	[0,793 0,920]	-	0,863
CON_PED	0,880	0,336	[0,163 0,504]	-	0,308	[0,111 0,493]	-	0,759
CON_ETH	0,876	0,745	[0,654 0,819]	-	0,744	[0,653 0,819]	-	0,751
APP_CREAT	0,723	0,198	[0,102 0,303]	-	0,117	[-0,017 0,250]	-	0,446
APP_VAL	0,888	0,773	[0,697 0,844]	-	0,770	[0,690 0,843]	-	0,775
APP_CODE	0,948	0,872	[0,795 0,933]	-	0,872	[0,795 0,933]	-	0,896
APP_COLL	0,952	0,867	[0,787 0,932]	-	0,867	[0,785 0,932]	-	0,904

Nota: Po = Percentuale di accordo osservato; κ = Cohen's kappa; α = Krippendorff's alpha; PABAK = Prevalence-Adjusted Bias-Adjusted Kappa

I valori di Cohen's κ e Krippendorff's α mostrano generalmente un accordo da sostanziale a quasi perfetto (0.745-0.872) per la maggior parte dei codici, con l'eccezione di CON_PED (accordo discreto, κ = 0.336) e APP_CREAT (accordo debole, κ = 0.198). Questi due codici con affidabilità inferiore hanno richiesto un'attenzione particolare nell'interpretazione dei risultati, considerando potenziali ambiguità nella loro applicazione. Gli intervalli di confidenza al 95% confermano la significatività statistica dei valori di κ e α per tutti i codici tranne APP_CREAT, il cui intervallo di confidenza per α include valori negativi.

L'analisi integrata di questo ricco e multiforme materiale testuale si è articolata attraverso un processo iterativo che ha combinato l'interpretazione umana con l'analisi computazionale, seguendo le fasi principali qui descritte:

Pre-elaborazione dei Dati: Il primo passo, fondamentale per qualsiasi analisi testuale computazionale, ha riguardato la preparazione del corpus. I testi grezzi, raccolti tramite *Google Forms* ed esportati in formato CSV, sono stati importati in un ambiente di analisi dati (specificamente, Google Colab con l'ausilio della libreria Pandas). È stata effettuata una pulizia preliminare per rimuovere eventuali artefatti di raccolta dati o informazioni non pertinenti. Successivamente, il testo è stato sottoposto a procedure standard di pre-processing NLP, implementate principalmente tramite la libreria NLTK (Natural Language Toolkit) per la lingua italiana: normalizzazione (conversione sistematica in minuscolo per evitare duplicazioni), rimozione della punteggiatura e dei caratteri speciali, eliminazione dei numeri (quando non semanticamente rilevanti) e degli spazi bianchi multipli. È stata quindi effettuata la tokenizzazione, ovvero la suddivisione del testo continuo in unità lessicali discrete

(token, corrispondenti approssimativamente alle parole), e la rimozione delle stopwords, cioè delle parole grammaticali molto frequenti e poco informative (articoli, preposizioni, congiunzioni), utilizzando una lista standard per l'italiano. In alcuni casi, per analisi specifiche come il calcolo delle frequenze lessicali di base, si è proceduto anche alla lemmatizzazione, ovvero alla riduzione delle parole flesse alla loro forma base (lemma), per aggregare varianti morfologiche dello stesso concetto (es. "algoritmo", "algoritmi" ricondotti al lemma "algoritmo"). Questi passaggi, pur comportando una minima perdita di informazione contestuale, sono essenziali per ridurre il "rumore" lessicale e rendere il testo trattabile dagli algoritmi NLP in modo efficiente e significativo (Manning, Raghavan & Schütze, 2008).

Una volta preparato il corpus, sono state applicate diverse tecniche computazionali, utilizzando un ecosistema di librerie Python specializzate (tra cui NLTK, Scikit-learn, Pandas per la manipolazione dati, Matplotlib e Seaborn per le visualizzazioni, NetworkX per le reti e WordCloud per le nuvole di parole), al fine di estrarre pattern quantitativi e strutturali:

- *Analisi di Frequenza Lessicale e Type-Token Ratio (TTR)*: È stata calcolata la frequenza di ciascun lemma significativo nel corpus pulito (e per ciascun diario separatamente) per identificare i termini più ricorrenti, offrendo una prima indicazione dei temi centrali. È stato inoltre calcolato il TTR (rapporto tra numero di tipi, cioè parole uniche, e numero di token, parole totali) come indice semplice della diversità lessicale, fornendo una misura della ricchezza e della ripetitività del vocabolario utilizzato in ciascuna fase (cfr. Tabella 5.21).
- *TF-IDF (Term Frequency-Inverse Document Frequency)*: Questa tecnica è stata impiegata per andare oltre la semplice frequenza e pesare l'importanza relativa dei termini (sia unigrammi che bigrammi significativi) all'interno di ciascuna sezione del diario (Credenze, Conoscenza, Applicazioni) e per ciascun diario nel suo complesso rispetto all'intero corpus. Il TF-IDF assegna un peso maggiore ai termini che sono frequenti in un documento specifico ma rari nel resto del corpus, identificando così le parole chiave più distintive e caratterizzanti per ogni fase della riflessione e per ogni dimensione dell'AI literacy (cfr. Tabelle I, II, III, IV nell'annex).
- *Analisi di Similarità (Coseno e Jaccard)*: Per quantificare l'evoluzione del discorso e valutare l'ipotesi di un apprendimento trasformativo, sono state condotte due tipologie complementari di analisi di similarità che offrono prospettive distinte ma convergenti sui processi di cambiamento attivati dal percorso formativo.
 - 1) *Analisi INTER-diario (trasformazione linguistica collettiva)*: sono state calcolate misure di similarità tra i corpus aggregati delle diverse sezioni (Credenze vs Conoscenza, Conoscenza vs Applicazioni, Credenze vs Applicazioni) per ciascuno dei quattro momenti. La similarità coseno, calcolata sui vettori TF-IDF dei testi delle sezioni aggregate, misura la somiglianza semantica e concettuale basandosi sulla co-occorrenza pesata dei termini a livello di corpus complessivo. L'indice di Jaccard (Jaccard, 1901), calcolato sugli insiemi di parole uniche (tipi) presenti in ciascuna sezione aggregata, misura invece la sovrapposizione puramente lessicale tra i vocabolari collettivi utilizzati nelle diverse fasi. Valori bassi in entrambe le metriche (vicini a 0) sono stati interpretati come indicatori di un significativo cambiamento nel linguaggio e nei temi trattati tra una fase e l'altra della riflessione collettiva, suggerendo una ristrutturazione cognitiva profonda a livello di gruppo piuttosto che un semplice apprendimento additivo (Mezirow, 1991).

- 2) Analisi INTRA-diario (coerenza individuale del ragionamento): parallelamente, sono state calcolate misure di similarità all'interno dei singoli diari individuali per valutare la coerenza logica e tematica dei ragionamenti personali attraverso le diverse sezioni. Questa analisi complementare è cruciale per distinguere tra una trasformazione autentica (che dovrebbe mantenere coerenza interna pur cambiando il linguaggio collettivo) e una possibile confusione cognitiva (che produrrebbe incoerenza sia a livello individuale che collettivo).

La combinazione di bassa similarità INTER-diario (indicativa di trasformazione linguistica collettiva) e alta coerenza INTRA-diario (indicativa di ragionamenti individuali logicamente strutturati) fornisce una prova metodologicamente robusta dell'apprendimento trasformativo, dimostrando che i cambiamenti osservati non sono dovuti a confusione ma rappresentano un'autentica ristrutturazione concettuale accompagnata da capacità di elaborazione coerente e articolata (cfr. Tabelle annex).

○ *Topic Modeling con LDA (Latent Dirichlet Allocation):*

Per identificare i temi latenti presenti nel corpus, è stata utilizzata la tecnica del Topic Modeling basata su LDA (Blei, Ng & Jordan, 2003). L'LDA è un metodo probabilistico che permette di scoprire strutture tematiche nascoste in grandi collezioni di testi, rappresentando ciascun documento come una miscela di topic latenti e ciascun topic come una distribuzione probabilistica di parole.

Per ottimizzare la selezione del numero di topic (k), sono stati testati diversi valori da k=3 a k=15, valutando la coerenza semantica di ciascun modello attraverso la metrica di coerenza 'c_v', che misura il grado di interpretabilità semantica dei topic estratti. In base all'analisi di coerenza, è stato selezionato un modello con k=5 topic come ottimale (coerenza: 0.3276), in quanto forniva il miglior equilibrio tra specificità tematica e interpretabilità.

I cinque topic identificati dall'analisi LDA, con le relative parole chiave e le loro probabilità, sono riportati nella Tabella 5.21.

Tabella 5.21: Topic identificati dall'analisi LDA sul corpus complessivo

Topic	Parole chiave (con peso probabilistico)
Topic 0	"algoritmi" (0.007), "storia" (0.007), "studenti" (0.007), "algoritmo" (0.007), "pensiero" (0.007), "istruzioni" (0.006), "esempio" (0.005), "computazionale" (0.005), "computer" (0.005), "base" (0.005)
Topic 1	"machine" (0.011), "learning" (0.010), "bias" (0.008), "esempio" (0.006), "dati" (0.005), "gruppi" (0.004), "modello" (0.004), "pensato" (0.004)
Topic 2	"dati" (0.012), "bias" (0.011), "pensiero" (0.010), "critico" (0.010), "immagini" (0.009), "pregiudizi" (0.007), "macchine" (0.007), "informazioni" (0.005), "stereotipi" (0.005)
Topic 3	"molto" (0.006), "l'intelligenza" (0.006), "potrebbe" (0.006), "cosa" (0.006), "proprio" (0.005), "macchina" (0.005), "lezioni" (0.005), "dimensione" (0.005)
Topic 4	"etica" (0.018), "dimensione" (0.010), "studenti" (0.009), "gruppo" (0.009), "etici" (0.009), "potrebbe" (0.008), "responsabilità" (0.007), "dati" (0.007), "etico" (0.007)

- *Analisi di Co-occorrenza e Reti Semantiche:* È stata calcolata la matrice di co-occorrenza

dei termini più frequenti e significativi all'interno di finestre di testo definite (es. la singola risposta) per identificare quali parole tendono ad apparire insieme nello stesso contesto. Queste relazioni di prossimità semantica sono state visualizzate tramite heatmap e, soprattutto, tramite reti semantiche (utilizzando la libreria NetworkX). In queste reti, i nodi rappresentano i termini e gli archi (spesso pesati in base alla frequenza di co-occorrenza) indicano la forza della loro associazione nel discorso dei partecipanti, permettendo di esplorare visivamente cluster tematici e connessioni concettuali.

- *Sentiment Analysis*: Per valutare l'evoluzione della tonalità emotiva, è stata condotta un'analisi del sentiment utilizzando un modello Transformer pre-addestrato per la lingua italiana (neuraly/bert-base-italian-cased-sentiment). È stato calcolato un punteggio di polarità (-1 a +1) per ciascuna delle sezioni principali dei diari. Questo approccio, simile a metodi come SO-CAL (Taboada et al., 2011), permette di stimare il tono emotivo prevalente. Per i punteggi medi di sentiment tra i quattro diari e valutare la significatività statistica delle differenze osservate, è stato utilizzato il test non parametrico di Kruskal-Wallis (Kruskal & Wallis, 1952), adatto per confrontare più gruppi indipendenti senza assumere la normalità dei dati, seguito da test post-hoc di Dunn con correzione di Bonferroni per identificare quali coppie di diari differissero significativamente.
- *Word Clouds*: Infine, sono state generate nuvole di parole per ciascuna sezione e per l'intero corpus di ogni diario, offrendo una rappresentazione visuale immediata e intuitiva dei termini più frequenti, utile come strumento esplorativo e comunicativo.

In parallelo e in dialogo costante con le analisi computazionali, è stata condotta un'analisi tematica interpretativa sull'intero corpus testuale. Seguendo le fasi canoniche proposte da Braun & Clarke (2006, 2012) – familiarizzazione approfondita con i dati attraverso letture multiple; generazione sistematica di codici descrittivi iniziali per segmenti di testo significativi; ricerca attiva di temi sovraordinati raggruppando i codici correlati; revisione critica e affinamento dei temi per assicurarne coerenza interna e distinzione esterna; definizione precisa e denominazione appropriata dei temi finali – si è cercato di cogliere la ricchezza semantica e le prospettive soggettive. Particolare attenzione è stata dedicata all'analisi interpretativa delle metafore come finestre sulle credenze implicite, all'analisi del contenuto dei commenti alle definizioni/citazioni, alla categorizzazione delle tipologie di attività didattiche proposte e all'analisi delle motivazioni sottostanti la scelta delle immagini rappresentative degli scenari d'uso.

Nella fase finale, cruciale per un approccio misto, i risultati emersi dalle diverse analisi NLP (frequenze, termini chiave TF-IDF, topic LDA, pattern di similarità, reti semantiche, punteggi di sentiment, visualizzazioni) sono stati sistematicamente confrontati, incrociati e integrati con i temi e le interpretazioni derivanti dall'analisi qualitativa. L'obiettivo non era semplicemente giustapporre i risultati, ma utilizzare le evidenze quantitative e strutturali per validare, arricchire, sfumare o approfondire le interpretazioni qualitative, e, viceversa, utilizzare la comprensione qualitativa del contesto e del significato per dare senso ai pattern computazionali emersi. Questa triangolazione metodologica (Denzin, 1978) ha permesso di costruire una narrazione dei risultati più robusta, dettagliata, sfaccettata e, in ultima analisi, attendibile.

Nelle sezioni seguenti, verranno presentate le analisi dettagliate per ciascuno dei quattro diari, strutturate secondo questo approccio integrato, evidenziando l'evoluzione dei partecipanti attraverso le dimensioni conoscitiva, operativa, critica ed etica dell'AI literacy.

5.2.1 Analisi del diario 1: dimensione conoscitiva - demistificare l'ignoto

Il primo modulo formativo, e il diario riflessivo ad esso associato, erano dedicati alla dimensione conoscitiva dell'AI literacy. L'intento era di gettare le fondamenta concettuali, accompagnando i partecipanti nella familiarizzazione con la terminologia essenziale, nella comprensione dell'evoluzione storica dell'IA – caratterizzata da cicliche "primavere" di entusiasmo e "inverni" di disillusione –, nella disamina delle principali definizioni proposte nel tempo da studiosi e istituzioni, e nell'avvio di una riflessione cruciale sulla distinzione, spesso sfumata nel discorso comune, tra intelligenza umana e intelligenza artificiale. Il diario mirava, quindi, a catturare le conoscenze pregresse, le percezioni iniziali e le prime rielaborazioni concettuali su questi aspetti fondamentali.

L'analisi di frequenza condotta sul corpus testuale del primo diario (N=63 partecipanti, 25.092 parole pulite) conferma immediatamente, come era lecito attendersi, l'assoluta centralità dei termini "intelligenza" (con 559 occorrenze) e "artificiale" (488 occorrenze). La loro elevata frequenza testimonia la pertinenza del corpus rispetto al tema indagato. Tuttavia, risulta particolarmente significativo osservare come, fin da questo primo momento di riflessione focalizzato sulla conoscenza teorica, emergano con forza termini inequivocabilmente legati al contesto umano e educativo: "essere" (252 occorrenze), "bambini" (150), "attività" (149), "immagine" (140), "studenti" (94) e "classe" (87). Questa precoce compenetrazione tra il lessico dell'IA e quello pedagogico suggerisce una tendenza spontanea e radicata nei partecipanti a calare immediatamente le riflessioni sull'IA, anche quelle più astratte, all'interno della propria cornice professionale. Accanto a questi poli tematici, sono molto presenti anche termini legati alla modalità stessa della riflessione e dell'esplorazione concettuale in corso: verbi modali come "può" (190) e "potrebbe" (109), la ricerca di esemplificazioni e un confronto costante con le conoscenze pregresse, come indicato dalla frequenza del termine "prima" (104). Il lessico complessivo suggerisce quindi un discorso di natura introduttiva ed esplorativa. La ricchezza lessicale, misurata tramite il Type-Token Ratio (TTR) è pari a 0.2139 (cfr. Tabella 5.21), indicando una moderata diversità lessicale. Ciò suggerisce che i partecipanti utilizzano un vocabolario sufficientemente vario, ma convergono su termini chiave condivisi, necessari per discutere di un argomento complesso e nuovo, risultando in un linguaggio focalizzato ma non eccessivamente ripetitivo.

Una finestra privilegiata sulle percezioni e sulle credenze iniziali è offerta dall'analisi delle metafore utilizzate. L'analisi qualitativa rivela una chiara predominanza di immagini concettuali riconducibili all'ignoto e alla vastità complessa (es. "oceano", "viaggiatore", "mappa", "giardino segreto") o alla difficoltà di comprensione (es. "labirinto", "puzzle", "nebbia", "scatola chiusa"). Accanto a questo tema dominante dell'ignoto, si manifesta con forza anche la percezione della potenza dell'IA, una potenza spesso connotata in termini ambivalenti o potenzialmente incontrollabili (es. "cavallo imbizzarrito", "onda minacciosa", "vita aliena"). Infine, un terzo gruppo di metafore, meno frequente ma comunque presente, suggerisce un atteggiamento iniziale più pragmatico, vedendo l'IA come uno strumento o una risorsa informativa (es. "biblioteca", "lanterna", "bussola"). Questa eterogeneità delle percezioni iniziali, che spaziano dal timore reverenziale dell'ignoto all'interesse pragmatico per lo strumento, riflette probabilmente la natura pervasiva ma ancora nebulosa dell'IA nel discorso

pubblico e l'assenza, per molti insegnanti, di precedenti esperienze formative strutturate e approfondite sull'argomento.

Il passaggio alla sezione del diario dedicata alla "Conoscenza Attuale" rivela un netto cambiamento nel linguaggio e nel focus tematico, indicando un'effettiva acquisizione e rielaborazione dei contenuti specifici del modulo. L'analisi TF-IDF, che identifica i termini più distintivi di questa sezione rispetto al resto del corpus, mostra che il discorso si concentra ora su concetti legati alla definizione formale e al dibattito concettuale sull'IA. Termini come intelligenza (TF-IDF score aggregato: 6.36), umano (5.11), artificiale (4.83), intelligente (4.62), macchina (4.37), definizione (4.21), e grado (4.07) diventano particolarmente rilevanti. Questo lessico suggerisce che i partecipanti hanno attivamente interiorizzato le definizioni fornite durante il modulo (spesso richiamando esplicitamente fonti specifiche come McCarthy, le linee guida dell'Unione Europea, o citazioni chiave discusse), hanno riflettuto sul fondamentale confronto tra intelligenza umana e artificiale (il riferimento ad Alan Turing e al suo celebre test è frequente e significativo), e hanno colto il concetto cruciale che l'IA contemporanea opera principalmente attraverso la simulazione o l'imitazione di specifiche capacità cognitive umane, piuttosto che attraverso un pensiero autonomo, cosciente ed emotivamente connotato analogo al nostro. L'analisi lessicale comparativa, identificando i termini utilizzati quasi esclusivamente in questa sezione, rafforza questa interpretazione: parole come "citazione", "obiettivi" (specifici dell'IA), "interpretare", "test", "estensione" (delle capacità umane), "simulazione", "filosofica", "mccarthy", "leacun" (figure chiave nella storia dell'IA) appartengono chiaramente al dominio concettuale, storico e filosofico dell'IA, sostituendo il lessico più vago, metaforico e personale delle credenze iniziali. Un altro aspetto cruciale che emerge con forza dai commenti alle definizioni scelte e dalle riflessioni libere è una rinnovata e più articolata consapevolezza dei limiti intrinseci dell'IA. Molti partecipanti, dopo averne esplorato le definizioni e le capacità, sottolineano esplicitamente l'assenza di emozioni autentiche, empatia, coscienza soggettiva, creatività genuina e capacità di giudizio etico nelle macchine attuali, operando una distinzione netta tra le loro pur notevoli prestazioni in compiti specifici (spesso superiori a quelle umane) e la complessità irriducibile dell'esperienza e dell'intelligenza umana nella sua interezza.

La sezione dedicata alle "Applicazioni Future" dimostra la notevole capacità dei partecipanti di iniziare immediatamente a tradurre le conoscenze concettuali appena acquisite in ipotesi concrete di intervento didattico. Il lessico identificato dal TF-IDF è dominato da termini pedagogici (bambino (5.65), attività (5.23), affiancati da intelligenza (4.03) e artificiale (3.90), e ulteriormente specificato da classe (3.49), scuola (3.01), utilizzare (2.89) e alunno), segnalando una forte e immediata proiezione nel contesto professionale. L'analisi dei termini utilizzati esclusivamente in questa sezione è particolarmente eloquente nel rivelare le due principali direttrici delle strategie didattiche immaginate in questa fase iniziale. Da un lato, emerge con prepotenza la metodologia della Philosophy for Children (P4C), identificata da termini esclusivi come "insegnante", "proporre", "gruppi", "discussione", "philosophy", "children". Questa metodologia viene vista come strumento elettivo per introdurre il tema complesso e potenzialmente divisivo dell'IA in modo riflessivo, per stimolare il pensiero critico sulle diverse definizioni e sugli impatti sociali, e per facilitare discussioni aperte e argomentate in classe. Dall'altro lato, si manifesta un forte interesse per attività mirate a demistificare il funzionamento di base dell'IA e a sviluppare le fondamenta del pensiero computazionale, anche nei più piccoli. Termini esclusivi come "introdurre", "percorsi", "piccoli", "gioco", "coding", "robot" puntano in questa direzione. Il coding, spesso menzionato in riferimento a

strumenti specifici come il robot programmabile Bee-Bot o ad attività *unplugged* (senza computer), e l'uso di giochi educativi che simulano in modo semplificato l'apprendimento automatico (come *AI for Oceans* o il gioco dei fiammiferi, entrambi sperimentati durante il corso), vengono visti come strategie chiave per rendere tangibili concetti altrimenti astratti. Accanto a questi due filoni principali, si delinea anche l'idea, seppur meno frequente, di utilizzare l'IA come strumento di supporto diretto per l'insegnante (ad esempio, per la creazione di materiali didattici) o per gli studenti (per la personalizzazione dei percorsi o come supporto ai Bisogni Educativi Speciali). Questa combinazione di approcci suggerisce una visione dell'AI literacy che, fin dalla sua dimensione conoscitiva, integra strettamente la riflessione critica sui concetti con la comprensione preliminare dei meccanismi operativi.

Infine, la distribuzione delle immagini scelte nel primo diario (N=61 partecipanti) per rappresentare lo scenario d'uso didattico preferito offre uno spaccato interessante delle rappresentazioni mentali associate all'insegnamento e all'apprendimento dell'IA dopo averne affrontato la dimensione conoscitiva. Non emerge una preferenza netta, ma piuttosto un mosaico di approcci. La preferenza relativa va all'Immagine 3 (19 scelte), raffigurante uno studente immerso nello studio individuale tra libri e idee (simboleggiate da lampadine), suggerendo un'associazione iniziale dell'AI literacy con la comprensione concettuale profonda e la riflessione personale. Tuttavia, questa visione è immediatamente affiancata dall'Immagine 4 (14 scelte), che rappresenta un ambiente di laboratorio vivace e collaborativo, e dall'Immagine 2 (12 scelte), che mostra uno studente che utilizza uno schermo interattivo come supporto allo studio individuale. Seguono, con frequenze minori ma significative, l'Immagine 6 (7 scelte), interpretabile come un contesto di discussione di gruppo o presentazione, l'Immagine 5 (6 scelte), legata all'analisi critica di artefatti visivi o narrativi, e l'Immagine 1 (4 scelte), che evoca un contesto più formale, forse scientifico-disciplinare. Questa distribuzione variegata indica che i partecipanti, fin da questa fase iniziale, non convergono su un unico modello ideale di insegnamento/apprendimento per l'IA. Piuttosto, emerge una visione plurale e integrata, che riconosce l'importanza complementare sia della comprensione individuale e dello studio approfondito (Img 3), sia dell'attività pratica, della sperimentazione e della collaborazione tra pari (Img 4), sia dell'uso dell'IA come strumento di supporto personalizzato (Img 2), sia della discussione sociale e del confronto di idee (Img 6), sia dell'analisi critica di prodotti specifici (Img 5). Questa visione iniziale, già orientata verso un mosaico di approcci metodologici, pone le basi per l'adattamento flessibile che si osserverà nei diari successivi in risposta alle specificità delle altre dimensioni dell'AI literacy.

L'analisi della similarità tra le diverse sezioni del diario fornisce una delle evidenze quantitative più nette e significative dell'impatto del percorso formativo fin dal suo primo modulo. Analisi INTER-diario (trasformazione collettiva del linguaggio): le medie di similarità coseno calcolate tra i vettori TF-IDF delle sezioni risultano estremamente basse: il valore medio è 0.079 tra la sezione "Credenze Iniziali" e la sezione "Conoscenza Attuale", 0.067 tra "Conoscenza Attuale" e "Applicazioni Future", e addirittura 0.056 tra "Credenze Iniziali" e "Applicazioni Future". Questi valori, molto vicini allo zero (che indicherebbe una completa dissimilarità), segnalano una profonda divergenza nel linguaggio, nei temi e nei concetti espressi nelle diverse fasi della riflessione stimolata dal diario.

Analisi INTRA-diario (coerenza individuale del ragionamento): parallelamente, l'analisi della similarità all'interno dei singoli diari individuali rivela valori significativamente più elevati (similarità coseno media ~ 0.80), indicando che ciascun partecipante mantiene una buona coerenza logica e

tematica tra le proprie riflessioni nelle diverse sezioni. Questo dato è cruciale perché dimostra che la trasformazione linguistica collettiva non è dovuta a confusione o incoerenza individuale, ma rappresenta un autentico cambiamento strutturale nel modo di concettualizzare e discutere l'IA.

La combinazione di questi due pattern – bassa similarità INTER-diario e alta coerenza INTRA-diario – fornisce una prova robusta dell'apprendimento trasformativo. Suggerisce fortemente che i partecipanti non hanno semplicemente aggiunto nuove informazioni alle loro strutture cognitive preesistenti in un processo di apprendimento additivo, ma hanno vissuto un processo di ristrutturazione significativa delle loro concezioni iniziali, mantenendo al contempo la capacità di elaborare ragionamenti coerenti e articolati. Questo fenomeno è interpretabile, alla luce delle teorie sull'apprendimento adulto, come un apprendimento trasformativo (Mezirow, 1991), in cui le prospettive di significato iniziali vengono messe criticamente in discussione e rielaborate alla luce delle nuove conoscenze, producendo un cambiamento profondo nel linguaggio collettivo utilizzato per discutere l'IA, pur preservando la logica interna del ragionamento individuale.

La marcata discontinuità lessicale e semantica tra le fasi del percorso, quantificata da questi bassi indici di similarità INTER-diario, combinata con l'alta coerenza INTRA-diario, costituisce un indicatore particolarmente potente e metodologicamente robusto della trasformazione concettuale avvenuta già dalla prima fase del percorso.

Infine, anche l'analisi esplorativa tramite Topic Modeling (LDA), pur richiedendo cautela interpretativa data la natura non supervisionata dell'algoritmo, ha identificato topic latenti coerenti con gli obiettivi specifici della dimensione conoscitiva. I gruppi di parole più probabili emersi nei diversi topic (ad esempio, un topic con parole come *ia, umano, macchine, cosa, attività*; un altro con *bambini, attività, studenti, classe, pensiero*; un altro ancora con *sistema, strumento, bambini, grado*) riflettono le principali aree tematiche affrontate nel modulo: la natura dell'IA e il confronto tra intelligenza umana e artificiale, le implicazioni per il contesto educativo e la concettualizzazione dell'IA come sistema complesso o come strumento potenziale. La coerenza tra i topic estratti automaticamente e i contenuti didattici del modulo suggerisce che il modello computazionale è riuscito a catturare efficacemente le strutture semantiche sottostanti al discorso dei partecipanti, confermando la focalizzazione delle loro riflessioni sui nuclei concettuali della dimensione conoscitiva.

In sintesi, l'analisi integrata del primo diario disegna un quadro chiaro e incoraggiante dell'impatto iniziale del percorso formativo. I partecipanti, partendo da una comprensione dell'IA spesso vaga, metaforicamente ricca ma talvolta carica di timore, hanno sviluppato una base concettuale più solida, articolata e critica. Hanno familiarizzato con le definizioni fondamentali, l'evoluzione storica e il dibattito cruciale sul rapporto tra intelligenza umana e artificiale. Hanno acquisito un lessico più specifico e, soprattutto, una maggiore consapevolezza dei limiti attuali dell'IA, superando visioni ingenuie o eccessivamente antropomorfizzate. Questo processo di apprendimento si è dimostrato trasformativo, inducendo una significativa ristrutturazione del linguaggio e delle concezioni, come evidenziato dalle misure di similarità. È notevole osservare come, fin da subito, questa nuova comprensione sia stata proiettata sul piano didattico, con una chiara preferenza per approcci che integrano la riflessione critica (facilitata da metodologie come la P4C) con la demistificazione del funzionamento dell'IA (attraverso coding e giochi), e una visione degli scenari d'uso che bilancia apprendimento individuale, collaborazione tra pari e supporto strumentale. Il primo modulo ha quindi

efficacemente gettato basi solide e stimolato una riflessione profonda, preparando il terreno per affrontare le dimensioni successive dell'AI literacy con maggiore consapevolezza e strumenti concettuali più affinati.

5.2.2 Analisi del diario 2: dimensione operativa - svelare i meccanismi

Il secondo modulo formativo ha rappresentato un passaggio cruciale, spostando l'attenzione dalla comprensione concettuale del "cosa" l'IA sia, all'esplorazione del "come" essa funzioni concretamente. La dimensione operativa dell'AI literacy è stata quindi al centro della riflessione, focalizzandosi sui mattoni fondamentali che ne costituiscono l'architettura logica e funzionale: gli algoritmi, intesi come sequenze finite e non ambigue di istruzioni per risolvere un problema o eseguire un compito; il ruolo determinante dei dati, non più visti come semplice informazione ma come materia prima essenziale per l'addestramento dei sistemi di IA, specialmente nell'apprendimento automatico; i principi base dell'apprendimento automatico (machine learning), ovvero la capacità delle macchine di migliorare le proprie prestazioni su un compito specifico attraverso l'esperienza (i dati), senza essere state esplicitamente programmate per ogni singola eventualità; e infine, il pensiero computazionale, inteso non solo come abilità di programmazione, ma come approccio mentale alla scomposizione dei problemi (decomposition), al riconoscimento di pattern (pattern recognition), alla focalizzazione sugli aspetti rilevanti (abstraction) e alla progettazione di soluzioni algoritmiche (algorithm design). Il diario riflessivo associato mirava quindi a valutare la comprensione di questi concetti tecnici fondamentali e la capacità dei partecipanti di tradurli in attività didattiche appropriate e significative per i propri studenti.

L'analisi del lessico utilizzato nel secondo diario (N=60 partecipanti, 24.277 parole totali dopo pulizia) mostra una trasformazione evidente e significativa rispetto al primo. Sebbene i termini generali "intelligenza" (219 occorrenze) e "artificiale" (211) rimangano importanti come sfondo tematico, il palcoscenico è ora chiaramente dominato da un vocabolario tecnico specifico della dimensione operativa. I termini "algoritmo" (162) e la sua forma plurale "algoritmi" (132) emergono come le parole chiave più distintive e frequenti, segnalando l'assoluta centralità di questo concetto. Essi sono affiancati da altri termini tecnici fondamentali come "dati" (127), "istruzioni" (114), "macchine" (114) e "macchina" (110), nonché dal termine "computazionale" (110), quasi sempre associato a "pensiero". La predominanza di questo lessico indica in modo inequivocabile che la riflessione dei partecipanti si è concentrata sui processi logico-sequenziali che governano l'IA (algoritmi, istruzioni), sul materiale grezzo che alimenta l'apprendimento automatico (dati) e sull'approccio metodologico del pensiero computazionale. Il contesto educativo rimane una cornice importante, come testimoniato dalla persistente alta frequenza di "attività" (251) e "bambini" (202), e dalla rilevanza di "pensiero" (127), ma il discorso acquisisce una marcata e specifica connotazione tecnica. Il Type-Token Ratio (TTR), pari a 0.2033, risulta leggermente inferiore a quello del primo diario (0.2139). Questa lieve riduzione della diversità lessicale non va interpretata come un impoverimento, ma piuttosto come un segnale di maggiore focalizzazione e convergenza terminologica. È plausibile che la necessità di discutere con precisione concetti tecnici specifici come algoritmi, dati o machine learning abbia portato i partecipanti a utilizzare un linguaggio più condiviso e meno eterogeneo rispetto alla discussione più ampia e aperta della dimensione conoscitiva.

Le metafore utilizzate dai partecipanti per descrivere le loro credenze iniziali sulla dimensione operativa presentano anch'esse un carattere distintivo rispetto al primo diario. Le immagini più

frequenti sono "labirinto" (4 occorrenze), "mappa" (3), "puzzle" (2) e "bussola" (2). A differenza delle metafore del primo diario, spesso legate a una vastità imperscrutabile o a una potenza quasi magica e incontrollata, queste immagini suggeriscono una percezione degli algoritmi e del funzionamento interno dell'IA come una sfida complessa, sì, ma intrinsecamente strutturata, logica e potenzialmente risolvibile. Il "labirinto" e il "puzzle" evocano l'idea di un problema con regole definite, che richiede ingegno, strategia e passaggi sequenziali per essere risolto. La "mappa" e la "bussola", d'altro canto, suggeriscono la necessità di una guida, di istruzioni precise o di un metodo per navigare questa complessità procedurale. Sembra quindi prevalere un'idea di difficoltà legata alla logica formale, alla sequenzialità delle operazioni e alla necessità di comprendere le regole del gioco, piuttosto che al mistero o alla vaghezza concettuale. Questo potrebbe indicare un approccio iniziale alla dimensione operativa più analitico, forse meno carico emotivamente e più orientato al problem-solving rispetto alla dimensione conoscitiva generale.

L'analisi della sezione "Conoscenza Attuale", tramite la tecnica TF-IDF che ne evidenzia i termini più caratteristici, conferma l'acquisizione di una comprensione tecnica più approfondita e specifica. I termini che emergono come più rilevanti sono inequivocabilmente "dati", "algoritmi", "macchina", "machine learning", ma anche "simbolica", "apprende", "automatico". Questo lessico specialistico indica che i partecipanti hanno non solo consolidato la definizione di algoritmo come sequenza di istruzioni, ma hanno anche iniziato a distinguere tra diversi approcci storici e concettuali all'interno dell'IA (come l'approccio basato su regole esplicite, definito "simbolico", e quello più recente basato sull'apprendimento dai dati, o "connessionista", di cui il "machine learning" è l'esponente principale). Hanno inoltre afferrato il ruolo cruciale e attivo dei "dati" come "carburante" indispensabile per l'addestramento dei modelli di apprendimento "automatico" e hanno iniziato a comprendere i meccanismi fondamentali attraverso cui le macchine possono "imparare" a migliorare le proprie "prestazioni" attraverso l'esperienza fornita dai dati. L'analisi lessicale comparativa supporta ulteriormente questa interpretazione: termini utilizzati quasi esclusivamente in questa sezione come "citazione" (spesso riferita a definizioni tecniche di algoritmo o machine learning), "basa" (come nella frequente espressione "si basa sui dati"), "prestazioni" (dei modelli), "soluzione" (del problema tramite algoritmo), "quantità" (di dati necessaria), "umana" (in confronto alle capacità della macchina), "finito" (come attributo essenziale di una procedura algoritmica) dimostrano l'appropriazione di un vocabolario tecnico e concettuale specifico, che sostituisce le percezioni metaforiche più generali delle credenze iniziali.

La sezione dedicata alle "Applicazioni Future" mostra una chiara e diffusa volontà di tradurre i concetti tecnici della dimensione operativa in esperienze didattiche concrete, accessibili e coinvolgenti per gli studenti, anche per quelli dei livelli scolastici inferiori. Il lessico identificato dal TF-IDF, pur mantenendo la centralità di termini pedagogici generali come "attività", "bambini", "classe", evidenzia la particolare rilevanza attribuita al "coding". L'analisi dei termini esclusivi di questa sezione è particolarmente eloquente nel rivelare le strategie didattiche immaginate: "pixel", "art" (indicativi della frequente proposta di attività di pixel art, un modo semplice e visivamente gratificante per introdurre concetti di codifica binaria, coordinate e rappresentazione digitale delle immagini), "gruppi", "alunni", "dovranno" (suggerendo la progettazione di attività strutturate con istruzioni precise da seguire), "introdurre", "piccoli", "infanzia", "gioco". Emerge prepotentemente l'intenzione di rendere comprensibili concetti astratti come algoritmo, sequenza, istruzione condizionale e pensiero computazionale attraverso attività pratiche, manipolative, ludiche e visive,

ritenute adatte anche per i bambini della scuola dell'infanzia e primaria. Il coding, sia nella sua forma *plugged* (utilizzando robot programmabili come il popolare Bee-Bot, citato da molti) sia, e forse ancor più frequentemente, nella sua forma *unplugged* (attraverso giochi da tavolo con istruzioni, creazione di percorsi su griglie, attività motorie basate su sequenze), viene visto come la strategia didattica d'elezione per affrontare la dimensione operativa in modo efficace e motivante.

Coerentemente con questa enfasi sulla pratica, la distribuzione delle immagini scelte nel secondo diario per rappresentare lo scenario d'uso didattico conferma e accentua lo spostamento verso l'azione e la collaborazione. L'Immagine 4, che raffigura un ambiente di laboratorio collaborativo con studenti che lavorano insieme attorno a dispositivi o progetti, diventa la più scelta in assoluto (20 partecipanti su 60). Questo suggerisce che l'apprendimento dei concetti operativi dell'IA (algoritmi, coding, pensiero computazionale) è fortemente associato dai partecipanti a scenari di lavoro attivo, sperimentazione pratica e problem-solving di gruppo. Anche l'Immagine 2, che rappresenta uno studente che utilizza uno schermo interattivo come supporto allo studio individuale, guadagna terreno rispetto al primo diario (13 scelte), indicando l'importanza attribuita all'uso di strumenti tecnologici concreti (come computer, tablet, robot educativi) per supportare queste attività, anche a livello individuale. Al contrario, l'Immagine 3, che evoca la riflessione individuale e lo studio concettuale (9 scelte), perde nettamente la sua posizione dominante rispetto al primo diario. Ciò suggerisce che la pura riflessione astratta è percepita come meno centrale e meno efficace per la comprensione della dimensione operativa rispetto all'azione pratica, all'interazione con gli strumenti e alla collaborazione con i pari.

L'analisi della similarità tra le diverse sezioni del secondo diario conferma e approfondisce i pattern di trasformazione osservati nella dimensione conoscitiva. Nell'analisi INTER-diario (trasformazione collettiva) le medie di similarità coseno tra i corpus aggregati delle sezioni rimangono estremamente basse (Credenze-Conoscenza: 0.090; Conoscenza-Applicazioni: 0.066; Credenze-Applicazioni: 0.051), fornendo una conferma robusta e quantitativa del processo di apprendimento trasformativo anche per la dimensione operativa. Mentre nell'analisi INTRA-diario (coerenza individuale) la coerenza interna dei singoli diari mantiene valori elevati (~ 0.76), indicando che nonostante la natura più tecnica e forse più strutturata dell'argomento rispetto alla dimensione conoscitiva, i partecipanti hanno comunque vissuto una significativa ristrutturazione del loro linguaggio e del loro focus concettuale mantenendo ragionamenti logicamente coerenti.

La marcata discontinuità INTER-diario tra le credenze iniziali (spesso legate a una visione della complessità come "labirinto" o "puzzle"), la conoscenza acquisita (focalizzata su dati, algoritmi, machine learning) e le applicazioni proposte (centrate sul coding e attività pratiche) testimonia una profonda rielaborazione dei concetti, mentre l'alta coerenza INTRA-diario conferma che questo processo non è stato casuale ma rappresenta una trasformazione strutturata e consapevole. Infine, i topic identificati dall'analisi LDA per questo diario convergono in modo molto chiaro sui concetti fondamentali della dimensione operativa. Le parole chiave dominanti nei diversi topic estratti sono "algoritmo", "bambini", "intelligenza", "artificiale", "dati", "istruzioni", "macchine", "pensiero", "computazionale", "coding", "learning". Questa forte coerenza tematica, catturata dal modello LDA, conferma che le riflessioni dei partecipanti si sono effettivamente concentrate sui meccanismi di funzionamento dell'IA (algoritmi, dati, istruzioni, machine learning) e sulla loro stretta connessione con l'educazione al pensiero computazionale e al coding, visti come strumenti essenziali per rendere questi meccanismi comprensibili e accessibili nel contesto educativo.

In conclusione, l'analisi integrata del secondo diario evidenzia un consolidamento e un approfondimento tecnico significativo della comprensione dell'IA da parte dei partecipanti. Superata la fase di familiarizzazione generale, il focus si è spostato con successo sui meccanismi operativi fondamentali: algoritmi, dati, apprendimento automatico e pensiero computazionale. I partecipanti hanno acquisito un lessico tecnico specifico e hanno sviluppato una visione del funzionamento dell'IA come complesso ma logicamente strutturato e analizzabile. L'apprendimento si conferma trasformativo, indicando una rielaborazione profonda. La scoperta forse più significativa di questa fase è la chiara e quasi unanime traduzione di questi concetti tecnici in proposte didattiche eminentemente pratiche, attive, spesso ludiche e visive (coding, pixel art, giochi unplugged, robotica), ritenute adatte ed efficaci anche per i primi livelli scolastici. Coerentemente, la visione prevalente dello scenario d'uso didattico si orienta decisamente verso ambienti collaborativi, laboratoriali e supportati da strumenti concreti, sottolineando l'importanza cruciale del "fare" per comprendere la dimensione operativa dell'IA.

5.2.3 Analisi del diario 3: dimensione critica - riconoscere bias e coltivare giudizio

Il terzo modulo formativo ha segnato un ulteriore, importante passaggio di livello, introducendo i partecipanti alla dimensione critica dell'AI literacy. Questo aspetto è fondamentale per sviluppare una cittadinanza digitale consapevole e responsabile nell'era algoritmica. L'attenzione si è spostata dalla comprensione del funzionamento dell'IA (dimensione operativa) alla valutazione approfondita dei suoi impatti, dei suoi limiti intrinseci e dei rischi potenziali che il suo utilizzo comporta. Sono stati affrontati temi cruciali e spesso controversi, come i bias algoritmici, ovvero le distorsioni sistematiche nei risultati prodotti dall'IA che possono derivare da dati di addestramento non rappresentativi, incompleti o essi stessi distorti, oppure da scelte di progettazione che riflettono pregiudizi consci o inconsci degli sviluppatori. Strettamente connesse ai bias sono le questioni di equità e discriminazione, poiché gli algoritmi possono perpetuare o addirittura amplificare disuguaglianze sociali esistenti. Sono stati esplorati anche i problemi legati alla trasparenza e all'interpretabilità di molti modelli di IA complessi (le cosiddette "scatole nere" o *black boxes*), che rendono difficile comprendere come vengano prese certe decisioni e verificarne la correttezza o l'equità. È stata inoltre discussa l'affidabilità delle informazioni generate dall'IA (si pensi alle *hallucinations* dei modelli linguistici o alle immagini *deepfake*) e, di conseguenza, la necessità impellente di sviluppare un pensiero critico specifico per interagire con questi sistemi, valutarne gli output in modo informato e riconoscere potenziali manipolazioni o imprecisioni. Il diario riflessivo associato a questo modulo mirava quindi a sondare la consapevolezza acquisita su queste problematiche complesse e a esplorare le strategie didattiche immaginate dai partecipanti per promuovere un approccio critico all'IA nei propri studenti.

L'analisi di frequenza del terzo diario (N=62 partecipanti, 27.026 parole totali dopo pulizia) mostra una nuova, decisa caratterizzazione lessicale che riflette pienamente il focus sulla criticità. Accanto ai termini onnipresenti "intelligenza" (335) e "artificiale" (333), il vocabolario è ora dominato da parole chiave inequivocabilmente legate alla valutazione e all'analisi critica. Il termine tecnico "bias" (241 occorrenze) emerge come concetto centrale e altamente distintivo di questa fase. Esso è frequentemente associato a "dati" (228), sottolineando la comprensione del legame tra la qualità dei dati di addestramento e l'insorgere dei bias. Altrettanto centrali sono i termini "pensiero" (220) e "critico" (197), che appaiono spessissimo in coppia, indicando l'importanza attribuita allo sviluppo di questa competenza. Termini come "pregiudizi" (123) e "critica" (116) rafforzano ulteriormente questo

orientamento. È interessante notare anche l'altissima frequenza di "immagine" (189) e "immagini" (176), che suggerisce come l'analisi critica si sia spesso focalizzata sugli output visivi dell'IA (ad esempio, la valutazione di immagini generate artificialmente per rilevare stereotipi o incongruenze) o che le immagini stesse siano state utilizzate come potenti stimoli didattici per la riflessione critica. Il lessico complessivo delinea quindi un discorso fortemente orientato all'analisi valutativa dei sistemi IA, all'identificazione delle loro potenziali distorsioni e alla necessità fondamentale di coltivare un approccio riflessivo e critico nel loro utilizzo. Il Type-Token Ratio (TTR), che si attesta a 0.1968, risulta il più basso registrato finora. Questa ulteriore riduzione della diversità lessicale suggerisce un linguaggio ancora più specializzato e convergente rispetto ai diari precedenti. Tale maggiore focalizzazione potrebbe riflettere sia la complessità intrinseca dei concetti trattati (bias, pensiero critico, equità), che richiede un uso preciso e condiviso della terminologia, sia una possibile maggiore omogeneità nelle riflessioni e nelle preoccupazioni dei partecipanti riguardo a questi temi specifici e socialmente rilevanti.

Le metafore utilizzate per esprimere le credenze iniziali sulla dimensione critica presentano sfumature peculiari e indicative. Ricompaiono alcune immagini legate all'oscurità o alla difficoltà di visione, come "notte" (3 occorrenze) o "cielo" (2), forse evocando l'assenza di stelle o punti di riferimento chiari, suggerendo una percezione della criticità come un'area difficile da illuminare, dove i problemi reali (bias, implicazioni negative, manipolazioni) non sono immediatamente evidenti o facilmente riconoscibili. La necessità di uno strumento per "vedere" più chiaramente o per orientarsi in questa complessità è richiamata da metafore come "lanterna" (1) e "bussola" (1). Il "giardino" (3), già incontrato nei diari precedenti ma qui riletto in chiave critica, potrebbe simboleggiare un sistema complesso (l'IA e i suoi prodotti) che richiede una cura critica costante per evitare la crescita incontrollata di "erbacce" indesiderate, come i bias, le discriminazioni o la disinformazione. Metafore potenti come "iceberg" (menzionata esplicitamente nei testi anche se non tra le parole chiave più frequenti nelle liste automatiche) o "scatola chiusa" (1) rafforzano potentemente l'idea di problemi latenti, nascosti sotto la superficie apparentemente liscia degli output dell'IA, che richiedono uno sguardo attento, analitico e indagatore per essere svelati e compresi nella loro reale portata. Nel complesso, la percezione iniziale della dimensione critica sembra essere quella di un dominio che richiede uno sforzo attivo di indagine, analisi e interpretazione, un esercizio di scetticismo metodologico per portare alla luce ciò che non è immediatamente visibile o dichiarato.

L'analisi della sezione "Conoscenza Attuale", attraverso la tecnica TF-IDF, conferma l'acquisizione di una comprensione specifica e approfondita dei temi critici affrontati nel modulo. I termini che emergono come più rilevanti e distintivi sono inequivocabilmente "bias", "dati", "pensiero", "critico", "pregiudizi". Emerge anche con forza il ruolo dell'"uomo" (o della società) come fonte, spesso involontaria, dei bias che vengono introdotti nei sistemi, e la rilevanza del "modello" algoritmico nel recepire, perpetuare e talvolta amplificare tali distorsioni. Questo indica che i partecipanti hanno compreso il meccanismo fondamentale attraverso cui i bias presenti nei dati di addestramento o incorporati nelle scelte di progettazione possono generare risultati iniqui, distorti o discriminatori. Hanno inoltre afferrato l'importanza cruciale di sviluppare e applicare il pensiero critico – vengono fatti riferimenti espliciti a teorici come Robert Ennis e alla sua definizione di pensiero critico come "pensiero razionale e riflessivo focalizzato sul decidere cosa credere o fare" – come strumento fondamentale per valutare l'affidabilità, l'equità, la trasparenza e le implicazioni degli output dei sistemi IA. L'analisi lessicale comparativa evidenzia l'ingresso nel vocabolario attivo dei partecipanti

di termini esclusivi di questa sezione come "modello" (algoritmico), "ennis", "forniti" (spesso riferito ai dati), "riflessivo", "razionale", "focalizzato" (attributi del pensiero critico), "garantire" (la correttezza, l'equità), "gruppo" (sociale, come oggetto potenziale di bias), "ricerca" (necessaria per identificare i bias). Questo dimostra l'appropriazione di un vocabolario specifico e articolato per l'analisi critica dell'IA.

La sezione "Applicazioni Future" mostra come i partecipanti traducano immediatamente la comprensione della dimensione critica in strategie didattiche mirate a sviluppare la consapevolezza e le capacità critiche negli studenti. Il lessico identificato dal TF-IDF mantiene la centralità di "pensiero", "critico", "bias", "immagini", integrandoli strettamente con termini pedagogici ("attività", "bambini", "studenti", "insegnante", "classe"). I termini esclusivi di questa sezione ("attività", "classe", "insegnante", "proporre", "lavoro", "dovranno", "gioco", "piccoli", "art", "disegno") suggeriscono un forte focus su attività di analisi attiva, discussione guidata e produzione consapevole. Le proposte didattiche descritte qualitativamente nei diari mirano frequentemente a:

- Rendere visibili i bias: Attraverso l'analisi guidata di esempi concreti, come la valutazione critica di immagini generate dall'IA che mostrano evidenti stereotipi di genere o etnici, o la discussione su come dataset sbilanciati possano portare a risultati distorti in diversi ambiti (dal riconoscimento facciale alla selezione del personale).
- Sviluppare il pensiero critico: Utilizzando metodologie che promuovono il ragionamento, l'argomentazione e la valutazione delle fonti, come la già citata P4C, discussioni strutturate su scenari problematici, o giochi educativi che richiedono valutazione e giudizio informato.
- Analizzare criticamente gli output dell'IA: In particolare le immagini generate artificialmente (valutandone coerenza, plausibilità, eventuali bias incorporati), ma anche testi prodotti da modelli linguistici (verificandone l'accuratezza, le fonti, la presenza di "hallucinations") o altri prodotti algoritmici, per identificarne limiti, potenziali manipolazioni o rappresentazioni stereotipate.
- Promuovere un uso consapevole e responsabile: Educare gli studenti a non accettare passivamente i risultati forniti dai sistemi IA come verità oggettive, ma a interrogarli costantemente, a verificarne le fonti (quando possibile), a confrontarli con altre informazioni e a valutarli criticamente alla luce del contesto e delle possibili implicazioni.

La distribuzione delle immagini scelte per rappresentare lo scenario d'uso didattico preferito per la dimensione critica conferma e accentua ulteriormente la tendenza verso l'interazione sociale e la costruzione collettiva di significato. L'Immagine 6, che raffigura un gruppo di persone (studenti e/o insegnante) impegnate in una discussione, una presentazione o un lavoro collettivo attorno a uno schermo o una lavagna, diventa nettamente la più popolare (20 partecipanti su 62), distanziando significativamente tutte le altre opzioni. Seguono, con frequenze notevolmente inferiori, l'Immagine 3 (riflessione individuale, 13 scelte), che mantiene una sua rilevanza probabilmente per la componente introspettiva e personale intrinseca al pensiero critico, e l'Immagine 5 (analisi di foto/storia, 11 scelte), scelta forse per la sua diretta attinenza all'analisi critica di artefatti specifici come le immagini o i casi di studio. Le immagini che evocano un contesto primariamente laboratoriale-pratico (Img 4, 8 scelte) o di supporto strumentale individuale (Img 2, solo 2 scelte) perdono decisamente terreno rispetto ai diari precedenti. Questa marcata preferenza per l'Immagine 6 suggerisce in modo potente che la dimensione critica dell'IA viene percepita dai partecipanti come un dominio che si affronta e si sviluppa primariamente attraverso il dialogo, il confronto di

prospettive, la discussione collettiva e la costruzione condivisa di consapevolezza all'interno di una comunità di apprendimento. L'analisi critica non è vista (solo) come un esercizio intellettuale individuale, ma come una pratica sociale fondamentale per navigare collettivamente le complessità e le ambiguità dell'era algoritmica.

L'analisi della similarità per il terzo diario rivela pattern di trasformazione particolarmente interessanti per la complessa dimensione critica. Nell'analisi INTER-diario (trasformazione collettiva) le medie di similarità coseno rimangono molto basse (Credenze-Conoscenza: 0.089; Conoscenza-Applicazioni: 0.095; Credenze-Applicazioni: 0.067), ribadendo la natura trasformativa dell'apprendimento anche per la complessa dimensione critica. Mentre nell'analisi INTRA-diario (coerenza individuale) la coerenza interna dei singoli diari si mantiene su valori elevati (~0.82), indicando che i partecipanti hanno significativamente ristrutturato il loro linguaggio e le loro concezioni riguardo ai bias, alla valutazione dell'informazione algoritmica e alla necessità del pensiero critico, pur mantenendo ragionamenti logicamente articolati.

Dato particolarmente interessante è il valore della similarità INTER-diario tra "Conoscenza Attuale" e "Applicazioni Future" (0.095) che, pur rimanendo basso in assoluto, è il più alto registrato tra tutte le comparazioni Conoscenza-Applicazioni nei primi tre diari. Combinato con l'alta coerenza INTRA-diario, questo potrebbe suggerire una percezione di maggiore urgenza e legame diretto tra la comprensione dei concetti critici (la natura pervasiva dei bias, l'importanza cruciale del pensiero critico) e la necessità impellente di specifiche azioni didattiche volte a promuovere tali competenze negli studenti. Sembra che, una volta comprese chiaramente le problematiche e i rischi legati ai bias e alla valutazione critica dell'IA, i partecipanti vedano con maggiore chiarezza e urgenza come e perché sia necessario affrontare questi temi in classe, traducendo la consapevolezza in proposte operative mirate.

Infine, i topic identificati dall'analisi LDA per questo diario convergono sui concetti chiave della dimensione critica: "bias", "dati", "pensiero", "critico", "ia", "immagini", "informazioni", "studenti", "chatgpt", "pregiudizi". I diversi topic sembrano esplorare le interrelazioni tra questi concetti: l'origine dei bias nei dati e i loro effetti sugli output prodotti (Topic 1, 4), le strategie per sviluppare il pensiero critico nella valutazione di informazioni e immagini generate dall'IA (Topic 2, 5), e l'analisi critica di strumenti specifici (come ChatGPT, menzionato esplicitamente) e delle loro performance (Topic 3). L'analisi LDA conferma quindi, a livello strutturale, la focalizzazione del discorso sulla valutazione critica dell'IA, dei suoi processi interni (dati, algoritmi) e dei suoi prodotti esterni (immagini, informazioni), e sulla necessità di educare gli studenti a questa competenza.

In conclusione, l'analisi integrata del terzo diario mostra che i partecipanti hanno sviluppato una comprensione approfondita, sfaccettata e matura della dimensione critica dell'AI literacy. Hanno acquisito non solo un lessico specifico, ma anche una consapevolezza dei meccanismi sottili attraverso cui bias e pregiudizi possono infiltrarsi e venire amplificati dai sistemi IA, e hanno riconosciuto l'importanza fondamentale del pensiero critico come antidoto e strumento di navigazione consapevole. L'apprendimento si conferma trasformativo. Le proposte didattiche riflettono pienamente questa consapevolezza acquisita, concentrandosi sull'analisi attiva dei bias attraverso esempi concreti, sulla valutazione critica degli output dell'IA (con un focus particolare su quelli visivi) e sulla promozione esplicita del pensiero critico attraverso metodologie dialogiche e partecipative (P4C, discussioni guidate, analisi di casi). La forte preferenza per scenari didattici basati

sull'interazione sociale e la discussione collettiva (Immagine 6) sottolinea la percezione della criticità come una competenza essenzialmente sociale, da costruire e praticare all'interno di una comunità di apprendimento. La maggiore coerenza percepita tra conoscenza e applicazioni suggerisce infine un forte legame tra la comprensione dei rischi critici e l'urgenza sentita della risposta educativa.

5.2.4 Analisi del diario 4: dimensione etica - navigare valori e responsabilità

Il quarto e ultimo modulo formativo, con il diario riflessivo conclusivo, ha accompagnato i partecipanti nell'esplorazione della dimensione etica dell'AI literacy. Questa dimensione rappresenta forse il livello più complesso, astratto e normativo dell'alfabetizzazione all'IA, spostando l'attenzione dalle questioni di funzionamento (operativa) e valutazione critica (critica) alle questioni fondamentali relative ai valori, ai principi morali, alla responsabilità individuale e collettiva, all'equità, alla trasparenza, alla privacy e alla governance nello sviluppo e nell'impiego delle tecnologie di Intelligenza Artificiale. L'obiettivo era stimolare nei partecipanti una riflessione approfondita sui dilemmi etici, spesso inediti, posti dall'IA (dalla responsabilità per le decisioni algoritmiche all'impatto sul lavoro, dalla privacy nell'era dei big data all'autonomia umana) e sulla necessità di promuovere un approccio allo sviluppo e all'uso di queste tecnologie che sia intrinsecamente centrato sull'uomo, rispettoso dei diritti fondamentali e orientato al benessere collettivo e alla giustizia sociale. Il diario mirava quindi a cogliere la loro comprensione dei principi etici fondamentali, la loro consapevolezza delle sfide normative e la loro capacità di progettare attività educative volte a promuovere la riflessione etica e lo sviluppo del giudizio morale negli studenti.

L'analisi di frequenza sul corpus del quarto diario (N=63 partecipanti, 26.560 parole totali dopo pulizia) evidenzia in modo inequivocabile il focus tematico sulla dimensione etica. Il termine "etica" (336 occorrenze) diventa la parola chiave più distintiva e caratterizzante di questa fase, raggiungendo una frequenza paragonabile a quella dei termini generali "intelligenza" (384) e "artificiale" (387). Altri termini ad alta frequenza includono "dimensione" (207, quasi sempre nella locuzione "dimensione etica"), "essere" (263), "può" (165), "modo" (158), ma soprattutto un nutrito gruppo di termini specifici del dominio etico-normativo: "dati" (122), la cui gestione solleva cruciali questioni di privacy e consenso, "persone" (108), poste al centro della riflessione etica, "etici" (97), "robot" (96), frequentemente utilizzati come soggetti o attori in scenari e dilemmi etici ipotetici, e termini rilevati come particolarmente significativi dall'analisi TF-IDF quali "responsabilità", "privacy", "umanità", "leggi", "decisioni". Il contesto educativo rimane, come sempre, una cornice centrale ("attività": 198, "bambini": 159, "classe": 137, "studenti": 120, "gruppo": 122), indicando che la riflessione sui principi etici astratti è costantemente orientata alla sua traducibilità didattica e alla sua discussione con gli studenti. Il Type-Token Ratio (TTR), pari a 0.1926, risulta il più basso registrato tra i quattro diari. Questo indica un linguaggio estremamente focalizzato e convergente. Tale elevata specificità lessicale potrebbe riflettere sia la complessità intrinseca dei concetti etici, che richiede un uso attento e preciso della terminologia per evitare ambiguità, sia una possibile maggiore condivisione di linguaggio e di preoccupazioni tra i partecipanti nell'affrontare questioni normative fondamentali e universalmente rilevanti.

Le metafore utilizzate per descrivere le credenze iniziali sulla dimensione etica presentano un pattern peculiare e particolarmente significativo. Le due immagini concettuali più frequenti sono il "giardino" (ben 7 occorrenze) e la "bussola" (6 occorrenze), seguite a distanza da "fonte" (3), "scatola" (2), "oceano" (2), "labirinto" (2). La notevole popolarità della metafora del "giardino" suggerisce una

percezione dell'etica dell'IA non (o non solo) come un campo minato irto di pericoli o un mistero oscuro e indecifrabile, ma piuttosto come un ecosistema potenzialmente fertile e benefico, che tuttavia richiede una cura attiva, una coltivazione attenta e una responsabilità costante da parte dell'uomo per garantirne uno sviluppo positivo, armonioso e allineato ai valori umani, evitando la crescita incontrollata di "erbacce" (conseguenze negative impreviste, usi impropri, derive disumanizzanti). Analogamente, la frequenza della "bussola" indica una forte consapevolezza, fin dall'approccio iniziale a questa dimensione, della necessità imprescindibile di principi guida chiari, di un orientamento valoriale solido per navigare la complessità ("oceano", "labirinto" suggeriscono ancora questa complessità) delle decisioni etiche in un campo così nuovo e in rapida evoluzione. Rispetto alle percezioni iniziali delle altre dimensioni, qui sembra prevalere un senso di responsabilità attiva e la consapevolezza che l'etica non è un dato di fatto o un insieme di regole predefinite da applicare passivamente, ma qualcosa che va costruito, orientato, negoziato e curato nel tempo attraverso un impegno consapevole.

L'analisi della sezione "Conoscenza Attuale", tramite TF-IDF, conferma l'acquisizione di una comprensione specifica dei concetti etici fondamentali e delle sfide correlate. Termini come "etica", "ia", "decisioni", "dati", "privacy", "responsabilità", "etici", "sistemi", "umanità", "leggi" (spesso con riferimento alle celebri Tre Leggi della Robotica di Isaac Asimov, utilizzate come modello paradigmatico per la riflessione sulla necessità di principi guida), "benefico", "robustezza", "legalità", "equità" (questi ultimi evidenziati come termini esclusivi particolarmente rilevanti dal TF-IDF) diventano centrali nel discorso. Questo indica che i partecipanti hanno approfondito la loro comprensione dei principi etici cardine che dovrebbero guidare lo sviluppo e l'uso dell'IA (come la non maleficenza, la giustizia/equità, il rispetto dell'autonomia umana, la trasparenza, la spiegabilità), delle questioni cruciali legate alla governance (la necessità di quadri normativi e leggi appropriate, la difficile attribuzione delle responsabilità in caso di errori algoritmici) e delle sfide specifiche come la tutela della privacy dei dati personali nell'era dei Big Data, la prevenzione delle discriminazioni sistemiche (equità) e la garanzia di affidabilità, sicurezza e controllo dei sistemi IA (robustezza). L'analisi lessicale comparativa supporta pienamente questa lettura: termini utilizzati quasi esclusivamente in questa sezione come "citazione" (spesso di principi etici fondamentali o delle leggi di Asimov), "sistema" (IA come sistema socio-tecnico con implicazioni etiche), "legalità", "bluwashing" (termine critico che indica una patina etica superficiale), "robustezza", "rispettare" (valori, diritti), "zero" (riferimento alla Legge Zero di Asimov, che pone il benessere dell'umanità al di sopra di tutto), "benefico", "umanità", "giustizia", "asimov", "equità", "maleficenza" dimostrano l'appropriazione di un vocabolario normativo, filosofico e concettuale specifico per la discussione etica sull'IA.

La sezione "Applicazioni Future" mostra come i partecipanti traducano la complessa riflessione sui principi etici in metodologie didattiche attive, dialogiche e centrate sulla deliberazione morale e sulla costruzione di giudizi di valore. Il lessico identificato dal TF-IDF mantiene la centralità di "etica", "ia", "dimensione", ma la abbina in modo significativo a termini legati all'interazione sociale e a specifiche pratiche didattiche: "bambini", "attività", "robot" (spesso utilizzato come protagonista di scenari e dilemmi etici da discutere), "classe", "studenti", "gruppo", "gruppi", e in modo particolarmente rilevante, "debate". L'analisi dei termini esclusivi di questa sezione è particolarmente rivelatrice delle strategie immaginate: "insegnante", "primaria", "discutere", "proporre", "argomentazioni", "storia", "giuria", "quinta" (classe), "chiedere", "gioco", "debate", "scenari", "pro",

"contro". Emerge con forza l'intenzione di affrontare la dimensione etica in classe non tramite lezioni frontali o la semplice trasmissione di regole astratte, ma attraverso metodologie che promuovono il confronto attivo, la discussione critica, l'argomentazione ragionata e la costruzione attiva di giudizi di valore. Le attività proposte includono frequentemente:

- Discussioni guidate su dilemmi etici specifici legati all'IA (es. dilemma del carrello autonomo, responsabilità per errori medici di IA, uso di IA nella sorveglianza).
- Analisi critica di casi reali o fittizi (storie, articoli di cronaca, scenari ipotetici) che presentano implicazioni etiche.
- Simulazioni e role-playing, in cui gli studenti possono impersonare diversi attori (es. sviluppatori, utenti, robot, legislatori, membri di una giuria) chiamati a prendere decisioni eticamente rilevanti.
- E, in modo particolarmente ricorrente e significativo, il debate strutturato, in cui gli studenti sono chiamati a ricercare, preparare e sostenere diverse posizioni ("pro", "contro") su questioni etiche controverse legate all'IA, sviluppando "argomentazioni" basate su principi etici e confrontandosi dialetticamente.

L'obiettivo implicito e condiviso di queste proposte sembra essere quello di educare alla responsabilità, alla riflessione morale autonoma e alla capacità di prendere decisioni etiche consapevoli e argomentate in relazione alle sfide poste dall'Intelligenza Artificiale.

La distribuzione delle immagini scelte per rappresentare lo scenario d'uso didattico preferito per la dimensione etica è la più polarizzata e significativa tra tutti e quattro i diari. L'Immagine 6, che rappresenta un gruppo di persone impegnate in una discussione, una presentazione o un lavoro collettivo attorno a uno schermo o una lavagna, è scelta da una maggioranza schiacciante (26 partecipanti su 61). Segue, a notevole distanza, l'Immagine 4 (laboratorio collaborativo, 16 scelte), che potrebbe essere interpretata come un contesto adatto per simulazioni etiche di gruppo o analisi collaborative di dati e casi, e l'Immagine 5 (analisi foto/storia, 10 scelte), pertinente per la discussione approfondita di casi specifici presentati visivamente o narrativamente. Le immagini che evocano primariamente l'apprendimento individuale (Immagine 3, solo 5 scelte; Immagine 1 e 2, appena 2 scelte ciascuna) diventano quasi irrilevanti per questa dimensione. Questa netta preferenza per l'Immagine 6 suggerisce in modo potente che la dimensione etica dell'IA viene percepita dai partecipanti come intrinsecamente sociale, dialogica e deliberativa. L'apprendimento e la riflessione sull'etica non sembrano essere visti come processi solitari o puramente intellettuali, ma come attività che richiedono il confronto collettivo, la discussione aperta di valori e principi, la negoziazione di significati e la costruzione condivisa di un orientamento morale all'interno di una comunità di apprendimento come la classe.

L'analisi della similarità nel quarto e ultimo diario completa il quadro della trasformazione lungo l'intero percorso formativo. Nell'analisi INTER-diario (trasformazione collettiva) le medie di similarità coseno si mantengono molto basse (Credenze-Conoscenza: 0.081; Conoscenza-Applicazioni: 0.075; Credenze-Applicazioni: 0.066), confermando in modo definitivo che l'apprendimento trasformativo ha permeato l'intero percorso formativo, inclusa la complessa e sfidante dimensione etica. Mentre nell'analisi INTRA-diario (coerenza individuale) l'alta coerenza interna (~0.79) dimostra che i partecipanti hanno ristrutturato il loro modo di esprimersi e di concettualizzare l'etica dell'IA, integrando principi, sfide e un vocabolario normativo specifico,

mantenendo al contempo la capacità di elaborare ragionamenti etici articolati e logicamente strutturati.

La similarità INTER-diario tra "Conoscenza Attuale" e "Applicazioni Future" (0.075) si colloca in una posizione intermedia rispetto ai diari precedenti. Questa posizione, interpretata insieme all'alta coerenza INTRA-diario, potrebbe suggerire un collegamento relativamente efficace tra la comprensione dei principi etici fondamentali (pur complessi) e la capacità di ideare attività didattiche pertinenti e specifiche (come i debate, l'analisi di scenari, le discussioni guidate). La traduzione dell'etica in pratica didattica sembra essere percepita come fattibile e necessaria, sebbene forse richieda metodologie specifiche e non sempre dirette come quelle per le dimensioni più tecniche.

Infine, i topic identificati dall'analisi LDA per questo diario ruotano coerentemente attorno ai concetti chiave della dimensione etica: "dimensione", "ia", "etica", "etici", "può", "potrebbe", "bambini", "studenti", "classe", "gruppo", "robot", "dati", "privacy", "decisioni", "responsabilità". I diversi topic sembrano catturare le sfaccettature della riflessione etica applicata: le implicazioni generali e le potenzialità/rischi da bilanciare (Topic 1, 2), le questioni legate all'uso concreto, alle decisioni algoritmiche e all'attribuzione della responsabilità (Topic 3), il ruolo fondamentale del contesto sociale ed educativo nel plasmare il dibattito etico (Topic 4), e le preoccupazioni specifiche relative alla gestione dei dati, alla tutela della privacy e alle scelte etiche concrete da compiere (Topic 5). L'analisi LDA conferma quindi il focus sulla riflessione etica applicata all'IA, con una forte enfasi sulle sue implicazioni sociali, educative e sulla necessità di un approccio normativo consapevole.

In conclusione, l'analisi integrata del quarto e ultimo diario dimostra che i partecipanti hanno completato il loro percorso formativo sviluppando una comprensione profonda, articolata e matura della cruciale dimensione etica dell'AI literacy. Hanno acquisito consapevolezza dei principi etici fondamentali, delle complesse sfide legate alla responsabilità, alla privacy, all'equità e alla governance, superando una percezione iniziale che, pur riconoscendone la complessità, vedeva l'etica come un campo bisognoso di guida e cura attiva ("giardino", "bussola"). L'apprendimento trasformativo è confermato per l'ultima volta dalle metriche di similarità. Le proposte didattiche si concentrano in modo deciso su metodologie attive, dialogiche e deliberative, come il debate, l'analisi di casi, il role-playing e la discussione guidata di dilemmi etici, con l'obiettivo esplicito di costruire negli studenti una consapevolezza critica e una capacità di giudizio morale autonomo. La preferenza schiacciante per scenari didattici basati sull'interazione sociale e la discussione collettiva (Immagine 6) sottolinea la percezione dell'educazione etica come un processo eminentemente comunitario, dialogico e volto alla costruzione condivisa di valori nell'era dell'IA.

5.3 Discussione integrata dei risultati della sperimentazione

5.3.1 Discussione dei risultati quantitativi

I risultati dell'analisi statistica forniscono evidenze robuste dell'efficacia dell'intervento formativo nel migliorare l'AI literacy dei partecipanti. L'osservazione di incrementi statisticamente significativi in tutte le aree e dimensioni valutate, accompagnati da dimensioni dell'effetto eccezionalmente grandi, suggerisce un impatto sostanziale dell'intervento sulle conoscenze e competenze relative all'IA.

Particolarmente significativo appare il miglioramento nella dimensione Conoscitiva, che mostra la dimensione dell'effetto più elevata ($d = 2.46$). Questo risultato indica che l'intervento è stato particolarmente efficace nel trasmettere le conoscenze teoriche e matematiche relative all'IA, aree in cui i partecipanti mostravano inizialmente le maggiori lacune. La comprensione dei fondamenti teorici dell'IA e delle funzioni matematiche alla base degli algoritmi rappresenta un elemento cruciale dell'AI literacy, in quanto fornisce la base concettuale necessaria per comprendere criticamente e utilizzare consapevolmente le tecnologie basate sull'IA (Long & Magerko, 2020).

Anche la dimensione Operativa ha mostrato un miglioramento molto consistente ($d = 2.13$), suggerendo che l'intervento ha efficacemente sviluppato nei partecipanti la capacità di applicare le conoscenze acquisite e di comprendere i processi di progettazione delle applicazioni di IA. Questo aspetto è particolarmente rilevante in una prospettiva di alfabetizzazione che non si limiti alla conoscenza teorica, ma includa anche la capacità di interagire attivamente con le tecnologie di IA (Ng, 2021).

La dimensione Critica ($d = 2.04$) e la dimensione Etica ($d = 1.31$), pur mostrando dimensioni dell'effetto relativamente più contenute rispetto alle altre dimensioni, presentano comunque valori che eccedono ampiamente la soglia convenzionale per un effetto grande ($d = 0.8$). La capacità di valutare criticamente le applicazioni dell'IA e di comprenderne le implicazioni etiche rappresenta un aspetto fondamentale per un utilizzo consapevole e responsabile di queste tecnologie (Floridi et al., 2018). È interessante notare che la dimensione Etica presentava il punteggio iniziale più elevato, suggerendo una maggiore sensibilità di base verso le questioni etiche dell'IA, probabilmente alimentata dal dibattito pubblico su questi temi.

L'indice globale dell'AI literacy mostra un effetto straordinariamente ampio ($d = 2.48$), che si colloca ben oltre i valori tipicamente osservati in interventi educativi. A questo proposito, è importante sottolineare che il Reliable Change Index (RCI) indica che il 100% dei partecipanti ha mostrato un cambiamento clinicamente significativo, suggerendo che l'intervento ha prodotto miglioramenti reali e sostanziali a livello individuale, non attribuibili all'errore di misurazione.

I risultati dell'Equivalence Testing (TOST) confermano ulteriormente la rilevanza pratica degli effetti osservati. Per tutte le aree e dimensioni, il test TOST rifiuta con decisione l'ipotesi di equivalenza a un effetto trascurabile ($p \approx 10^{-12}$), indicando che i miglioramenti osservati sono non solo statisticamente significativi, ma anche praticamente rilevanti.

L'analisi della distribuzione dei cambiamenti offre ulteriori spunti interpretativi. La percentuale estremamente elevata di partecipanti che hanno mostrato un miglioramento in tutte le aree e dimensioni (dal 79% al 100%) suggerisce che l'intervento formativo è stato efficace trasversalmente, indipendentemente dal livello iniziale di conoscenza. Particolarmente notevole è il fatto che il 100% dei partecipanti ha mostrato un miglioramento nell'area dei fondamenti teorici, con totale assenza di invariati o peggiorati, indicando che l'intervento non ha generato confusione o misconcezioni in questa area particolarmente tecnica.

Il valore Gamma di Rosenbaum ($\Gamma = 1.4$) suggerisce che l'effetto dell'intervento rimarrebbe significativo anche in presenza di un fattore confondente non misurato che aumenti del 40% le odds di esposizione all'intervento. Questo indicatore fornisce un'ulteriore evidenza della robustezza dei risultati rispetto a potenziali bias di selezione o variabili confondenti non considerate nell'analisi.

I risultati sulle variabili aggiuntive, come la familiarità con il concetto di IA e l'utilizzo di applicazioni di IA, indicano che l'intervento ha avuto un impatto non solo sulle conoscenze e competenze esplicite, ma anche sulla consapevolezza più generale dell'IA e delle sue applicazioni nella vita quotidiana. Questo aspetto è particolarmente rilevante in una prospettiva di alfabetizzazione funzionale, che mira a sviluppare la capacità di riconoscere e interagire consapevolmente con l'IA in contesti reali.

I cambiamenti nelle percezioni dell'IA, sebbene meno marcati in termini di dimensione dell'effetto rispetto ai cambiamenti nelle conoscenze e competenze, sono comunque significativi e coerenti nella direzione. L'intervento sembra aver promosso una visione più equilibrata e informata dell'IA, riducendo sia l'eccessivo entusiasmo sia l'eccessiva preoccupazione, in linea con l'obiettivo di sviluppare un approccio critico ma costruttivo a queste tecnologie.

Le analisi correlazionali offrono ulteriori spunti sull'interazione tra diverse variabili. La correlazione negativa tra il livello iniziale e l'entità del miglioramento in alcune dimensioni suggerisce che l'intervento è stato particolarmente efficace nel colmare le lacune dei partecipanti con minori conoscenze di partenza. Allo stesso tempo, le correlazioni significative tra le dimensioni dell'AI literacy e le percezioni dell'IA indicano che lo sviluppo di conoscenze e competenze è accompagnato da un cambiamento nelle attitudini, verso una maggiore apertura e fiducia nelle potenzialità dell'IA.

Le analisi per sottogruppi hanno evidenziato la robustezza dell'efficacia dell'intervento attraverso diverse caratteristiche demografiche e professionali. L'assenza di differenze significative nell'entità del miglioramento tra educatori e non educatori, tra donne e uomini, tra gruppi di età diversa e tra livelli di istruzione differenti suggerisce che l'approccio adottato è sufficientemente versatile da rispondere alle esigenze di un pubblico eterogeneo. Questo è un risultato importante in una prospettiva di diffusione dell'AI literacy a livello sociale.

L'unica variabile che mostra un'associazione significativa con i livelli assoluti di AI literacy è la familiarità tecnologica generale, che risulta correlata sia con il livello iniziale che con quello finale, ma non con l'entità del miglioramento. Questo suggerisce che, sebbene un background tecnologico più solido possa fornire un vantaggio in termini di punto di partenza, non rappresenta un prerequisito per beneficiare dell'intervento formativo.

Nel complesso, i risultati dell'analisi statistica supportano fortemente l'efficacia dell'intervento formativo nel promuovere l'AI literacy nelle sue diverse dimensioni, e suggeriscono che l'approccio adottato è appropriato per un pubblico diversificato in termini di background professionale, caratteristiche demografiche e familiarità tecnologica.

I risultati di questa analisi hanno importanti implicazioni per lo sviluppo dell'AI literacy. In primo luogo, dimostrano che anche interventi formativi relativamente brevi (12 ore complessive) possono produrre miglioramenti significativi nella comprensione dell'IA, con dimensioni dell'effetto che superano notevolmente quelle tipicamente osservate in contesti educativi. Questo risultato è particolarmente incoraggiante in un contesto in cui la rapida diffusione delle tecnologie di IA richiede interventi formativi accessibili e scalabili (Roll & Wylie, 2016).

La particolare efficacia dell'intervento nel migliorare la dimensione Conoscitiva suggerisce l'importanza di focalizzare gli sforzi formativi sulla trasmissione dei concetti fondamentali dell'IA, che possono apparire inizialmente astratti o complessi. La comprensione delle basi teoriche e matematiche dell'IA rappresenta infatti un prerequisito per lo sviluppo di un'alfabetizzazione più

ampia, che includa anche gli aspetti applicativi, critici ed etici (Wang, 2019). Le dimensioni dell'effetto eccezionalmente elevate osservate in questa dimensione ($d = 2.46$) suggeriscono che anche concetti potenzialmente complessi possono essere efficacemente trasmessi attraverso approcci didattici appropriati.

Il consistente miglioramento nella dimensione Operativa ($d = 2.13$) evidenzia come la capacità di applicare le conoscenze acquisite possa essere efficacemente sviluppata nel contesto di interventi formativi strutturati. Questo aspetto è particolarmente rilevante in una prospettiva di alfabetizzazione funzionale, che miri non solo alla comprensione teorica, ma anche alla capacità di interagire concretamente con le tecnologie di IA (Bocconi et al., 2022). L'efficacia dell'intervento in questa dimensione suggerisce l'importanza di integrare attività pratiche e interattive nel processo formativo.

L'incremento nelle dimensioni Critica ($d = 2.04$) ed Etica ($d = 1.31$), pur con dimensioni dell'effetto relativamente più contenute rispetto alle altre dimensioni, conferma l'importanza di integrare negli interventi formativi sull'IA anche riflessioni sulle implicazioni sociali, culturali ed etiche di queste tecnologie. La capacità di valutare criticamente le applicazioni dell'IA e di comprenderne le implicazioni etiche rappresenta infatti un elemento essenziale per un utilizzo consapevole e responsabile di queste tecnologie, in linea con le recenti raccomandazioni dell'UNESCO (2021) sull'etica dell'IA.

Il fatto che il 100% dei partecipanti abbia mostrato un cambiamento significativo secondo il Reliable Change Index suggerisce che l'approccio formativo adottato ha una efficacia pervasiva, capace di produrre miglioramenti rilevanti in tutti i partecipanti, indipendentemente dalle loro caratteristiche individuali. Questo risultato ha importanti implicazioni per la progettazione di interventi formativi orientati alla diffusione dell'AI literacy su scala più ampia.

I cambiamenti osservati nelle percezioni dell'IA suggeriscono che lo sviluppo dell'AI literacy contribuisce a promuovere una visione più equilibrata e informata di queste tecnologie, riducendo sia l'eccessivo entusiasmo sia l'eccessiva preoccupazione. Questo risultato è particolarmente rilevante in un contesto caratterizzato da narrazioni spesso polarizzate sull'IA, che tendono a oscillare tra l'utopismo tecnologico e il catastrofismo (Cave et al., 2019). L'alfabetizzazione sull'IA può quindi svolgere un ruolo importante nel facilitare un dibattito pubblico più informato e costruttivo su queste tecnologie.

La riduzione osservata nel gap di genere nella percezione della pericolosità dell'IA suggerisce che l'AI literacy può contribuire a ridurre le disparità di genere nelle attitudini verso queste tecnologie. Considerando la persistente sottorappresentazione femminile nei campi relativi all'IA (UNESCO, 2020), questo risultato ha implicazioni potenzialmente rilevanti per promuovere una maggiore inclusività in questo ambito.

Le analisi per sottogruppi offrono indicazioni utili per la progettazione di interventi formativi sull'IA. L'efficacia trasversale dell'intervento, indipendentemente dal background professionale, genere, età e livello di istruzione, suggerisce la possibilità di adottare approcci relativamente standardizzati per pubblici diversificati. Allo stesso tempo, le differenze osservate in termini di livelli assoluti di AI literacy (ad esempio in relazione alla familiarità tecnologica) suggeriscono l'importanza di calibrare adeguatamente il livello di partenza degli interventi, per rispondere alle esigenze di partecipanti con diversi background.

L'impatto positivo dell'intervento sulle percezioni dell'IA nell'insegnamento evidenzia il potenziale dell'AI literacy nel contesto educativo. Lo sviluppo di conoscenze e competenze relative all'IA sembra promuovere una maggiore apertura verso l'integrazione di queste tecnologie nei processi di insegnamento e apprendimento, pur mantenendo una consapevolezza critica dei loro limiti e delle implicazioni della loro adozione. Questo suggerisce l'importanza di promuovere l'AI literacy tra gli educatori, come passo preliminare per un'integrazione consapevole dell'IA nei contesti educativi.

Nel complesso, i risultati di questa analisi supportano l'importanza di promuovere l'AI literacy come competenza trasversale, rilevante per cittadini con diversi profili professionali e formativi. In un contesto caratterizzato dalla crescente pervasività delle tecnologie di IA in ambiti sempre più diversificati della vita quotidiana e professionale, lo sviluppo di un'alfabetizzazione adeguata rappresenta una sfida educativa cruciale, a cui interventi formativi ben strutturati possono efficacemente contribuire.

Nonostante i risultati incoraggianti, è importante riconoscere alcuni limiti dell'analisi condotta. In primo luogo, l'assenza di un gruppo di controllo non permette di escludere completamente che i miglioramenti osservati siano dovuti a fattori esterni all'intervento formativo, come l'esposizione a informazioni sull'IA attraverso altre fonti o un effetto di apprendimento dovuto alla ripetizione del test (Shadish et al., 2002). Tali valori eccezionalmente alti possono essere spiegati dalla bassa conoscenza pregressa dei partecipanti (floor effect) e dalla stretta aderenza tra i contenuti del corso e gli item del questionario, trattandosi di conoscenze tecniche specifiche (es. reti neurali) difficilmente acquisibili altrove. Tuttavia, la magnitudine degli effetti osservati ($d > 2.0$ in molte dimensioni), la loro coerenza tra le diverse aree e dimensioni, e i risultati dell'analisi di sensibilità di Rosenbaum ($\Gamma = 1.4$) suggeriscono che l'intervento abbia effettivamente giocato un ruolo determinante.

Un secondo limite riguarda la natura auto-valutativa delle misure utilizzate. La percezione soggettiva delle proprie conoscenze e competenze potrebbe non riflettere accuratamente il reale livello di alfabetizzazione sull'IA (Kruger & Dunning, 1999). L'integrazione di misure oggettive, come test di conoscenza o valutazioni di performance, potrebbe fornire una visione più completa dell'efficacia dell'intervento.

La rilevazione post-intervento è stata effettuata immediatamente dopo la conclusione dell'intervento formativo. Sarebbe interessante valutare la persistenza dei miglioramenti osservati attraverso una rilevazione di follow-up a distanza di tempo, per verificare la stabilità dell'apprendimento e l'eventuale trasferimento delle conoscenze e competenze acquisite in contesti reali (Arthur et al., 1998).

Il campione utilizzato, sebbene diversificato in termini di caratteristiche demografiche e professionali, è relativamente limitato numericamente ($n = 62$) e basato su una partecipazione volontaria all'intervento formativo. Questo potrebbe introdurre un bias di selezione, con una sovrarappresentazione di individui particolarmente interessati all'IA e motivati ad apprendere su questo tema. La generalizzabilità dei risultati a popolazioni più ampie e diversificate richiede quindi cautela.

Un ulteriore limite metodologico riguarda la non normalità delle differenze pre-post per la maggior parte delle variabili analizzate. Sebbene l'utilizzo di test non parametrici (Wilcoxon) rappresenti un approccio appropriato per gestire questa violazione delle assunzioni, tali test sono generalmente

considerati meno potenti dei loro equivalenti parametrici. Tuttavia, la netta significatività dei risultati (tutti i $p < 0.001$) suggerisce che questo limite non abbia compromesso la capacità di rilevare gli effetti dell'intervento.

L'intervento formativo valutato rappresenta uno specifico approccio all'AI literacy, basato su determinati principi pedagogici e contenuti. I risultati potrebbero non essere direttamente trasferibili ad altri interventi con caratteristiche diverse, e sarebbe interessante condurre studi comparativi per identificare gli approcci più efficaci in relazione a diversi obiettivi e contesti formativi.

L'analisi statistica condotta, pur approfondita, ha dovuto necessariamente focalizzarsi su alcuni aspetti chiave, tralasciandone altri. Ad esempio, non abbiamo esplorato in dettaglio le interazioni tra le diverse dimensioni dell'AI literacy nel determinare i cambiamenti nelle percezioni dell'IA, o il ruolo potenziale di variabili moderatrici non misurate, come i tratti di personalità o gli stili di apprendimento.

Nella valutazione dei cambiamenti nelle percezioni dell'IA, abbiamo analizzato le affermazioni individualmente, senza aggregarle in scale composite. Questo approccio, pur fornendo una visione granulare, potrebbe risultare meno robusto rispetto a un'analisi basata su costrutti latenti derivati da più item, specialmente considerando il numero relativamente elevato di test condotti, che aumenta il rischio di errori di tipo I.

Nonostante questi limiti, che suggeriscono la necessità di ulteriori ricerche per consolidare e approfondire i risultati ottenuti, l'analisi condotta fornisce evidenze robuste dell'efficacia dell'intervento formativo nel promuovere l'AI literacy nelle sue diverse dimensioni, e offre spunti significativi per la progettazione di futuri interventi in questo ambito.

5.3.2 Discussione dei risultati qualitativi

L'analisi dettagliata e integrata dei quattro diari riflessivi, ciascuno dedicato a una dimensione specifica dell'AI literacy – conoscitiva, operativa, critica ed etica – consente ora di tracciare un quadro complessivo dell'evoluzione delle conoscenze, delle percezioni, delle attitudini e delle proiezioni didattiche dei partecipanti lungo l'intero percorso formativo. Facendo dialogare le evidenze quantitative emerse dall'analisi NLP con le interpretazioni qualitative dei testi, emergono pattern significativi e ricorrenti che illuminano sia le traiettorie di trasformazione individuale sia l'adattamento collettivo del focus tematico e metodologico alle diverse sfide poste da ciascuna dimensione. Questa visione d'insieme offre una comprensione ricca e sfumata dell'impatto della sperimentazione e dei processi di apprendimento attivati.

Innanzitutto, il risultato più metodologicamente robusto, confermato trasversalmente da tutte le analisi di similarità (sia coseno che Jaccard) condotte sui quattro diari, è la profonda e costante discontinuità lessicale e concettuale osservata tra le credenze iniziali espresse dai partecipanti all'inizio di ogni modulo e le conoscenze acquisite al termine dello stesso, nonché tra queste ultime e le applicazioni didattiche successivamente proposte. Le medie di similarità coseno, come discusso in dettaglio per ciascun diario, si attestano costantemente su valori estremamente bassi, molto vicini allo zero. Questa radicale dissimilarità tra le diverse fasi della riflessione fornisce una prova quantitativa forte e convergente dell'efficacia del percorso formativo nel promuovere non un apprendimento

superficiale o meramente additivo di nozioni, ma un processo di ristrutturazione cognitiva e linguistica profonda. I partecipanti sembrano aver vissuto, in relazione a ciascuna dimensione dell'AI literacy, un autentico apprendimento trasformativo, secondo la definizione di Mezirow (1991): un processo che implica la messa in discussione critica delle proprie prospettive di significato iniziali, l'integrazione di nuovi schemi concettuali e la conseguente modifica significativa del proprio modo di esprimersi – e, presumibilmente, di pensare e agire – riguardo all'Intelligenza Artificiale. L'analisi dei termini esclusivi per ciascuna sezione in ogni diario ha ulteriormente dettagliato questa trasformazione, mostrando il passaggio progressivo da un linguaggio spesso vago, carico di metafore legate all'ignoto o al timore iniziale, a un vocabolario via via più tecnico e preciso (Diario 2), poi più analitico e critico (Diario 3) e infine più normativo ed etico (Diario 4), sempre più orientato verso la riflessione sull'applicazione didattica.

In secondo luogo, l'analisi comparativa delle metafore utilizzate dai partecipanti all'inizio di ciascun modulo rivela un'interessante traiettoria evolutiva nella percezione dell'IA man mano che se ne esplorano le diverse dimensioni. Si parte da una visione dell'IA come concetto generale percepito come vasto, misterioso, complesso e talvolta minaccioso o incontrollabile (le metafore dominanti nel Diario 1 sono quelle del viaggio/esplorazione in territori ignoti e dell'entità autonoma potente e ambivalente). Si passa poi a una percezione della sua dimensione operativa come una sfida sì complessa, ma intrinsecamente logica, strutturata e decifrabile con gli strumenti appropriati (le metafore del Diario 2 sono quelle dello strumento/risorsa e del problema strutturato come labirinto o puzzle). Successivamente, la dimensione critica viene percepita come un'area con complessità nascoste, problemi latenti ("iceberg", "scatola chiusa") che richiedono uno sforzo attivo di svelamento, analisi e interpretazione critica (le metafore del Diario 3 sono quelle dell'amplificazione/riflessione e della necessità di cura). Infine, la dimensione etica viene vista primariamente come un campo che necessita di cura attiva, responsabilità costante e orientamento valoriale chiaro per garantirne uno sviluppo positivo (le metafore del Diario 4 sono quelle dell'organismo evolutivo/giardino e della bussola/guida). Questa progressione suggerisce un affinamento e una maturazione della percezione: dall'iniziale disorientamento e timore reverenziale di fronte al concetto generale, si passa a un approccio più analitico e fiducioso verso i meccanismi operativi, per poi sviluppare una consapevolezza delle criticità non immediatamente evidenti e, infine, riconoscere la necessità fondamentale di una guida etica e di una responsabilità attiva nella gestione di questa potente tecnologia.

In terzo luogo, l'analisi delle frequenze lessicali e dei termini TF-IDF attraverso i quattro diari conferma una progressiva specializzazione del vocabolario e un adattamento del focus tematico in stretta corrispondenza con la dimensione dell'AI literacy affrontata in ciascun modulo. Il linguaggio utilizzato dai partecipanti diventa via via più preciso, tecnico e specifico man mano che si addentrano negli aspetti operativi, critici ed etici. Questo è evidenziato sia dall'emergere di termini chiave specifici e distintivi in ciascun diario (come "algoritmo" e "dati" nel Diario 2, "bias" e "critico" nel Diario 3, "etica" e "responsabilità" nel Diario 4), sia dalla leggera ma costante diminuzione del Type-Token Ratio (TTR) attraverso i diari, che suggerisce una crescente convergenza su un lessico condiviso e specializzato necessario per discutere con precisione i concetti specifici di ciascuna dimensione. Questo indica che i partecipanti non solo hanno appreso nuovi contenuti, ma hanno anche sviluppato la capacità di utilizzare un linguaggio appropriato, differenziato e specifico per ciascuna delle complesse dimensioni dell'AI literacy.

In quarto luogo, emerge come uno dei risultati più significativi e incoraggianti l'evidente capacità dei partecipanti di adattare le metodologie didattiche proposte e le visioni d'uso (rappresentate dalla scelta delle immagini) in funzione della specifica dimensione dell'AI literacy trattata. Non si osserva l'applicazione di un approccio didattico unico e indifferenziato per tutti gli argomenti, ma piuttosto una notevole flessibilità e maturità pedagogica nel selezionare e privilegiare le strategie ritenute più efficaci per ciascun obiettivo di apprendimento. Come evidenziato nelle analisi dei singoli diari, si passa da un equilibrio tra riflessione individuale, pratica collaborativa e uso di strumenti per la dimensione conoscitiva (Diario 1), a una netta predominanza di attività pratiche, laboratoriali e collaborative (coding, giochi unplugged) per la dimensione operativa (Diario 2), per poi spostarsi verso approcci basati sulla discussione collettiva, l'analisi critica di artefatti e la riflessione guidata per la dimensione critica (Diario 3), fino a una quasi assoluta preferenza per la discussione di gruppo, il debate e la costruzione sociale di significato per affrontare la dimensione etica (Diario 4). Questa capacità di modulare consapevolmente l'approccio didattico in base alla natura epistemologica del contenuto (concettuale, procedurale, critico, valoriale) è un indicatore importante di una competenza professionale avanzata e suggerisce che il percorso formativo ha stimolato non solo l'apprendimento sull'IA, ma anche una riflessione meta-pedagogica profonda su come insegnarla efficacemente ai propri studenti.

In quinto luogo, nonostante la chiara specificità di ciascun modulo, emerge con forza la centralità trasversale attribuita dai partecipanti alla riflessione critica ed etica. La metodologia della Philosophy for Children (P4C), ad esempio, viene citata frequentemente non solo nel primo diario come strumento per introdurre il concetto di IA, ma anche come approccio privilegiato per affrontare le dimensioni critica ed etica nei diari successivi. Termini come "pensiero critico" mantengono una rilevanza significativa attraverso più diari. Questo suggerisce che i partecipanti hanno profondamente interiorizzato l'idea che l'AI literacy non possa essere ridotta a un insieme di conoscenze tecniche o abilità operative, ma richieda intrinsecamente lo sviluppo di una disposizione critica permanente, di una consapevolezza etica vigile e di una capacità di giudizio autonomo e responsabile. Questi aspetti sembrano essere percepiti come il cuore pulsante e il fine ultimo di una vera competenza di cittadinanza digitale nell'era dell'Intelligenza Artificiale.

Inoltre, per quanto riguarda l'evoluzione della tonalità emotiva, l'analisi del sentiment, condotta utilizzando il modello Transformer neuraly/bert-base-italian-cased-sentiment, ha rivelato ulteriori sfumature. Complessivamente, la tonalità emotiva dei diari tende al positivo, con le sezioni 'Applicazioni' (media sentiment: 0.360) e 'Conoscenza' (0.289) che mostrano un sentiment medio più elevato rispetto alle 'Credenze' (0.146). È emerso che il sentiment medio della sezione 'Applicazioni' è costantemente più positivo di quello della sezione 'Credenze' in tutti e quattro i momenti del percorso. Inoltre, mentre il sentiment delle 'Credenze' non ha mostrato variazioni statisticamente significative tra i quattro diari (Kruskal-Wallis, $p=0.906$), quello delle 'Applicazioni' ha registrato una differenza significativa (Kruskal-Wallis, $p=0.005$). Il test post-hoc di Dunn ha indicato che il sentiment relativo alle applicazioni nel Diario 1 (media HF: 0.476) era significativamente più positivo rispetto a quello del Diario 4 (Etica, media HF: 0.241). Questa evoluzione potrebbe suggerire un entusiasmo iniziale per le potenzialità applicative che si affina e si modera nel corso della formazione, parallelamente all'approfondimento delle dimensioni critica ed etica, oppure riflettere la diversa natura degli argomenti trattati nei moduli che possono aver evocato risposte emotive differenti quando si è trattato di ipotizzarne l'applicazione.

Infine, l'evoluzione complessiva osservata attraverso i quattro diari può essere efficacemente interpretata alla luce dei pattern di trasformazione concettuale identificati nell'analisi qualitativa che trovano qui una solida conferma empirica integrata:

- Da estraneità a familiarità: Il percorso ha chiaramente facilitato il superamento della percezione iniziale dell'IA come entità aliena, misteriosa o minacciosa, portando a una maggiore familiarità, confidenza e senso di auto-efficacia nel confrontarsi con essa.
- Da passività ad agency: Si osserva un'evoluzione da una posizione talvolta passiva o timorosa nei confronti della tecnologia verso un ruolo più attivo, propositivo e intenzionale da parte dell'insegnante, che si sente in grado di progettarne l'uso didattico.
- Da strumento a partnership: Sebbene la visione strumentale rimanga presente, specialmente nella dimensione operativa, si assiste a una transizione verso una concezione più collaborativa e interattiva, particolarmente evidente nella discussione sulla dimensione etica, dove l'IA è vista come un interlocutore (o l'oggetto di un dialogo) nel processo educativo e sociale.
- Da tecnica a etica: Il percorso mostra uno spostamento progressivo e consapevole dell'attenzione dagli aspetti puramente tecnici e funzionali dell'IA (pur importanti) alle sue più ampie e profonde implicazioni critiche, sociali ed etiche, percepite come dimensioni fondanti dell'AI literacy.

In conclusione, l'analisi integrata dei quattro diari riflessivi, supportata dalle evidenze quantitative dell'analisi NLP e illuminata dai framework teorici della trasformazione concettuale e della metacognizione, fornisce un quadro complessivo estremamente positivo dell'impatto del percorso formativo sullo sviluppo dell'AI literacy nei partecipanti. I risultati dimostrano in modo convergente e robusto l'efficacia del curriculum sperimentato nel promuovere un apprendimento significativo, profondo e trasformativo, che ha inciso non solo sulle conoscenze dichiarative, ma anche sulle credenze implicite, sul linguaggio utilizzato e sulla capacità di progettazione didattica in relazione alle diverse e complesse dimensioni dell'Intelligenza Artificiale. La traiettoria di apprendimento osservata appare logica e coerente, muovendosi dalla demistificazione iniziale e dalla costruzione di una base conoscitiva, all'esplorazione dei meccanismi operativi attraverso approcci pratici, per poi approdare alla valutazione critica dei sistemi e dei loro output e, infine, alla riflessione sui fondamenti etici e normativi che devono guidarne lo sviluppo e l'uso. La flessibilità e la maturità pedagogica dimostrate dai partecipanti nell'adattare le strategie didattiche alle specificità di ciascuna dimensione rappresentano un risultato particolarmente prezioso, indicando che il percorso ha stimolato anche una riflessione profonda sul "come" insegnare efficacemente l'IA. La costante e trasversale enfasi sulla dimensione critica ed etica conferma infine la percezione dell'AI literacy come una competenza di cittadinanza digitale essenziale e complessa, che va ben oltre l'alfabetizzazione tecnica per abbracciare la consapevolezza sociale, la responsabilità individuale e collettiva e la capacità di giudizio morale autonomo. Questi risultati, nel loro insieme, offrono implicazioni concrete e dirette per la progettazione di futuri interventi formativi e per la definizione delle Linee Guida che verranno elaborate nella sezione conclusiva di questo capitolo.

L'analisi approfondita e integrata dei diari metacognitivi dei partecipanti al percorso formativo sull'AI literacy ha permesso di delineare un quadro articolato e incoraggiante dei processi di apprendimento e trasformazione concettuale attivati. Attraverso l'integrazione di metodi quantitativi derivati dal Natural Language Processing e di un'analisi qualitativa attenta alle sfumature di significato e alle esperienze soggettive, è stato possibile identificare pattern linguistici, strutture tematiche, modelli

metaforici e proiezioni didattiche che riflettono una significativa evoluzione nella comprensione dell'Intelligenza Artificiale e del suo ruolo nel contesto educativo.

I risultati evidenziano in modo convergente un processo di trasformazione multidimensionale, che ha coinvolto non solo la sfera cognitiva, ma anche quella emotiva e prassica dei partecipanti. Dal punto di vista cognitivo, si è osservato un progressivo arricchimento e una complessificazione dei modelli concettuali relativi all'IA, con l'integrazione graduale di considerazioni tecniche, pedagogiche, critiche ed etiche. Sul piano emotivo, sebbene non misurato sistematicamente in tutte le sue sfaccettature, l'analisi delle metafore iniziali e l'interpretazione generale dei testi suggeriscono una transizione da un atteggiamento iniziale spesso caratterizzato da incertezza, disorientamento o timore, verso una maggiore fiducia nelle proprie capacità di comprendere, valutare e integrare l'IA nella pratica educativa. Sul piano prassico, infine, si è rilevata una chiara evoluzione nelle proiezioni didattiche: dalle applicazioni più semplici e talvolta isolate immaginate all'inizio, si è passati a scenari più complessi, integrati e pedagogicamente ambiziosi, che valorizzano il pensiero critico, la collaborazione tra pari e la riflessione etica come componenti essenziali dell'AI literacy.

Particolarmente significativa si è rivelata l'evoluzione delle metafore utilizzate dai partecipanti per concettualizzare l'IA nelle sue diverse dimensioni, poiché esse riflettono un processo di appropriazione progressiva e di ri-significazione di questa tecnologia complessa. Le metafore iniziali del viaggio in territori sconosciuti o dell'incontro con entità aliene o potenti rivelano la percezione iniziale dell'IA come qualcosa di distante, complesso e potenzialmente minaccioso. Le metafore successive dello strumento, della mappa o del puzzle evidenziano un passaggio verso una visione più funzionale, analitica e orientata all'uso controllato. Le metafore dell'amplificazione, dello specchio o dell'iceberg, emerse con forza nella dimensione critica, rivelano una crescente consapevolezza della non neutralità dell'IA e delle sue implicazioni sociali nascoste. Infine, le metafore della partnership, del giardino da curare o della bussola, che caratterizzano soprattutto la riflessione sulla dimensione etica, suggeriscono una visione più matura, responsabile e integrata, che riconosce sia le potenzialità che i limiti dell'IA e la colloca in una relazione di necessaria co-evoluzione con i valori e le finalità umane.

Questa evoluzione metaforica rispecchia fedelmente i quattro pattern fondamentali di trasformazione concettuale identificati attraverso la triangolazione dei dati: un movimento da estraneità a familiarità, che porta a superare l'ansia iniziale e a sviluppare confidenza; un passaggio da passività a agency, che vede l'insegnante assumere un ruolo attivo e propositivo nell'uso della tecnologia; una transizione da strumento a partnership (o comunque a oggetto di dialogo critico), che supera una visione meramente utilitaristica; e infine, uno spostamento cruciale del focus da tecnica a etica, che riconosce le implicazioni critiche e sociali come dimensioni fondanti dell'AI literacy. Questi pattern non vanno intesi come fasi sequenziali rigide di un percorso lineare e uguale per tutti, ma piuttosto come dimensioni interrelate e ricorsive di un processo complesso di trasformazione, che può manifestarsi con tempistiche e intensità diverse nei singoli partecipanti, ma la cui presenza diffusa nel corpus ne sottolinea la rilevanza generale.

Le implicazioni che derivano da questi risultati per la progettazione e l'implementazione di percorsi formativi efficaci per gli insegnanti (e, per estensione, per gli studenti) nell'ambito dell'Intelligenza Artificiale sono significative e concrete. Emerge con chiarezza la necessità di superare approcci

formativi meramente tecnici o, all'opposto, eccessivamente teorici e astratti. Si delinea piuttosto l'esigenza di percorsi integrati che:

Adottino un approccio olistico e multidimensionale all'AI literacy, affrontando in modo equilibrato e progressivo le conoscenze concettuali, le abilità operative, le capacità di pensiero critico e la riflessione etica, riconoscendone l'interdipendenza.

Privilegino metodologie didattiche attive, esperienziali, riflessive e collaborative, capaci di coinvolgere profondamente i partecipanti in processi di scoperta guidata, analisi critica, discussione argomentata e costruzione condivisa di significato. Esempi emersi come particolarmente promettenti includono la Philosophy for Children, il debate strutturato, il coding laboratoriale (anche unplugged), l'analisi di casi reali e scenari problematici, il project-based learning che integra diverse dimensioni.

- Valorizzino la riflessione metacognitiva come strumento essenziale non solo per monitorare, ma anche per facilitare l'apprendimento trasformativo. L'uso di diari riflessivi strutturati, come in questa sperimentazione, o di altre pratiche riflessive (portfolio, discussioni di gruppo focalizzate sull'apprendimento) appare cruciale per aiutare i partecipanti a esplicitare, confrontare e ristrutturare le proprie credenze e i propri modelli mentali.
- Mirino esplicitamente a un apprendimento trasformativo, che non si accontenti della trasmissione di informazioni ma sfidi attivamente le pre-concezioni, stimoli la riflessione critica sulle proprie prospettive e promuova una ristrutturazione profonda della comprensione dell'IA e del suo ruolo nella società e nell'educazione.
- Promuovano la flessibilità e l'adattabilità didattica negli insegnanti, fornendo loro non un'unica "ricetta", ma un repertorio diversificato di strategie, strumenti e approcci metodologici, sensibilizzandoli alla necessità di scegliere e adattare le strategie più appropriate in base agli specifici obiettivi di apprendimento, alla natura della dimensione dell'AI literacy trattata e al contesto specifico della propria classe.
- Mantengano una costante e pervasiva attenzione agli aspetti critici ed etici come elemento fondante e trasversale di qualsiasi percorso di AI literacy. Queste dimensioni non dovrebbero essere relegate a moduli separati o conclusivi, ma integrate in modo ricorsivo nell'esplorazione di tutte le altre dimensioni, per promuovere una consapevolezza continua delle implicazioni valoriali e sociali della tecnologia.
- Prestino attenzione ai modelli concettuali e alle metafore utilizzate sia dai formatori sia dai partecipanti, riconoscendo che esse non sono neutre ma influenzano profondamente la comprensione e la pratica. È opportuno stimolare una riflessione esplicita sulle metafore dominanti e utilizzare consapevolmente una pluralità di immagini concettuali per illuminare aspetti diversi e complementari dell'IA.

In conclusione, l'analisi approfondita dei diari metacognitivi ha evidenziato che l'appropriazione dell'Intelligenza Artificiale nel contesto educativo non è un semplice processo lineare di acquisizione di competenze tecniche, ma una complessa e sfidante trasformazione concettuale che coinvolge profondamente diverse dimensioni dell'esperienza cognitiva, emotiva e professionale degli insegnanti. Comprendere questa complessità offre basi solide per ripensare non solo la formazione specifica sull'IA, ma più in generale i processi di accompagnamento all'innovazione tecnologica in educazione, adottando una prospettiva che valorizzi la riflessività, la criticità, la contestualizzazione e la costruzione condivisa di significato come elementi chiave per un'integrazione della tecnologia che sia realmente consapevole, pedagogicamente fondata e orientata al miglioramento dell'esperienza educativa per tutti.

È importante riconoscere alcune limitazioni intrinseche a questa ricerca. In primo luogo, l'analisi si è basata esclusivamente sui dati testuali provenienti dai diari metacognitivi, che, pur offrendo una prospettiva ricca e profonda, rappresentano una visione necessariamente mediata e parziale dei processi di trasformazione concettuale e delle pratiche reali. Sarebbe auspicabile, in futuro, integrare questa analisi con altre fonti di dati, come osservazioni dirette delle pratiche didattiche implementate dai partecipanti dopo la formazione, interviste in profondità per esplorare ulteriormente le loro esperienze e prospettive, o analisi dei prodotti didattici realizzati, al fine di triangolare ulteriormente i risultati e ottenere una comprensione ancora più completa e sfaccettata del fenomeno.

In secondo luogo, l'approccio analitico adottato, pur sforzandosi di integrare metodi quantitativi e qualitativi, presenta alcune limitazioni operative. In particolare, le tecniche NLP utilizzate, sebbene abbiano fornito insight preziosi, presentano margini di perfezionamento. L'analisi LDA, ad esempio, ha prodotto un modello con 5 topic e una coerenza c_v di 0.3831; ulteriori esplorazioni con diversi numeri di topic o iperparametri potrebbero affinare ulteriormente la struttura tematica. Per l'analisi del sentiment, è stato impiegato un modello Transformer pre-addestrato per l'italiano (neuraly/bert-base-italian-cased-sentiment), che rappresenta un approccio più sofisticato rispetto a quelli puramente lessicali. Tuttavia, anche questi modelli hanno le loro specificità e l'integrazione con lessici di dominio o la validazione su corpora specifici potrebbe ulteriormente migliorarne l'accuratezza contestuale. Analogamente, l'analisi qualitativa delle metafore e dei temi, per sua natura interpretativa, potrebbe beneficiare di procedure di validazione intersoggettiva più strutturate, come la codifica indipendente da parte di più ricercatori o la discussione dei risultati con i partecipanti stessi (*member checking*).

In terzo luogo, la dimensione temporale dell'indagine, pur coprendo l'intero arco del percorso formativo, non consente di valutare la stabilità a lungo termine delle trasformazioni concettuali e delle proiezioni didattiche osservate. Sarebbe di grande interesse condurre analisi di follow-up a distanza di tempo (es. dopo sei mesi o un anno) per verificare se e come le concettualizzazioni dell'IA e le pratiche didattiche dichiarate si modificano ulteriormente con l'effettiva integrazione nella pratica quotidiana, l'esperienza sul campo e l'evoluzione stessa della tecnologia.

Queste limitazioni aprono interessanti direzioni per ricerche future. Una prima direzione riguarda l'estensione dell'analisi a contesti culturali, istituzionali e disciplinari diversi, per indagare l'influenza di fattori contestuali sui processi di trasformazione concettuale e sull'adozione dell'AI literacy. Una seconda direzione concerne l'approfondimento specifico della dimensione etica, emersa come particolarmente significativa e complessa, attraverso indagini mirate sulle concezioni degli insegnanti riguardo a responsabilità, equità, trasparenza, autonomia e giustizia algoritmica nell'uso dell'IA in educazione. Una terza direzione, infine, riguarda l'esplorazione sistematica delle relazioni tra la trasformazione concettuale dichiarata e il cambiamento effettivo nelle pratiche didattiche, per comprendere meglio come le nuove concettualizzazioni si traducano (o non si traducano) in innovazioni concrete e sostenibili nelle aule scolastiche.

In una prospettiva più ampia, questa ricerca si inserisce nel crescente corpus di studi sull'integrazione dell'Intelligenza Artificiale in educazione, offrendo una prospettiva specifica e, per certi versi, innovativa sui processi di trasformazione concettuale che mediano tale integrazione dal punto di vista degli insegnanti. Comprendere in profondità come gli educatori concettualizzano, interpretano e si appropriano dell'IA è un passo essenziale per poter progettare percorsi formativi realmente efficaci e

per promuovere un'integrazione della tecnologia nei contesti educativi che sia non solo tecnicamente fattibile, ma soprattutto consapevole, critica, eticamente fondata e pedagogicamente significativa.

5.3.3 Triangolazione tra risultati quantitativi e qualitativi e implicazioni per l'AI literacy

L'adozione di un disegno di ricerca misto (Creswell & Plano Clark, 2011) richiede, nella fase conclusiva dell'analisi, un momento di integrazione sistematica delle evidenze provenienti dalle diverse componenti dello studio. Questa sezione si propone di triangolare i risultati dell'analisi quantitativa (questionario pre-post) con quelli dell'analisi qualitativa (diari metacognitivi), al fine di costruire una comprensione più ricca e sfaccettata dell'impatto del curriculum sull'AI literacy dei partecipanti.

La triangolazione delle evidenze quantitative e qualitative permette di identificare importanti punti di convergenza. Il dato quantitativo del massimo miglioramento nella dimensione conoscitiva ($d = 2.46$) trova pieno riscontro nell'analisi qualitativa, che ha evidenziato la trasformazione più radicale proprio nel passaggio dalle credenze iniziali verso una comprensione strutturata dei concetti fondamentali dell'IA. Sia i dati quantitativi che quelli qualitativi evidenziano come le dimensioni critica ed etica assumano un ruolo trasversale: quantitativamente la dimensione etica presentava i punteggi iniziali più elevati, mentre qualitativamente metodologie come la Philosophy for Children sono state citate trasversalmente in tutti i diari.

Il fatto che il 100% dei partecipanti abbia mostrato un cambiamento clinicamente significativo (RCI) trova corrispondenza qualitativa nella radicale dissimilarità lessicale osservata tra le sezioni dei diari. Dove i dati quantitativi documentano l'entità del cambiamento, i dati qualitativi ne illuminano i meccanismi: il miglioramento nella dimensione operativa si accompagna all'emergere di metafore dell'IA come "strumento logico" e "puzzle risolvibile"; la crescita nella dimensione critica corrisponde all'adozione di un linguaggio analitico specifico (bias, trasparenza, equità algoritmica).

La dimensione etica presenta la dimensione dell'effetto quantitativo più contenuta ($d = 1.31$), ma l'analisi qualitativa rivela che questa area è stata percepita come la più complessa dai partecipanti. Questa apparente divergenza si spiega considerando che: (a) i punteggi iniziali erano già elevati, limitando il margine di miglioramento; (b) la complessità intrinseca delle questioni etiche richiede tempi di elaborazione più lunghi; (c) l'etica dell'IA è un campo dove le risposte sono meno definite rispetto alle altre dimensioni.

In conclusione, la triangolazione dei risultati quantitativi e qualitativi rafforza significativamente la validità delle conclusioni dello studio, offrendo una comprensione multidimensionale dell'impatto del curriculum sull'AI literacy dei partecipanti. Le convergenze osservate testimoniano la robustezza dei risultati, mentre le complementarità illuminano i processi sottostanti ai cambiamenti documentati.

CAPITOLO 6 Elaborazione delle linee guida per l'AI literacy

6.1 Dalla ricerca alla pratica - verso un'AI literacy per la cittadinanza attiva

L'evoluzione dell'Intelligenza Artificiale da tecnologia specialistica a strumento pervasivo della vita quotidiana ha reso sempre più importante lo sviluppo di competenze diffuse che permettano a ogni cittadino di interagire consapevolmente con questi sistemi. Come evidenziato dalla sperimentazione condotta e dai risultati emersi dalla ricerca presentata nei capitoli precedenti, l'AI literacy non può più essere considerata una competenza riservata a esperti del settore o a professionisti dell'educazione, ma deve configurarsi come un elemento fondamentale della cittadinanza attiva nel XXI secolo.

Il presente capitolo si propone di tradurre in linee guida operative i risultati teorici e empirici emersi dalla ricerca, offrendo un quadro sistematico per l'implementazione dell'AI literacy rivolto alla cittadinanza nella sua interezza. Queste linee guida nascono dall'evidenza che il framework quadridimensionale sviluppato da Cuomo, Ranieri e Biagini (2022) e validato empiricamente nella sperimentazione descritta nei capitoli precedenti, rappresenta non solo un modello teorico robusto, ma anche uno strumento operativo efficace per promuovere una comprensione articolata e critica dell'Intelligenza Artificiale.

L'approccio adottato in questo capitolo riflette una concezione democratica e inclusiva dell'AI literacy, intesa come diritto-dovere del cittadino moderno. Come sottolineato da Freire (1970) nel suo approccio alla pedagogia critica, l'alfabetizzazione autentica non è mai un processo neutrale o meramente tecnico, ma un atto profondamente politico che ha il potere di trasformare sia gli individui che le società. Analogamente, l'AI literacy si configura come uno strumento di empowerment che consente ai cittadini di non essere semplicemente consumatori passivi di tecnologie intelligenti, ma di diventare attori consapevoli e critici nel plasmare il futuro digitale della società.

Le linee guida qui presentate si fondano su tre pilastri metodologici emersi dalla ricerca: la multidimensionalità del framework teorico, che garantisce una comprensione olistica dell'IA; l'efficacia delle strategie didattiche attive e partecipative sperimentate nel corso della ricerca; e l'importanza dell'adattabilità dei contenuti ai diversi contesti e target di popolazione. Questo approccio consente di delineare percorsi formativi che, pur mantenendo rigore teoretico e coerenza metodologica, possano essere flessibilmente adattati alle esigenze specifiche di diverse comunità, fasce d'età e contesti socioculturali.

È opportuno però premettere che le linee guida qui delineate hanno carattere esplorativo e orientativo: esse sono state elaborate a partire da una prima sperimentazione condotta in un contesto formativo specifico - quello dei futuri insegnanti di scuola primaria - e con un campione contenuto e relativamente omogeneo per caratteristiche sociodemografiche e formative. Come discusso nei paragrafi dedicati ai limiti della ricerca, la generalizzabilità delle evidenze empiriche che sostengono queste proposte è pertanto circoscritta, e la loro consolidazione richiederà ulteriori iterazioni su popolazioni più ampie e diversificate. Le indicazioni che seguono andranno dunque intese non come prescrizioni da applicare in modo uniforme, ma come un quadro di riferimento da adattare, rinegoziare e rivalidare nei diversi contesti di implementazione. In particolare, la traducibilità di queste linee guida in scenari scolastici e sociali reali richiederà di confrontarsi con vincoli

organizzativi (tempi curricolari, risorse umane, assetti formativi), culturali (rappresentazioni diffuse della tecnologia, resistenze o aspettative verso l'innovazione) e tecnologici (accesso alle infrastrutture, disponibilità di strumenti adeguati, livello di competenze dei docenti e degli operatori) che possono incidere sensibilmente sulla fattibilità e sull'efficacia degli interventi, come evidenziato dalla letteratura sull'implementazione dell'innovazione nei sistemi educativi. Il passaggio dal framework teorico alle pratiche d'aula richiede dunque un dialogo continuo tra la cornice proposta e la specificità dei contesti in cui viene applicata, e le formulazioni operative utilizzate nei paragrafi successivi vanno lette alla luce di questa premessa.

6.1.1 Principi fondamentali per l'implementazione dell'AI literacy nella società

L'implementazione efficace dell'AI literacy nella società richiede l'adozione di principi guida che orientino non solo i contenuti formativi, ma anche le modalità attraverso cui questi vengono trasmessi e appropriati. Tali principi, emersi dall'analisi dei risultati della sperimentazione e consolidati attraverso la riflessione teorica, costituiscono il fondamento su cui costruire ogni iniziativa di alfabetizzazione all'Intelligenza Artificiale.

Il primo principio fondamentale è quello dell'universalità inclusiva, che riconosce il diritto di ogni cittadino ad acquisire competenze di AI literacy indipendentemente dal background educativo, socioeconomico o tecnologico. Questo principio trova la sua legittimazione non solo in considerazioni di giustizia sociale, ma anche nell'evidenza empirica emersa dalla sperimentazione, che ha dimostrato come partecipanti con diversi livelli di familiarità tecnologica abbiano potuto beneficiare efficacemente del percorso formativo. L'universalità inclusiva non implica uniformità di approcci, ma piuttosto la necessità di sviluppare strategie differenziate che tengano conto delle specificità di ogni target, mantenendo però la coerenza con il framework quadridimensionale di riferimento.

Il secondo principio è quello della multiliteracy integrata, che riconosce l'AI literacy non come una competenza isolata, ma come parte integrante di un più ampio ecosistema di alfabetizzazioni necessarie nella società digitale. Come evidenziato dall'analisi interdisciplinare presentata nel capitolo 2, l'AI literacy si intreccia profondamente con la digital literacy, la media literacy, la data literacy e l'information literacy. L'implementazione efficace deve quindi evitare approcci compartimentalizzati, favorendo invece percorsi formativi che evidenzino le connessioni e le sinergie tra queste diverse competenze.

Il terzo principio è quello dell'apprendimento trasformativo, mutuato dalla teoria di Mezirow (1991) e validato dall'analisi dei diari riflessivi raccolti durante la sperimentazione. I risultati hanno mostrato come l'apprendimento significativo dell'AI literacy non avvenga attraverso un semplice accumulo di informazioni, ma attraverso processi di riflessione critica che portano alla messa in discussione e alla trasformazione delle proprie prospettive interpretative. Questo principio implica la necessità di progettare esperienze formative che non si limitino alla trasmissione di conoscenze, ma stimolino la riflessione metacognitiva e il questionamento delle proprie assunzioni implicite sull'Intelligenza Artificiale.

Il quarto principio è quello della pratica riflessiva situata, che riconosce l'importanza di ancorare l'apprendimento dell'AI literacy a esperienze concrete e contestualizzate. La sperimentazione ha evidenziato come le attività hands-on, dalla Philosophy for Children al coding unplugged, dalle

simulazioni con Teachable Machine al debate etico, abbiano facilitato una comprensione più profonda e duratura dei concetti. Questo principio non si limita al contesto educativo formale, ma si estende a tutti gli ambiti della vita sociale, richiedendo che l'AI literacy sia promossa attraverso esperienze significative e rilevanti per i diversi contesti di vita e di lavoro dei cittadini.

Il quinto principio è quello della responsabilità collettiva, che riconosce l'AI literacy non come una competenza puramente individuale, ma come una risorsa sociale che richiede la collaborazione di diverse istituzioni e attori sociali. I risultati della ricerca hanno mostrato come l'apprendimento collaborativo, attraverso discussioni di gruppo, progetti condivisi e confronti tra pari, abbia accelerato e arricchito il processo di acquisizione delle competenze. Questo principio implica la necessità di costruire reti di partnership tra istituzioni educative, organizzazioni della società civile, enti locali, imprese e altri stakeholder per creare un ecosistema formativo integrato e sostenibile.

Il sesto principio è quello dell'adattabilità evolutiva, che riconosce la natura dinamica dell'Intelligenza Artificiale e la necessità di sviluppare competenze che siano flessibili e aggiornabili nel tempo. L'analisi della dimensione conoscitiva nel framework ha evidenziato come la comprensione dell'evoluzione storica dell'IA aiuti a sviluppare una prospettiva critica sui suoi sviluppi futuri. Questo principio richiede che i percorsi di AI literacy non siano concepiti come interventi puntuali, ma come processi di apprendimento permanente che preparino i cittadini ad affrontare le trasformazioni future della tecnologia.

6.1.2 Linee guida per la dimensione conoscitiva: costruire le fondamenta concettuali

La dimensione conoscitiva dell'AI literacy, come dimostrato dai risultati della sperimentazione, costituisce il fondamento su cui si sviluppano tutte le altre competenze. L'analisi dei questionari pre-post e dei diari riflessivi ha evidenziato come una solida base concettuale non solo faciliti la comprensione degli aspetti operativi, critici ed etici dell'IA, ma contribuisca anche a demistificare la tecnologia, riducendo ansie e resistenze che spesso caratterizzano l'approccio iniziale dei cittadini a questi temi.

La prima linea guida per lo sviluppo della dimensione conoscitiva riguarda la costruzione di un linguaggio condiviso sull'IA. La ricerca ha mostrato come la proliferazione di definizioni spesso contraddittorie o parziali dell'Intelligenza Artificiale generi confusione e fraintendimenti. È quindi essenziale che ogni percorso di AI literacy parta da un'analisi critica delle diverse concettualizzazioni esistenti, non per imporre una definizione univoca, ma per sviluppare la capacità di navigare consapevolmente la molteplicità di prospettive. L'attività del "prisma delle definizioni" sperimentata nel curriculum ha dimostrato la sua efficacia nel favorire questa competenza meta-concettuale, stimolando i partecipanti a riflettere sulle diverse enfasi (razionalità versus umanità, pensiero versus azione) e sulle loro implicazioni.

La seconda linea guida concerne l'approccio storico-evolutivo alla comprensione dell'IA. I risultati hanno evidenziato come la conoscenza della storia dell'IA, dalle sue radici teoretiche agli sviluppi contemporanei, contribuisca significativamente allo sviluppo di una prospettiva critica e non deterministica sulla tecnologia. È fondamentale che i cittadini comprendano che l'IA non è un fenomeno monolitico emerso improvvisamente, ma il risultato di decenni di ricerca caratterizzata da successi, fallimenti, "primavere" ed "inverni". Questa prospettiva storica aiuta a sviluppare quello che

la letteratura definisce "technological wisdom" (Jonas, 1984), ovvero la capacità di valutare le innovazioni tecnologiche in una prospettiva temporale ampia e critica.

La terza linea guida riguarda la comprensione delle diverse tipologie e applicazioni dell'IA. È essenziale che i cittadini sviluppino una mappa concettuale delle diverse manifestazioni dell'IA, dalle applicazioni più familiari (assistenti vocali, sistemi di raccomandazione) a quelle più specialistiche (IA per la ricerca scientifica, sistemi di diagnosi medica). Questa comprensione deve andare oltre la semplice catalogazione, promuovendo la capacità di riconoscere i principi comuni sottostanti e le specificità di ciascun dominio applicativo. L'approccio adottato nella sperimentazione, che ha integrato esempi quotidiani con riflessioni sui principi generali, si è rivelato particolarmente efficace in questo senso.

La quarta linea guida concerne la riflessione sul confronto tra intelligenza umana e artificiale. I diari riflessivi hanno mostrato come questa riflessione non sia solo un esercizio filosofico, ma un elemento fondamentale per sviluppare una comprensione realistica delle capacità e dei limiti dell'IA. È importante che i cittadini comprendano che l'intelligenza artificiale, pur eccellendo in compiti specifici, opera secondo principi fondamentalmente diversi dall'intelligenza umana, mancando di attributi come la coscienza, l'esperienza soggettiva, la comprensione contestuale profonda e la creatività genuina. Questa comprensione è cruciale per evitare sia l'antropomorfizzazione eccessiva dell'IA sia la sottovalutazione delle sue potenzialità.

La quinta linea guida riguarda l'integrazione di prospettive interdisciplinari. La ricerca ha evidenziato come l'AI literacy benefici enormemente dall'integrazione di prospettive provenienti da diverse discipline. La dimensione conoscitiva non deve essere approcciata solo da un punto di vista tecnico-informatico, ma deve incorporare contributi dalla filosofia (questioni epistemologiche sull'intelligenza e la conoscenza), dalla psicologia cognitiva (modelli della mente e dell'apprendimento), dalla storia della scienza (evoluzione del pensiero computazionale), dalla sociologia (impatto sociale delle tecnologie). Questa interdisciplinarietà non solo arricchisce la comprensione, ma favorisce lo sviluppo di quella "saggezza tecnologica" necessaria per navigare la complessità dell'IA.

Per quanto riguarda l'adattamento ai diversi target, la dimensione conoscitiva può essere scalata efficacemente attraverso diversi livelli di complessità mantenendo la coerenza concettuale. Per i bambini della scuola dell'infanzia e primaria, l'approccio può focalizzarsi su metafore e analogie concrete, storie stimolo e giochi di ruolo che introducano i concetti di base in modo ludico e accessibile. Per gli adulti con diversi background educativi, è possibile modulare il livello di formalismo teorico mantenendo l'enfasi sulla rilevanza pratica e sociale dei concetti. L'esperienza della sperimentazione ha mostrato come anche concetti complessi possano essere resi accessibili attraverso un uso appropriato di esempi, visualizzazioni e connessioni con l'esperienza quotidiana.

6.1.3 Linee guida per la dimensione operativa: sviluppare competenze pratiche e pensiero computazionale

La dimensione operativa dell'AI literacy, come emerso dalla sperimentazione, rappresenta il ponte tra la comprensione teorica e l'applicazione pratica, permettendo ai cittadini di interagire efficacemente con i sistemi di IA e di comprenderne i meccanismi di funzionamento. L'analisi dei risultati ha evidenziato come le competenze operative non si limitino all'utilizzo strumentale delle tecnologie,

ma comprendano lo sviluppo del pensiero computazionale e della capacità di decodificare i processi algoritmici che sottendono le applicazioni di IA.

La prima linea guida per lo sviluppo della dimensione operativa concerne la promozione del pensiero computazionale come competenza trasversale. Wing (2006, 2008) ha identificato il pensiero computazionale come una competenza fondamentale per il XXI secolo, e la nostra ricerca ha confermato la sua centralità nell'AI literacy. Tuttavia, è importante che il pensiero computazionale non sia concepito semplicemente come una competenza tecnica, ma come un modo di approcciare e risolvere problemi che ha valore in molteplici contesti della vita quotidiana. L'esperienza del coding unplugged nella sperimentazione ha dimostrato come concetti come la scomposizione di problemi, il riconoscimento di pattern, l'astrazione e la progettazione algoritmica possano essere sviluppati anche al di fuori di contesti informatici tradizionali.

La seconda linea guida riguarda l'esperienza diretta con i meccanismi di machine learning. L'utilizzo di piattaforme educative come "AI per gli Oceani" e "Teachable Machine" durante la sperimentazione ha mostrato l'importanza di far sperimentare direttamente ai cittadini come l'IA "apprende" dai dati. Questa esperienza hands-on non solo demistifica il machine learning, ma sviluppa una comprensione intuitiva del ruolo cruciale dei dati nella determinazione delle prestazioni dei sistemi di IA. È essenziale che i cittadini comprendano attraverso l'esperienza diretta come la qualità, la quantità e la rappresentatività dei dati di addestramento influenzino profondamente i risultati dell'IA.

La terza linea guida concerne lo sviluppo di competenze di data literacy. L'analisi dei diari riflessivi ha evidenziato come molti partecipanti abbiano sviluppato una nuova consapevolezza del valore e della complessità dei dati attraverso le attività pratiche. I cittadini devono acquisire la capacità di comprendere come vengono raccolti, processati e utilizzati i dati, di riconoscere l'importanza della loro qualità e rappresentatività, e di sviluppare competenze di base per l'analisi e l'interpretazione di dataset semplici. Questa competenza è fondamentale non solo per comprendere il funzionamento dell'IA, ma anche per esercitare una cittadinanza critica in una società sempre più datificata.

La quarta linea guida riguarda la familiarizzazione con gli strumenti e le interfacce di IA. È importante che i cittadini sviluppino competenze pratiche nell'utilizzo di strumenti di IA comuni, dai motori di ricerca intelligenti agli assistenti virtuali, dagli strumenti di traduzione automatica alle piattaforme di IA generativa. Tuttavia, questa familiarizzazione deve andare oltre l'apprendimento procedurale, promuovendo la comprensione dei principi di funzionamento e delle limitazioni di questi strumenti. L'esperienza della sperimentazione con l'inquiry-based learning utilizzando ChatGPT ha mostrato come l'uso guidato e critico di questi strumenti possa diventare un potente veicolo di apprendimento.

La quinta linea guida concerne lo sviluppo di competenze collaborative uomo-IA. Come evidenziato da Long e Magerko (2020), l'AI literacy contemporanea richiede non solo la capacità di utilizzare l'IA, ma anche di collaborare efficacemente con essa. I cittadini devono imparare a sfruttare i punti di forza dell'IA complementandoli con le capacità unicamente umane. Questo include la capacità di formulare prompt efficaci, di interpretare criticamente gli output dell'IA, di verificare e contestualizzare le informazioni fornite, e di integrare l'intelligenza artificiale in processi decisionali più ampi che tengano conto di fattori che l'IA potrebbe non considerare.

Per quanto riguarda l'implementazione pratica di queste linee guida, l'esperienza della sperimentazione ha mostrato l'efficacia di approcci gradualisti e diversificati. Per i più giovani, attività

di coding unplugged come pixel art, body coding e costruzione di alberi decisionali si sono rivelate particolarmente efficaci. Per gli adulti, laboratori pratici con tool di machine learning accessibili e progetti di gruppo che richiedono l'analisi di semplici dataset hanno favorito lo sviluppo delle competenze operative. È fondamentale che queste attività siano sempre accompagnate da momenti di riflessione metacognitiva che aiutino i partecipanti a collegare l'esperienza pratica ai principi teorici sottostanti.

La dimensione operativa richiede anche la creazione di spazi di pratica sicuri dove i cittadini possano sperimentare senza timore di commettere errori. L'esperienza della sperimentazione ha mostrato come l'apprendimento per trial and error con le piattaforme educative di IA abbia non solo favorito la comprensione dei concetti, ma anche ridotto l'ansia tecnologica e aumentato la fiducia nell'utilizzo di questi strumenti. Questi spazi di pratica dovrebbero essere caratterizzati da un ambiente non giudicante, dalla possibilità di sperimentare liberamente e dalla presenza di facilitatori preparati che possano guidare l'apprendimento senza essere direttivi.

6.1.4 Linee guida per la dimensione critica: sviluppare pensiero analitico e capacità valutative

La dimensione critica dell'AI literacy si è rivelata, attraverso l'analisi dei risultati della sperimentazione, uno degli aspetti più trasformativi del percorso formativo. I diari riflessivi hanno evidenziato come lo sviluppo di competenze critiche non solo abbia migliorato la capacità dei partecipanti di valutare l'IA, ma abbia anche promosso una più generale attitudine al pensiero critico applicabile in diversi contesti. Questa dimensione è particolarmente cruciale in un'epoca caratterizzata dalla rapida diffusione di sistemi di IA sempre più sofisticati ma spesso opachi nei loro meccanismi decisionali.

La prima linea guida per lo sviluppo della dimensione critica concerne la comprensione e il riconoscimento dei bias algoritmici. La sperimentazione ha dimostrato l'efficacia dell'approccio experiential learning attraverso l'utilizzo di Teachable Machine per far sperimentare direttamente ai partecipanti come dataset sbilanciati possano portare a classificazioni distorte. I cittadini devono acquisire la capacità di identificare le diverse tipologie di bias (di selezione, di conferma, rappresentazionale) e di comprendere come questi possano originarsi sia dai dati di addestramento sia dalle scelte di progettazione degli algoritmi. Questa competenza è fondamentale per sviluppare una consapevolezza critica delle potenziali iniquità e discriminazioni che i sistemi di IA possono perpetuare o amplificare.

La seconda linea guida riguarda lo sviluppo di strategie per la valutazione critica degli output dell'IA. L'attività di analisi comparativa di testi e immagini, alcune generate da umani e altre dall'IA, implementata durante la sperimentazione, ha mostrato come sia possibile sviluppare competenze di discernimento anche in assenza di indicatori espliciti. I cittadini devono imparare a porsi domande critiche di fronte agli output dell'IA: qual è la fonte dell'informazione? Quanto è attendibile? Esistono perspective alternative? Quali potrebbero essere le limitazioni del sistema che ha prodotto questo output? Lo sviluppo di questi "critical thinking skills" richiede pratica costante e deve essere integrato nell'uso quotidiano delle tecnologie.

La terza linea guida concerne la comprensione delle problematiche di trasparenza e interpretabilità. Il problema della "black box" rappresenta una sfida fondamentale nell'era dell'IA deep learning, e i

cittadini devono sviluppare una comprensione delle implicazioni di questa opacità. Durante la sperimentazione, l'attività di confronto tra modelli trasparenti (come gli alberi decisionali semplici) e modelli opachi ha aiutato i partecipanti a comprendere il trade-off tra accuratezza e interpretabilità. I cittadini devono essere in grado di valutare quando la trasparenza è critica (ad esempio, in decisioni che riguardano la giustizia, la sanità o l'istruzione) e di richiedere spiegazioni appropriate quando necessario.

La quarta linea guida riguarda lo sviluppo di competenze di information literacy nell'era dell'IA generativa. La proliferazione di contenuti generati artificialmente richiede l'evoluzione delle tradizionali competenze di information literacy. I cittadini devono imparare non solo a verificare l'accuratezza delle informazioni, ma anche a identificare possibili "allucinazioni" dell'IA, a riconoscere contenuti sintetici e a utilizzare strumenti di fact-checking potenziati dall'IA. L'esperienza dell'inquiry-based learning con ChatGPT durante la sperimentazione ha mostrato come l'approccio critico all'uso di questi strumenti possa trasformarli da potenziali fonti di disinformazione in strumenti di apprendimento e verifica.

La quinta linea guida concerne l'analisi dell'impatto sociale dell'IA e delle sue implicazioni. I cittadini devono sviluppare la capacità di analizzare criticamente come l'IA stia trasformando diversi settori della società, dall'economia al lavoro, dalla politica all'istruzione. Questo include la comprensione dei meccanismi attraverso cui l'IA può creare o esacerbare disuguaglianze sociali, l'impatto sul mercato del lavoro e le implicazioni per la democrazia e la partecipazione civica. L'analisi dei diari riflessivi ha mostrato come la discussione di questi temi in gruppo favorisca lo sviluppo di una visione più articolata e sfumata delle trasformazioni in atto.

Per quanto riguarda l'implementazione metodologica della dimensione critica, la ricerca ha evidenziato l'efficacia di approcci basati sul Case-Based Learning e sul Problem-Based Learning. L'analisi di casi reali di errori o bias dell'IA, di fallimenti di sistemi automatici o di controversie algoritmiche si è rivelata particolarmente efficace nel stimolare la riflessione critica. È importante che questi casi siano selezionati per la loro rilevanza e accessibilità, e che siano presentati in modo da stimolare domande piuttosto che fornire risposte preconfezionate.

La dimensione critica richiede anche lo sviluppo di comunità di pratica critica, dove i cittadini possano condividere esperienze, dubbi e scoperte relative all'uso dell'IA. L'esperienza della sperimentazione ha mostrato come la discussione in gruppo e il confronto tra diverse prospettive abbiano arricchito significativamente l'apprendimento. Queste comunità possono assumere forme diverse, dalle associazioni civiche ai gruppi online, dai caffè scientifici ai laboratori urbani, e dovrebbero essere caratterizzate da un approccio inclusivo e non giudicante che favorisca la partecipazione di cittadini con diversi background e livelli di competenza.

6.1.5 Linee guida per la dimensione etica: promuovere responsabilità e consapevolezza valoriale

La dimensione etica dell'AI literacy ha assunto, nel corso della sperimentazione, un ruolo centrale non solo nella formazione di competenze specifiche, ma anche nello sviluppo di una più ampia consapevolezza civica e sociale. L'analisi dei diari riflessivi ha evidenziato come l'esplorazione delle questioni etiche legate all'IA abbia stimolato una riflessione profonda sui valori individuali e

collettivi, promuovendo quello che Kohlberg (1984) definisce come "sviluppo del ragionamento morale" in contesti tecnologici complessi.

La prima linea guida per lo sviluppo della dimensione etica concerne l'esplorazione dei principi etici fondamentali nel contesto dell'IA. La sperimentazione ha mostrato l'efficacia di un approccio che parte dall'identificazione e dalla discussione di principi etici universali (giustizia, autonomia, beneficenza, non maleficenza, trasparenza) per poi esplorare le loro applicazioni specifiche nel contesto dell'Intelligenza Artificiale. I cittadini devono sviluppare la capacità di riconoscere come questi principi possano entrare in tensione tra loro (ad esempio, il conflitto tra privacy e sicurezza, o tra efficienza e equità) e di navigare questi dilemmi in modo riflessivo e argomentato.

La seconda linea guida riguarda lo sviluppo di competenze per l'analisi di dilemmi etici complessi. L'attività di debate strutturato implementata durante la sperimentazione si è rivelata particolarmente efficace nel promuovere la capacità di considerare multiple prospettive e di argomentare posizioni etiche in modo rigoroso. I cittadini devono imparare a identificare gli stakeholder coinvolti in una questione etica, a riconoscere i valori in conflitto, a valutare le possibili conseguenze di diverse scelte e a formulare giudizi etici informati e giustificati. Questa competenza è fondamentale non solo per questioni legate all'IA, ma per la partecipazione democratica più in generale.

La terza linea guida concerne la comprensione delle implicazioni etiche della raccolta e dell'uso dei dati. In un'epoca caratterizzata dalla datificazione pervasiva, i cittadini devono sviluppare una profonda consapevolezza di come i loro dati vengano raccolti, processati e utilizzati dai sistemi di IA. Questo include la comprensione dei concetti di consenso informato, di profilazione algoritmica, di diritto all'oblio e di trasparenza nell'uso dei dati. L'analisi di casi come quello del "giocattolo intelligente" durante la sperimentazione ha evidenziato come scenari concreti possano facilitare la comprensione di questioni etiche complesse.

La quarta linea guida riguarda lo sviluppo di una coscienza critica rispetto alle implicazioni sociali dell'IA. I cittadini devono essere in grado di analizzare criticamente come l'IA stia ridefinendo concetti fondamentali come il lavoro, la creatività, le relazioni sociali e la partecipazione democratica. Questo include la capacità di riconoscere e interrogare le narrative dominanti sui benefici e i rischi dell'IA, di identificare chi trae vantaggio e chi può essere danneggiato dalle trasformazioni in atto, e di immaginare alternative più eque e sostenibili. L'esperienza della sperimentazione ha mostrato come l'approccio della pedagogia critica possa essere efficacemente applicato a questi temi.

La quinta linea guida concerne la promozione della responsabilità individuale e collettiva nell'era dell'IA. Come evidenziato dall'analisi dei diari riflessivi, uno degli outcomes più significativi della dimensione etica è stato lo sviluppo di un senso di agency e responsabilità nei confronti delle tecnologie. I cittadini devono comprendere che non sono semplicemente utilizzatori passivi dell'IA, ma hanno un ruolo attivo nel plasmare il suo sviluppo e la sua applicazione. Questo include la responsabilità nelle scelte individuali di utilizzo, nella partecipazione a dibattiti pubblici, nel supporto a politiche e regolamentazioni appropriate, e nell'educazione di altri cittadini.

Per quanto riguarda le strategie metodologiche per l'implementazione della dimensione etica, la ricerca ha evidenziato l'efficacia di approcci esperienziali e dialogici. Il debate strutturato, che ha coinvolto i partecipanti in discussioni su dilemmi etici reali, si è rivelato particolarmente efficace nel promuovere sia le competenze argomentative sia la capacità di considerare prospettive diverse.

L'utilizzo di piattaforme interattive come Moral Machine del MIT ha permesso di rendere tangibili dilemmi etici altrimenti astratti, facilitando la riflessione e la discussione.

La dimensione etica richiede anche la creazione di spazi di dialogo etico nella società, dove cittadini con diversi background e prospettive possano confrontarsi su questioni legate all'IA in modo costruttivo e rispettoso. Questi spazi possono assumere forme diverse: caffè etici, forum civici, panel cittadini, piattaforme online moderate. È fondamentale che questi spazi siano caratterizzati da principi di inclusività, pluralismo e rispetto per la diversità di opinioni, e che siano facilitati da persone con competenze sia in ambito tecnologico sia in ambito etico-filosofico.

Un aspetto particolarmente importante emerso dalla sperimentazione è la connessione tra etica dell'IA e questioni ambientali. I partecipanti hanno mostrato un crescente interesse per l'impatto ambientale dell'IA, dalle emissioni di carbonio legate al training di modelli complessi alle implicazioni ecologiche della società digitale. Questa connessione evidenzia come l'etica dell'IA non possa essere considerata isolatamente, ma debba essere integrata in una più ampia riflessione sulla sostenibilità e sulla giustizia globale.

6.2 Strategie trasversali e approcci metodologici per l'AI literacy

L'analisi dei risultati della sperimentazione ha evidenziato come l'efficacia dell'AI literacy non dipenda solo dai contenuti specifici delle singole dimensioni, ma soprattutto dall'adozione di strategie metodologiche trasversali che favoriscano l'integrazione delle competenze e la loro applicazione in contesti reali. Queste strategie, validate empiricamente nel corso della ricerca, rappresentano un patrimonio metodologico fondamentale per l'implementazione dell'AI literacy a livello sociale.

La prima strategia trasversale è quella dell'apprendimento attraverso la riflessione metacognitiva. L'utilizzo sistematico dei diari riflessivi durante la sperimentazione ha dimostrato come la riflessione strutturata sui propri processi di apprendimento non solo faciliti la consolidazione delle conoscenze, ma promuova anche lo sviluppo di una maggiore consapevolezza delle proprie credenze e assunzioni sull'IA. Questa strategia può essere implementata in diversi contesti attraverso portfolii digitali, blog di riflessione, sessioni di peer reflection o journal guidati. È fondamentale che queste attività siano strutturate ma non prescrittive, lasciando spazio all'esplorazione personale mentre offrendo framework per orientare la riflessione.

La seconda strategia concerne l'integrazione di metodologie attive e partecipative. La sperimentazione ha confermato l'efficacia di approcci come la Philosophy for Children, il debate strutturato, il problem-based learning e l'inquiry-based learning nell'promuovere un apprendimento profondo e trasformativo dell'AI literacy. Queste metodologie condividono alcuni principi fondamentali: il coinvolgimento attivo dei partecipanti, l'enfasi sulla costruzione collaborativa di conoscenze, il questionamento delle assunzioni implicite e la connessione tra apprendimento e azione. L'implementazione di queste metodologie richiede la formazione di facilitatori competenti e la creazione di ambienti di apprendimento che favoriscano la partecipazione e il dialogo.

La terza strategia riguarda l'approccio narrativo e la pedagogia delle storie. L'analisi dei diari ha evidenziato come l'uso di storie, metafore e narrazioni abbia facilitato la comprensione di concetti complessi e astratti. Le storie non solo rendono i contenuti più accessibili e memorabili, ma favoriscono anche lo sviluppo di competenze di perspective-taking e di comprensione contestuale. Questa strategia può essere implementata attraverso l'uso di racconti di fantascienza, casi studio

narrativi, biographical learning (storie di pionieri dell'IA), fiction interattive e storytelling collaborativo con e sull'IA.

La quarta strategia concerne l'apprendimento esperienziale e l'uso strategico della tecnologia. La sperimentazione ha mostrato l'importanza di far sperimentare direttamente ai partecipanti le tecnologie di IA, ma sempre in modo guidato e riflessivo. Le platform educative come Teachable Machine, le attività di coding unplugged, l'interazione critica con chatbot e la creazione collaborativa di contenuti con IA generativa si sono rivelate strumenti potenti per demistificare la tecnologia e sviluppare competenze operative. Tuttavia, è fondamentale che l'uso della tecnologia sia sempre accompagnato da momenti di riflessione critica che aiutino a collegare l'esperienza pratica ai principi teorici e alle implicazioni etiche.

La quinta strategia riguarda la promozione di una cultura dell'investigazione e del questionamento. L'inquiry-based learning implementato durante la sperimentazione ha dimostrato come l'educazione all'AI literacy debba favorire lo sviluppo di un atteggiamento investigativo e critico piuttosto che la trasmissione di conoscenze preconfezionate. I cittadini devono essere incoraggiati a porsi domande, a mettere in discussione le affermazioni sull'IA, a cercare evidenze, a considerare fonti multiple e a formare opinioni informate. Questa cultura dell'investigazione è particolarmente importante in un campo in rapida evoluzione come l'IA, dove le conoscenze diventano rapidamente obsolete e dove la capacità di apprendere autonomamente è più importante della memorizzazione di fatti specifici.

La sesta strategia concerne l'approccio interdisciplinare e la contestualizzazione sociale. L'AI literacy non può essere sviluppata in isolamento, ma deve essere integrata con altre literacy e collegata a questioni sociali, economiche, politiche e culturali più ampie. La sperimentazione ha mostrato l'efficacia di approcci che collegano l'IA a temi come la sostenibilità ambientale, la giustizia sociale, la democrazia partecipativa e l'innovazione culturale. Questa contestualizzazione non solo rende l'apprendimento più rilevante e significativo, ma favorisce anche lo sviluppo di una visione sistemica delle trasformazioni in atto.

Per quanto riguarda l'adattamento di queste strategie a diversi target e contesti, l'esperienza della sperimentazione suggerisce alcuni principi guida. Per i bambini, l'approccio deve privilegiare la concretezza, il gioco e la manipolazione, utilizzando metafore semplici e storie accessibili. Per gli adolescenti, può essere efficace l'uso di social media, gaming e challenge collaborative. Per gli adulti, l'enfasi dovrebbe essere sulla rilevanza pratica, sulla connessione con l'esperienza professionale e personale, e sulla possibilità di applicare immediatamente le competenze apprese. Per tutti i target, è fondamentale rispettare i diversi stili di apprendimento, offrire multiple modalità di partecipazione e creare un ambiente di apprendimento inclusivo e non giudicante.

6.2.1 Implementazione in diversi contesti e per diverse popolazioni

L'universalizzazione dell'AI literacy richiede un approccio differenziato che tenga conto della diversità di contesti sociali, culturali e educativi in cui questa competenza deve essere sviluppata. L'analisi dei risultati della sperimentazione, pur essendo stata condotta in un contesto specifico (futuri insegnanti di scuola primaria), offre importanti indicazioni su come il framework quadridimensionale possa essere adattato e implementato in una varietà di contesti sociali per raggiungere l'obiettivo di una cittadinanza digitale consapevole e critica. Come precedentemente accennato, tale trasferibilità, tuttavia, non può essere assunta in modo automatico, ma deve essere valutata alla luce delle

condizioni concrete dei contesti scolastici reali. Vincoli organizzativi, come la disponibilità di tempo curricolare e la formazione dei docenti, vincoli tecnologici, come l'accesso a dispositivi, connettività e strumenti digitali adeguati, e vincoli culturali, legati agli atteggiamenti della comunità scolastica verso l'IA, possono incidere significativamente sulle modalità di implementazione delle linee guida. Per questo motivo, le indicazioni proposte devono essere intese come orientamenti flessibili, da adattare attraverso processi di contestualizzazione locale, sperimentazione progressiva e valutazione continua.

Nel sistema educativo formale, l'implementazione dell'AI literacy deve essere concepita come un processo integrato e progressivo che accompagni gli studenti dalla scuola dell'infanzia all'università. Per la scuola dell'infanzia, l'approccio deve essere prevalentemente ludico ed esplorativo, utilizzando storie, giochi di ruolo e attività manipolative per introdurre i concetti di base. Ad esempio, il concetto di algoritmo può essere introdotto attraverso ricette di cucina o istruzioni per giochi, mentre la distinzione tra intelligenza umana e artificiale può essere esplorata attraverso discussioni sui personaggi preferiti dei cartoni animati.

Nella scuola primaria, come dimostrato dalle attività progettate dai partecipanti alla sperimentazione, è possibile introdurre gradualmente tutti e quattro le dimensioni del framework attraverso metodologie appropriate all'età. Il coding unplugged, la Philosophy for Children, i giochi di classificazione e le discussioni guidate su temi etici semplici si sono rivelati strumenti efficaci. È importante che l'AI literacy sia integrata trasversalmente nelle diverse discipline piuttosto che essere confinata a ore dedicate di "educazione digitale".

Nella scuola secondaria, l'approccio può diventare più analitico e critico, incorporando progetti interdisciplinari, debate strutturati, analisi di casi studio complessi e sperimentazione diretta con strumenti di IA. Gli studenti possono essere coinvolti in ricerche sui bias algoritmici, nell'analisi critica di notizie generate dall'IA, nella progettazione di semplici sistemi di machine learning per risolvere problemi reali.

Nell'educazione superiore, l'AI literacy deve essere integrata in tutti i corsi di studio, non solo in quelli tecnici. Ogni disciplina deve esplorare come l'IA stia trasformando il proprio campo, le opportunità e i rischi specifici, e le competenze necessarie per i futuri professionisti. La formazione degli educatori, come dimostrato dalla presente ricerca, rappresenta un nodo cruciale per la diffusione dell'AI literacy nel sistema educativo.

Le biblioteche pubbliche rappresentano spazi ideali per l'implementazione di programmi di AI literacy per la cittadinanza. Possono ospitare cicli di conferenze divulgative, laboratori pratici per diverse fasce d'età, gruppi di lettura e discussione su temi legati all'IA, e spazi maker dove sperimentare con tecnologie accessibili. L'esperienza della sperimentazione suggerisce che programmi di peer education, dove cittadini più esperti guidano altri in percorsi di apprendimento, possono essere particolarmente efficaci.

I centri culturali e le case del popolo possono promuovere dibattiti pubblici sui temi etici dell'IA, cineforum su film che trattano questi temi, e workshop pratici per diverse comunità. È importante che questi spazi mantengano un approccio inclusivo e attento alle diverse competenze digitali di partenza dei partecipanti.

Le organizzazioni della società civile possono integrare l'AI literacy nelle loro attività di advocacy e educazione alla cittadinanza, collegando i temi dell'IA alle loro aree di intervento specifiche (diritti umani, ambiente, lavoro, genere). Questo approccio contestualizzato può rendere l'AI literacy più rilevante e motivante per diversi gruppi di cittadini.

Diversi settori professionali richiedono adattamenti specifici dell'AI literacy. Nel settore sanitario, l'enfasi deve essere sulla comprensione dei sistemi di supporto alle decisioni cliniche, sui temi di privacy e consenso, e sulle questioni etiche legate all'automazione in medicina. Nel settore legale, focus sulla comprensione degli algoritmi decisionali, sui bias nel sistema giudiziario e sulla responsabilità legale degli errori dell'IA.

Per il settore privato, programmi di formazione aziendale possono integrare l'AI literacy nella formazione continua dei dipendenti, adattando i contenuti alle specifiche esigenze del settore e del ruolo professionale. È importante che questi programmi non si limitino agli aspetti tecnici, ma includano anche riflessioni etiche sul utilizzo responsabile dell'IA nell'ambiente lavorativo.

L'implementazione dell'AI literacy deve prestare particolare attenzione alle popolazioni vulnerabili o a rischio di esclusione digitale. Per gli anziani, programmi specificatamente progettati possono utilizzare un linguaggio non tecnico, riferimenti culturali appropriati, e un ritmo più lento nell'introduzione dei concetti. L'enfasi dovrebbe essere sulla rilevanza per la vita quotidiana e sulla sicurezza nell'uso delle tecnologie.

Per le persone con disabilità, l'AI literacy deve essere resa accessibile attraverso tecnologie assistive appropriate e metodologie inclusive. È importante esplorare anche come l'IA possa migliorare l'accessibilità e l'autonomia, ma senza trascurare i rischi di esclusione o discriminazione algoritmica.

Per le comunità marginali o con barriere linguistiche, programmi culturalmente sensibili possono utilizzare mediatori culturali, materiali tradotti, e metodologie che valorizzino le conoscenze e le esperienze locali. È fondamentale che l'AI literacy non diventi uno strumento di ulteriore marginalizzazione, ma un'opportunità di empowerment.

La ricerca ha evidenziato come la flessibilità nell'adattamento dei contenuti e delle metodologie sia cruciale per raggiungere efficacemente diverse popolazioni. Tuttavia, è importante mantenere la coerenza con il framework quadridimensionale per assicurare che tutti i cittadini, indipendentemente dal contesto di apprendimento, sviluppino competenze complete e integrate di AI literacy.

6.2.2 Valutazione e monitoraggio dell'AI literacy

Lo sviluppo di sistemi efficaci di valutazione e monitoraggio dell'AI literacy rappresenta una sfida metodologica complessa, che richiede l'adozione di approcci multidimensionali e strumenti diversificati in grado di catturare la ricchezza e la complessità di questa competenza emergente. L'esperienza maturata durante la sperimentazione, attraverso l'uso combinato di questionari strutturati e diari riflessivi, offre importanti indicazioni per lo sviluppo di framework valutativi che possano essere adattati a diversi contesti e popolazioni.

Il primo principio fondamentale per la valutazione dell'AI literacy è quello della multidimensionalità. Come dimostrato dalla validazione del framework quadridimensionale durante la sperimentazione, l'AI literacy non può essere ridotta a una singola competenza o conoscenza, ma deve essere valutata nelle sue quattro dimensioni interconnesse. Questo richiede lo sviluppo di strumenti che possano

misurare separatamente ciascuna dimensione mentre allo stesso tempo catturano le sinergie e le connessioni tra di esse.

Il secondo principio è quello della valutazione autentica, che privilegia la misurazione delle competenze in contesti significativi e rilevanti per i valutati. L'analisi dei diari riflessivi ha mostrato come le competenze di AI literacy si manifestino più chiaramente quando i soggetti sono chiamati ad applicarle a problemi reali o a riflettere su situazioni concrete. Questo suggerisce la necessità di integrare forme di assessment tradizionali con valutazioni basate su performance, progetti, e riflessioni guidate.

Il terzo principio è quello della progressività e adattabilità, che riconosce come l'AI literacy si sviluppi in modo graduale e non lineare, e come le competenze debbano essere valutate in relazione al livello di sviluppo e al contesto di riferimento. Questo implica la necessità di sviluppare scale di valutazione progressive che possano essere adattate a diverse età, background educativi e contesti di apprendimento.

La sperimentazione ha validato l'efficacia di questionari strutturati basati sul framework quadridimensionale per la valutazione quantitativa dell'AI literacy. Il questionario sviluppato e utilizzato nella ricerca (Biagini et al., 2023) rappresenta un modello replicabile che può essere adattato a diverse popolazioni. Tuttavia, è importante che questi strumenti siano continuamente aggiornati per riflettere l'evoluzione delle tecnologie e delle competenze richieste.

I diari riflessivi strutturati si sono rivelati strumenti preziosi per la valutazione qualitativa, permettendo di catturare non solo il livello delle competenze, ma anche i processi di apprendimento e trasformazione concettuale. L'analisi di questi diari attraverso metodologie miste (qualitative e NLP) ha mostrato come sia possibile estrarre indicatori significativi sui pattern di apprendimento e sulle aree di sviluppo più critiche.

Le performance assessment basate su compiti autentici rappresentano un complemento importante agli strumenti tradizionali. Esempi includono: progetti di analisi critica di sistemi di IA reali, progettazione di semplici algoritmi per risolvere problemi specifici, sviluppo di proposte etiche per l'uso dell'IA in contesti particolari, creazione di materiali educativi sull'IA per target specifici.

I portfolio digitali possono fornire una documentazione longitudinale dello sviluppo dell'AI literacy, raccogliendo evidenze di apprendimento da diverse fonti e contesti. Questi portfolio possono includere riflessioni personali, progetti realizzati, contributi a discussioni online, e materiali prodotti durante percorsi formativi.

L'analisi dei risultati della sperimentazione suggerisce alcuni indicatori chiave per il monitoraggio dello sviluppo dell'AI literacy:

Per la dimensione conoscitiva: accuratezza nel riconoscere diverse tipologie di IA, capacità di spiegare concetti fondamentali con linguaggio appropriato, evoluzione delle metafore utilizzate per descrivere l'IA (da antropomorfe a più tecnicamente accurate), capacità di contestualizzare storicamente gli sviluppi dell'IA.

Per la dimensione operativa: competenza nell'uso di strumenti di IA comuni, capacità di identificare e risolvere problemi algoritmici semplici, sviluppo del pensiero computazionale in contesti non digitali, capacità di collaborazione efficace con sistemi di IA.

Per la dimensione critica: abilità nel riconoscere bias e limitazioni dei sistemi di IA, competenza nella valutazione critica di output generati dall'IA, capacità di porre domande appropriate sull'affidabilità e la trasparenza dei sistemi, sviluppo di strategie per la verifica delle informazioni generate dall'IA.

Per la dimensione etica: consapevolezza delle implicazioni etiche dell'IA, capacità di argomentare posizioni etiche in dibattiti strutturati, sviluppo di una prospettiva critica sull'impatto sociale dell'IA, comportamenti responsabili nell'uso quotidiano delle tecnologie.

Oltre alla valutazione individuale, è necessario sviluppare sistemi di monitoraggio che possano fornire indicazioni sullo stato dell'AI literacy a livello di popolazione. Questo può includere:

Survey periodiche su campioni rappresentativi della popolazione per monitorare l'evoluzione delle competenze e delle percezioni relative all'IA. Analisi dei comportamenti digitali attraverso dati aggregati e anonimizzati sull'uso di tecnologie AI. Valutazione dell'efficacia di programmi formativi attraverso studi longitudinali e comparativi. Monitoraggio della qualità e della diffusione di iniziative educative sull'AI literacy in diversi contesti.

6.2.3 Sfide e considerazioni etiche nella valutazione

La valutazione dell'AI literacy solleva importanti questioni etiche, particolarmente legate alla privacy, al consenso e all'uso dei dati raccolti. È fondamentale che i sistemi di valutazione rispettino i diritti dei soggetti valutati e utilizzino i dati raccolti esclusivamente per finalità educative e di ricerca.

Inoltre, esiste il rischio che la valutazione dell'AI literacy possa diventare un nuovo fattore di esclusione sociale se non progettata e implementata in modo inclusivo e equo. È importante che gli strumenti di valutazione siano culturalmente sensibili, accessibili a persone con disabilità, e non penalizzino coloro che hanno avuto minori opportunità di accesso alle tecnologie.

L'esperienza della sperimentazione suggerisce che un approccio efficace alla valutazione dell'AI literacy debba combinare rigore metodologico e sensibilità contestuale, utilizzando multiple fonti di evidenza e strumenti diversificati per catturare la complessità di questa competenza emergente. Solo attraverso sistemi di valutazione robusti e inclusivi sarà possibile monitorare efficacemente i progressi verso l'obiettivo di un'AI literacy universale e promuovere lo sviluppo di competenze che preparino tutti i cittadini a navigare consapevolmente nell'era dell'Intelligenza Artificiale.

L'implementazione dell'AI literacy a livello di società presenta sfide significative che richiedono approcci innovativi e una visione sistemica delle trasformazioni necessarie. L'analisi dei risultati della sperimentazione, integrata con le riflessioni teoriche e le evidenze dalla letteratura internazionale, permette di identificare i principali ostacoli da superare e di delineare prospettive promettenti per lo sviluppo futuro di questa competenza essenziale.

La prima sfida riguarda la rapida evoluzione tecnologica che caratterizza il campo dell'Intelligenza Artificiale. Come emerso dall'analisi delle trasformazioni concettuali documentate nei diari riflessivi, l'AI literacy deve essere concepita come una competenza dinamica e adattabile piuttosto che come un insieme statico di conoscenze. Questo richiede lo sviluppo di framework educativi che privilegino principi fondamentali e competenze trasferibili rispetto a conoscenze specifiche destinate a diventare rapidamente obsolete.

La seconda sfida è rappresentata dalle disuguaglianze nell'accesso alle tecnologie e alle opportunità formative. La ricerca ha evidenziato come background socioeconomici diversi possano influenzare le modalità di approccio all'IA. Per realizzare una vera AI literacy universale, è necessario affrontare sistematicamente le disuguaglianze digitali esistenti, garantendo non solo l'accesso alle tecnologie ma anche la disponibilità di percorsi formativi di qualità per tutti i cittadini.

La terza sfida concerne la formazione dei formatori. Come dimostrato dalla sperimentazione con futuri insegnanti, la qualità dell'AI literacy dipende crucialmente dalla preparazione di coloro che la insegnano. Questo richiede investimenti significativi nella formazione iniziale e continua di educatori, bibliotecari, operatori sociali e altri professionisti che possono contribuire alla diffusione di queste competenze.

La quarta sfida riguarda la resistenza culturale e psicologica al cambiamento. L'analisi delle metafore iniziali raccolte attraverso i diari ha mostrato come molti partecipanti avessero inizialmente rappresentazioni ansiogene o mistificate dell'IA. Superare queste resistenze richiede approcci pedagogici empatici e incrementali che rispettino i timori legittimi pur promuovendo un approccio critico e informato.

Una sfida significativa è rappresentata dalla complessità intrinseca dei sistemi di IA contemporanei, che rende difficile sviluppare spiegazioni accessibili senza sacrificare il rigore concettuale. L'esperienza del coding unplugged e dell'uso di piattaforme educative durante la sperimentazione ha mostrato l'importanza di trovare metafore e analogie efficaci, ma questo rimane un equilibrio delicato da mantenere.

Un'altra sfida metodologica riguarda la valutazione di competenze complesse e multidimensionali come l'AI literacy. Come discusso nella sezione precedente, sviluppare strumenti di assessment che siano al contempo rigorosi, autentici e praticabili su larga scala rappresenta una sfida metodologica ancora da affrontare pienamente.

La personalizzazione dei percorsi formativi per rispondere alla diversità di bisogni, stili di apprendimento e contesti rappresenta un'ulteriore sfida. La ricerca ha mostrato l'efficacia di approcci flessibili e adattabili, ma la loro implementazione sistematica richiede risorse e competenze non sempre disponibili.

6.2.4 Prospettive promettenti e innovazioni emergenti

Nonostante le sfide, emergono prospettive promettenti per lo sviluppo dell'AI literacy universale. L'utilizzo pedagogico dell'IA stessa per personalizzare e migliorare l'apprendimento rappresenta una frontiera innovativa. Sistemi di tutoring intelligenti, piattaforme adattive per l'autoapprendimento, e strumenti di IA per la creazione di contenuti educativi possono contribuire a scalare l'AI literacy mantenendo qualità e personalizzazione.

Lo sviluppo di comunità di pratica online e reti peer-to-peer per l'AI literacy offre opportunità per l'apprendimento collaborativo e continuo. L'esperienza della sperimentazione ha mostrato l'importanza dell'apprendimento sociale e della condivisione di esperienze, principi che possono essere amplificati attraverso piattaforme digitali appropriate.

L'integrazione dell'AI literacy in iniziative di cittadinanza digitale più ampie rappresenta un'opportunità per collegare queste competenze a questioni di democrazia, partecipazione civica e

giustizia sociale. Questo approccio può aumentare la rilevanza e la motivazione all'apprendimento collegando l'IA a preoccupazioni concrete dei cittadini.

Lo sviluppo di partnership multi-stakeholder tra istituzioni educative, organizzazioni della società civile, settore privato e istituzioni pubbliche può creare ecosistemi formativi più ricchi e sostenibili. L'esperienza della sperimentazione universitaria può essere scalata attraverso collaborazioni strategiche che mettano insieme diverse competenze e risorse.

Diverse direzioni di ricerca emergono come prioritarie per il futuro sviluppo dell'AI literacy. La ricerca longitudinale sui processi di sviluppo dell'AI literacy può fornire evidenze cruciali sull'efficacia a lungo termine di diversi approcci formativi e sui fattori che influenzano la persistenza delle competenze acquisite.

Lo sviluppo di metodologie di assessment innovative che utilizzino tecnologie emergenti (come l'analisi automatica di testi, la gamification valutativa, o la realtà virtuale per la simulazione di scenari) può migliorare significativamente la nostra capacità di monitorare e supportare lo sviluppo dell'AI literacy.

La ricerca cross-culturale sull'AI literacy può evidenziare come diverse culture e contesti socioeconomici influenzino la comprensione e l'appropriazione dell'IA, informando lo sviluppo di approcci più inclusivi e culturalmente sensibili.

L'analisi dell'impatto dell'AI literacy sui comportamenti reali dei cittadini, sulle loro scelte lavorative, sulla partecipazione democratica e sul benessere sociale può fornire evidenze cruciali sul valore sociale di questi investimenti formativi.

Guardando al futuro, l'obiettivo di un'AI literacy universale richiede una trasformazione sistematica dell'approccio all'educazione e alla formazione continua. Questo include l'integrazione dell'AI literacy nei curricula educativi a tutti i livelli, lo sviluppo di infrastrutture formative distribuite e accessibili, la creazione di incentivi per la formazione continua, e l'evoluzione di una cultura sociale che valorizzi il pensiero critico e la competenza tecnologica.

La ricerca presentata in questa tesi offre un contributo importante a questa visione, fornendo un framework teorico robusto, evidenze empiriche sulla sua efficacia, e linee guida operative per l'implementazione. Tuttavia, realizzare l'obiettivo di cittadini competenti e critici nell'era dell'IA richiederà sforzi sostenuti e coordinati da parte di tutta la società.

L'AI literacy non è solo una competenza tecnica, ma una componente essenziale della cittadinanza democratica nel XXI secolo. Come tale, il suo sviluppo rappresenta non solo una necessità educativa, ma un investimento nel futuro democratico e sostenibile delle nostre società. Le sfide sono significative, ma le prospettive aperte dalla ricerca e dall'innovazione pedagogica offrono motivi di ottimismo ragionato per il raggiungimento di questo obiettivo cruciale.

6.2.5 Considerazioni: verso un futuro consapevole nell'era dell'Intelligenza Artificiale

L'elaborazione delle linee guida per l'AI literacy presentate in questo capitolo rappresenta la sintesi operativa di un percorso di ricerca che ha attraversato dimensioni teoriche, metodologiche ed empiriche per delineare un framework comprensivo e applicabile di competenze essenziali per la cittadinanza nell'era dell'Intelligenza Artificiale. Tale percorso, dallo sviluppo del framework

quadridimensionale alla sua validazione attraverso la sperimentazione, dall'analisi critica delle metodologie didattiche alla proposta di strategie di implementazione sistematica, testimonia la possibilità concreta di trasformare la ricerca educativa in strumenti operativi per la trasformazione sociale.

Le linee guida articolate lungo le quattro dimensioni del framework – conoscitiva, operativa, critica ed etica – non costituiscono semplicemente un catalogo di competenze da acquisire, ma delineano un percorso di empowerment che trasforma i cittadini da consumatori passivi di tecnologie intelligenti in attori consapevoli e critici nel plasmare il futuro digitale della società. Questo approccio riflette una concezione democratica dell'educazione tecnologica che affonda le sue radici nella pedagogia critica e nel costruttivismo sociale, riconoscendo nell'apprendimento un processo collettivo di costruzione di significati e trasformazione della realtà.

L'efficacia delle strategie metodologiche sperimentate – dalla Philosophy for Children al debate strutturato, dal coding unplugged all'inquiry-based learning con IA generativa – conferma l'importanza di approcci attivi e partecipativi che coinvolgano pienamente i soggetti nel processo di apprendimento. Queste metodologie non solo facilitano l'acquisizione di conoscenze e competenze specifiche, ma promuovono lo sviluppo di quelle meta-competenze (pensiero critico, capacità argomentativa, consapevolezza metacognitiva) che rappresentano il vero valore aggiunto dell'AI literacy.

La prospettiva universalista che informa queste linee guida riconosce l'AI literacy non come privilegio di élite tecnologiche o educative, ma come diritto fondamentale di ogni cittadino. Questa visione comporta la responsabilità collettiva di sviluppare sistemi formativi inclusivi, accessibili e culturalmente sensibili che possano raggiungere efficacemente tutti i segmenti della popolazione, dai bambini dell'infanzia agli anziani, dalle comunità marginalizzate ai professionisti altamente qualificati.

Le sfide identificate – dalla rapida evoluzione tecnologica alle disuguaglianze nell'accesso, dalla formazione dei formatori alle resistenze culturali – non sono ostacoli insormontabili ma piuttosto aree di intervento prioritarie che richiedono approcci innovativi e investimenti sostenuti. Le prospettive emergenti, dall'uso pedagogico dell'IA stessa allo sviluppo di comunità di pratica online, dalle partnership multi-stakeholder alla ricerca interdisciplinare, offrono ragioni fondate di ottimismo per il raggiungimento dell'obiettivo di un'AI literacy universale.

L'importanza sociale di questo obiettivo trascende la dimensione puramente educativa per toccare questioni fondamentali di democrazia, giustizia e sostenibilità. In un'epoca in cui l'Intelligenza Artificiale influenza sempre più profondamente decisioni che riguardano ogni aspetto della vita sociale – dall'accesso ai servizi all'amministrazione della giustizia, dal mercato del lavoro alle relazioni interpersonali – la diffusione di competenze critiche per comprendere, valutare e orientare queste tecnologie diventa un requisito essenziale per il mantenimento di società democratiche e eque.

Le linee guida qui presentate costituiscono un primo passo verso questa visione, fornendo un quadro concettuale robusto e strumenti operativi validati empiricamente. Tuttavia, la loro implementazione richiederà adattamenti continui, sperimentazioni locali, valutazioni sistematiche e aggiustamenti basati sull'evidenza emergente. In questo senso, le linee guida non devono essere intese come

prescrizioni rigide ma come framework flessibili che possano orientare l'azione educativa pur permettendo l'innovazione e l'adattamento alle specificità contestuali.

Il futuro dell'AI literacy dipenderà dalla capacità di costruire coalizioni ampie che coinvolgano educatori, policymaker, tecnologi, organizzazioni della società civile e cittadini comuni in un progetto condiviso di alfabetizzazione critica. Solo attraverso questo sforzo collettivo sarà possibile realizzare la visione di una società in cui ogni cittadino possieda le competenze necessarie per navigare consapevolmente nell'era dell'Intelligenza Artificiale, contribuendo a orientarne lo sviluppo verso obiettivi di benessere comune, giustizia sociale e sostenibilità ambientale.

In conclusione, l'AI literacy rappresenta molto più di una semplice competenza tecnologica: è uno strumento di empowerment democratico, un veicolo di giustizia sociale e un investimento nel futuro dell'umanità. Le linee guida elaborate in questo capitolo offrono una roadmap per iniziare questo viaggio trasformativo, ma il successo del progetto dipenderà dalla volontà collettiva di abbracciare la sfida e impegnarsi per la sua realizzazione. L'era dell'Intelligenza Artificiale è già iniziata: spetta a noi decidere se subirla passivamente o modellarla consciamente per il bene di tutti.

Bibliografia

- Aldrich, J. H. (1995). Correlations genuine and spurious in Pearson and Yule. *Statistical Science*, 10(4), 364-376.
- Allison, P. D. (2001). *Missing data*. Sage Publications.
- Altman, D. G. (1991). *Practical statistics for medical research*. London: Chapman and Hall.
- American Educational Research Association. (2011). *Code of ethics*. *Educational Researcher*, 40(3), 145–156.
- Amershi, S., Weld, D., Vorvoreanu, M., Fournery, A., Nushi, B., Collisson, P., Suh, J., Iqbal, S., Bennett, P. N., Inkpen, K., Teevan, J., Kikin-Gil, R., & Horvitz, E. (2019). Guidelines for human-AI interaction. In *Proceedings of the 2019 CHI Conference on Human Factors in Computing Systems* (pp. 1-13).
- Anderson, J. C., & Gerbing, D. W. (1991). Predicting the performance of measures in a confirmatory factor analysis with a pretest assessment of their substantive validities. *Journal of Applied Psychology*, 76(5), 732-740.
- Anderson, L. W., & Krathwohl, D. R. (2001). *A taxonomy for learning, teaching, and assessing: A revision of Bloom's taxonomy of educational objectives*. Longman.
- Andreoli, R. A. (1996). *Micromondi linguistici. L'uso di LOGO nella didattica dell'italiano*. Firenze: La Nuova Italia.
- Aoun, J. E. (2017). *ROBOT-PROOF: Higher education in the age of artificial intelligence*. Cambridge, MA: MIT Press.
- Arthur Jr, W., Bennett Jr, W., Stanush, P. L., & McNelly, T. L. (1998). Factors that influence skill decay and retention: A quantitative review and analysis. *Human Performance*, 11(1), 57-101.
- Asimov, I. (1942). Runaround. *Astounding Science Fiction*, March 1942.
- Asimov, I. (1991). "Chissà come si divertivano!". Tutti i racconti. Arnoldo Mondadori, (Titolo originale: *The Fun They Had!*, pubblicato in *Magazine of Fantasy and S.F.*, 1954).
- Atlas.ti Scientific Software Development GmbH. (2023). *ATLAS.ti Mac* (version 23.2.1) [Qualitative data analysis software]. <https://atlasti.com>
- Bandura, A. (1982). Self-efficacy mechanism in human agency. *American Psychologist*, 37(2), 122-147.
- Barab, S., & Squire, K. (2004). Design-based research: Putting a stake in the ground. *The Journal of the Learning Sciences*, 13(1), 1-14.
- Bawden, D. (2008). Origins and concepts of digital literacy. In C. Lankshear & M. Knobel (Eds.), *Digital literacies: Concepts, policies and practices* (pp. 17–32). Peter Lang.
- Benanti, P. (2018). *Homo Faber: The Techno-Human Condition*. Bologna: EDB.
- Benjamini, Y., & Hochberg, Y. (1995). Controlling the false discovery rate: a practical and powerful approach to multiple testing. *Journal of the Royal Statistical Society: Series B (Methodological)*, 57(1), 289-300.

- Bennato, D. (2002). *Le metafore del computer. La costruzione sociale dell'informatica*. Roma: Meltemi.
- Bennato, D. (2015). *Il computer come macroscopio. Big data e approccio computazionale per comprendere i cambiamenti sociali e culturali*. Milano: Franco Angeli.
- Bennett, M. S. (2023). *A brief history of intelligence: evolution, AI, and the five breakthroughs that made our brains*. HarperCollins.
- Bers, M. U. (2018). *Coding as a playground: Programming and computational thinking in the early childhood classroom*. Routledge.
- Bhargava, R., & D'Ignazio, C. (2015). Designing tools and activities for data literacy learners. In Workshop on Data Literacy, Webscience.
- Biagini, G. (2024). Assessing the assessments: Toward a multidimensional approach to AI literacy. *Media Education*, 15(1), 91-101. doi: 10.36253/me-15831
- Biagini, G., Cuomo, S., & Ranieri, M. (2023). Developing and validating a multidimensional AI literacy questionnaire: Operationalizing AI literacy for higher education. *Proceedings of the First International Workshop on High-performance Artificial Intelligence Systems in Education*. Rome, Italy, November 6, 2023.
- Biagini, G., Cuomo, S., & Ranieri, M. (2024). Assessing AI literacy: A framework-based approach. In T. Minerva, A. De Santis, Proceedings of the Italian Symposium on Digital Education, ISYDE2023, pp. 1-14, Reggio Emilia (Italy), September 13-15, Pearson.
- Biggs, J. (1996). Enhancing teaching through constructive alignment. *Higher Education*, 32(3), 347-364.
- Biggs, J., & Tang, C. (2011). *Teaching for quality learning at university* (4th ed.). Open University Press.
- Blei, D. M., Ng, A. Y., & Jordan, M. I. (2003). Latent Dirichlet allocation. *Journal of Machine Learning Research*, 3, 993-1022.
- Bocconi, S., Chiocciariello, A., Kampylis, P., Dagienė, V., Wastiau, P., Engelhardt, K., Earp, J., Horvath, M., Jasutė, E., Malagoli, C., Masiulionytė-Dagienė, G., & Stupurienė, G. (2022). Reviewing computational thinking in compulsory education. Publications Office of the European Union.
- Boden, M. (2019). *Artificial Intelligence: A Very Short Introduction*. Oxford: Oxford University Press.
- Bogliolo, A. (2018). *Coding in Your Classroom, Now! Il pensiero computazionale è per tutti, come la scuola* (2a ed.). Firenze: Giunti. (Opera originale pubblicata nel 2016).
- Bonaiuti, G., Calvani, A., & Ranieri, M. (2016). *Fondamenti di didattica. Teoria e prassi dei dispositivi formativi* (3a ed.). Roma: Carocci. (Opera originale pubblicata nel 2007).
- Bossuet, G. (1983). *L'ordinateur à l'école: Le système LOGO* (S. Papert, Pref.). PUF.
- Bostrom, N. (2014). *Superintelligence: Paths, Dangers, Strategies*. Oxford: Oxford University Press. (Trad. it. Superintelligenza. Tendenze, pericoli, strategie, Torino: Bollati Boringhieri, 2014).

- Bostrom, N., & Yudkowsky, E. (2014). The ethics of artificial intelligence. In K. Frankish & W. M. Ramsey (Eds.), *The Cambridge handbook of artificial intelligence* (pp. 316–334). Cambridge University Press.
- Boud, D. (2001). Using journal writing to enhance reflective practice. *New Directions for Adult and Continuing Education*, 2001(90), 9-18.
- Braun, H. (2005). Value-added modeling: What does due diligence require? In R. Lissitz (Ed.), *Value added models in education: Theory and applications* (pp. 19-38). JAM Press.
- Braun, V., & Clarke, V. (2006). Using thematic analysis in psychology. *Qualitative Research in Psychology*, 3(2), 77-101.
- Braun, V., & Clarke, V. (2012). Thematic analysis. In H. Cooper, P. M. Camic, D. L. Long, A. T. Panter, D. Rindskopf, & K. J. Sher (Eds.), *APA handbook of research methods in psychology* (Vol. 2, pp. 57-71). American Psychological Association.
- British Educational Research Association. (2018). *Ethical guidelines for educational research* (4th ed.). BERA.
- Brookfield, S. D. (1995). *Becoming a critically reflective teacher*. Jossey-Bass.
- Bruinsma, M., & Jansen, E. P. (2010). Is the motivation to become a teacher related to pre-service teachers' intentions to remain in the profession? *European Journal of Teacher Education*, 33(2), 185-200.
- Bruner, J. S. (1960). *The process of education*. Harvard University Press.
- Bruner, J. S. (1961). The act of discovery. *Harvard Educational Review*, 31, 21-32.
- Buckingham, D. (2003). *Media education: Literacy, learning, and contemporary culture*. Polity Press.
- Buckingham, D. (2019). *The media education manifesto*. Polity Press.
- Burgsteiner, H., Kandlhofer, M., & Steinbauer, G. (2016a). Artificial intelligence programming in secondary schools: Experiences and lessons learned from secondary schools. *International Journal of Artificial Intelligence in Education*, 26(1), 519–541. <https://doi.org/10.1007/s40593-016-0112-5>
- Burgsteiner, H., Kandlhofer, M., & Steinbauer, G. (2016b). IRobot: Teaching the basics of artificial intelligence in high schools. *Proceedings of the AAAI Conference on Artificial Intelligence*, 30(1).
- Calvani, A., Cartelli, A., Fini, A., & Ranieri, M. (2009). Models and instruments for assessing digital competence at school. *Journal of E-Learning and Knowledge Society*, 4(3), 183-193.
- Calvani, A., Fini, A., & Ranieri, M. (2011). *Valutare la competenza digitale. Prove per la scuola primaria e secondaria*. Trento: Edizioni Erikson.
- Calvani, A., Fini, A., & Ranieri, M. (2010). Digital competence in K-12: theoretical models, assessment tools and empirical research. *Analisi: Quaderns de Comunicació i Cultura*, (40), 157-171.
- Campbell, D. T., & Stanley, J. C. (1963). *Experimental and quasi-experimental designs for research*. Rand McNally.

- Cappello, G. (2017). Literacy, media literacy and social change. Where do we go from now?. *Italian Journal of Sociology of Education*, 9(1), 31-44.
- Capponi, M. (2008). *Un giocattolo per la mente. L'«informatica cognitiva» di Seymour Papert*. Perugia: Morlacchi.
- Cardon, P., Fleischmann, C., Aritz, J., Logemann, M., & Heidewald, J. (2023). The Challenges and Opportunities of AI-Assisted Writing: Developing AI literacy for the AI Age. *Business and Professional Communication Quarterly*, 232949062311765. <https://doi.org/10.1177/23294906231176517>
- Carolus, A., Augustin, Y., Markus, A., & Wienrich, C. (2023). Digital interaction literacy model -- Conceptualizing competencies for literate interactions with voice-based AI systems. *Computers and Education: Artificial Intelligence*, 4, 100114. <https://doi.org/10.1016/j.caeai.2022.100114>
- Carolus, A., Koch, M. J., Straka, S., Latoschik, M. C., & Wienrich, C. (2023). MAILS – Meta AI literacy scale: Development and testing of an AI literacy questionnaire based on well-founded competency models and psychological change- and meta-competencies. *Computers in Human Behavior: Artificial Humans*, 1(2), 1-10.
- Carretero, S., Vuorikari, R., & Punie, Y. (2017). DigComp 2.1: The Digital Competence Framework for Citizens with eight proficiency levels and examples of use. Publications Office of the European Union.
- Casal-Otero, L., Catala, A., Fernández-Morante, C., Taboada, M., Cebreiro, B., & Barro, S. (2023). AI literacy in K-12: A systematic literature review. *International Journal of STEM Education*, 10(1), 29. <https://doi.org/10.1186/s40594-023-00418-7>
- Cave, S., Dihal, K., & Dillon, S. (Eds.). (2019). *AI narratives: A history of imaginative thinking about intelligent machines*. Oxford University Press.
- Cetindamar, D., Kitto, K., Wu, M., Zhang, Y., Abedin, B., & Knight, S. (2022). Explicating AI literacy of Employees at Digital Workplaces. *IEEE Transactions on Engineering Management*, 1–14. <https://doi.org/10.1109/TEM.2021.3138503>
- Charmaz, K. (2014). *Constructing grounded theory* (2nd ed.). Sage Publications.
- Chen, L., Chen, P., & Lin, Z. (2020). Artificial intelligence in education: A review. *IEEE Access*, 8, 75264-75278.
- Chevallard, Y. (1991). *La transposition didactique: Du savoir savant au savoir enseigné*. La Pensée Sauvage.
- Chignard, S. (2013). *A brief history of Open Data*. Paris: ParisTech Review.
- Chinese State Council. (2017). Notice of the State Council on issuing the next generation artificial intelligence development plan (Guofa [2017] No. 35). http://www.gov.cn/zhengce/content/2017-07/20/content_5211996.htm
- Chomsky, N. (2023, March 8). The False Promise of ChatGPT. *The New York Times*.
- Cobb, P., Confrey, J., DiSessa, A., Lehrer, R., & Schauble, L. (2003). Design experiments in educational research. *Educational Researcher*, 32(1), 9-13.

- Cohen, J. (1960). A coefficient of agreement for nominal scales. *Educational and Psychological Measurement*, 20, 37–46. <https://doi.org/10.1177/001316446002000104>
- Cohen, J. (1988). *Statistical power analysis for the behavioral sciences* (2nd ed.). Hillsdale, NJ: Lawrence Erlbaum Associates.
- Cohen, L., Manion, L., & Morrison, K. (2018). *Research methods in education* (8th ed.). Routledge.
- Connelly, F. M., & Clandinin, D. J. (1990). Stories of experience and narrative inquiry. *Educational Researcher*, 19(5), 2-14.
- Cook, T. D., & Campbell, D. T. (1979). *Quasi-experimentation: Design and analysis issues for field settings*. Houghton Mifflin.
- Cooper, H., Hedges, L. V., & Valentine, J. C. (2019). *The handbook of research synthesis and meta-analysis* (3rd ed.). Russell Sage Foundation.
- Cope, B., & Kalantzis, M. (2000). *Multiliteracies: Literacy learning and the design of social futures*. South Yarra, Victoria, Australia: Macmillan.
- Cortina, J. M. (1993). What is coefficient alpha? An examination of theory and applications. *Journal of Applied Psychology*, 78(1), 98-104.
- Crawford, K. (2021). *Atlas of AI: Power, Politics, and the Planetary Costs of Artificial Intelligence*. New Haven: Yale University Press.
- Creswell, J. W., & Creswell, J. D. (2018). *Research design: Qualitative, quantitative, and mixed methods approaches* (5th ed.). Sage Publications.
- Creswell, J. W., & Plano Clark, V. L. (2018). *Designing and conducting mixed methods research* (3rd ed.). Sage Publications.
- Crippa, M., & Girgenti, G. (2024). *Umano poco umano. Esercizi spirituali contro l'intelligenza artificiale*. Milano: La nave di Teseo.
- Cristianini, N. (2024). *Machina Sapiens. L'algoritmo che ci ha rubato il segreto della conoscenza*. Milano: Feltrinelli.
- Cronbach, L. J. (1951). Coefficient alpha and the internal structure of tests. *Psychometrika*, 16, 297–334.
- Cuomo, S., Biagini, G., & Ranieri, M. (2022). Artificial Intelligence Literacy, che cos'è e come promuoverla. Dall'analisi della letteratura ad una proposta di Framework. *Media Education*, 13(2), 19-28. <https://doi.org/10.36253/me-13374>
- Darling-Hammond, L. (2017). Teacher education around the world: What can we learn from international practice? *European Journal of Teacher Education*, 40(3), 291-309.
- Denzin, N. K. (1978). *The research act: A theoretical introduction to sociological methods* (2nd ed.). McGraw-Hill.
- Design-Based Research Collective. (2003). Design-based research: An emerging paradigm for educational inquiry. *Educational Researcher*, 32(1), 5-8.

- Deuze, M., & Beckett, C. (2022). Imagination, Algorithms and News: Developing AI literacy for Journalism. *Digital Journalism*, 10(10), 1913–1918. <https://doi.org/10.1080/21670811.2022.2119152>
- DeVellis, R. F. (2016). *Scale development: Theory and applications* (Vol. 26). Sage publications.
- Dewey, J. (1938). *Experience and education*. Macmillan.
- Dimitrov, D. M., & Rumrill Jr, P. D. (2003). Pretest-posttest designs and measurement of change. *Work*, 20(2), 159-165.
- Dreyfus, H. L. (1992). *What computers still can't do: A critique of artificial reason*. MIT Press.
- Drudy, S. (2008). Gender balance/gender bias: The teaching profession and the impact of feminisation. *Gender and Education*, 20(4), 309-323.
- Druga, S., Christoph, F. L., & Ko, A. J. (2022). Family as a Third Space for AI Literacies: How do children and parents learn about AI together? CHI Conference on Human Factors in Computing Systems, 1–17. <https://doi.org/10.1145/3491102.3502031>
- Druga, S., Vu, S. T., Likhith, E., & Qiu, T. (2019). Inclusive AI literacy for kids around the world. *Proceedings of FabLearn 2019*, 104-111. <https://doi.org/10.1145/3311890.3311904>
- Druga, S., Williams, R., Park, H. W., & Breazeal, C. (2019). How smart are the smart toys? Children and parents' agent interaction and intelligence attribution. *Proceedings of the 17th ACM Conference on Interaction Design and Children*, 231-240.
- Eisinga, R., Te Grotenhuis, M., & Pelzer, B. (2013). The reliability of a two-item scale: Pearson, Cronbach, or Spearman-Brown? *International Journal of Public Health*, 58(4), 637-642.
- Eisner, E. W. (2017). *The enlightened eye: Qualitative inquiry and the enhancement of educational practice*. Teachers College Press.
- Elliott, A. (2021). *Vita quotidiana e rivoluzione digitale*. Torino: Codice edizioni. (Opera originale pubblicata nel 2019 con il titolo *The Culture of AI. Everyday Life and the Digital Revolution*).
- Elliott, S. (2021). *Artificial Intelligence in Education: Where is It Now, and What Can It Deliver?* Oxford: BERA.
- Ennis, R. H. (2011). Critical thinking: Reflection and perspective Part I. *Inquiry: Critical Thinking Across the Disciplines*, 26(1), 4-18.
- Erceg-Hurn, D. M., & Mirosevich, V. M. (2008). Modern robust statistical methods: an easy way to maximize the accuracy and power of your research. *American Psychologist*, 63(7), 591-601.
- Ertmer, P. A., Ottenbreit-Leftwich, A. T., Sadik, O., Sendurur, E., & Sendurur, P. (2012). Teacher beliefs and technology integration practices: A critical relationship. *Computers & Education*, 59(2), 423-435.
- Eubanks, V. (2018). *Automating inequality: How high-tech tools profile, police, and punish the poor*. St. Martin's Press.
- European Commission, Joint Research Centre (JRC) & Organisation for Economic Co-operation and Development (OECD). (2021). *AI watch, national strategies on artificial intelligence: a European perspective*. Publications Office of the European Union. <https://doi.org/10.2760/069178>

- European Commission, Joint Research Centre (JRC). (2018). *The impact of Artificial Intelligence on learning, teaching, and education*. Publications Office. <https://data.europa.eu/doi/10.2760/12297>
- European Commission. (2018). *High-Level Expert Group on Artificial Intelligence*. <https://digital-strategy.ec.europa.eu/en/policies/expert-group-ai>
- European Commission. (2019). *Digital Education Action Plan (2021-2027)*. Luxembourg: Publications Office of the European Union.
- European Commission. (2019). *Pilot the Assessment List of the Ethics Guidelines for Trustworthy AI*. <https://ec.europa.eu/futurium/en/ethics-guidelines-trustworthy-ai/register-piloting-process-0.html>
- European Commission. (2020). *On Artificial Intelligence - A European approach to excellence and trust*. Technical report, Brussels.
- European Commission. (2021). *Shaping Europe's digital future - European strategy for data*.
- Facione, P. A. (2011). *Critical thinking: What it is and why it counts*. Measured Reasons LLC.
- Faggin, F. (2022). *Irriducibile. La coscienza, la vita, i computer e la nostra natura*. Milano: Mondadori.
- Faggin, F. (2024). *Oltre l'invisibile*. Milano: Edizioni Mondadori.
- Feenberg, A. (1991). *Critical theory of technology*. Oxford University Press.
- Feenberg, A. (2002). *Transforming technology: A critical theory revisited*. Oxford University Press.
- Ferrari, A., Brecko, B., & Punie, Y. (2014). *DIGCOMP: a Framework for Developing and Understanding Digital Competence in Europe*. Luxembourg: Publications Office of the European Union.
- Ferraris, M. (2023). A chi fa paura l'intelligenza artificiale. Gli improbabili timori circa la presa di potere delle macchine. *la Repubblica*, 21 maggio 2023.
- Field, A. (2013). *Discovering statistics using IBM SPSS statistics*. Sage.
- Field, A., Miles, J., & Field, Z. (2012). *Discovering statistics using R*. Sage Publications.
- Floridi, L. (2012). *La rivoluzione dell'informazione* (M. Durante, Trad.; J. C. De Martin, Pref.). Codice Edizioni. (Opera originale pubblicata 2010)
- Floridi, L. (2014). *The Fourth Revolution. How the infosphere is reshaping human reality*. Oxford: Oxford University Press. (Trad. it. *La quarta rivoluzione. Come l'infosfera sta trasformando il mondo*, Milano: Cortina, 2017).
- Floridi, L. (2017). *La quarta rivoluzione: Come l'intelligenza sta trasformando il mondo* (M. Durante, Trad.). R. Cortina Editore. (Opera originale pubblicata 2014)
- Floridi, L. (2020). AI and its new winter: From myths to realities. *Philosophy & Technology*, 33(1), 1-3.
- Floridi, L. (2021). Introduction—The importance of an ethics-first approach to the development of AI. In *Ethics, governance, and policies in artificial intelligence* (pp. 1-4).
- Floridi, L. (2022). Against digital ethics washing. In *Handbook of Digital Ethics* (pp. 31-49). Routledge.

- Floridi, L. (2022). *Etica dell'intelligenza artificiale: Sviluppi, opportunità, sfide* (M. Durante, Trad.). R. Cortina Editore. (Opera originale pubblicata 2022)
- Floridi, L., & Cabitza, F. (2021). *Intelligenza artificiale. L'uso delle nuove macchine*. Milano: Bompiani.
- Floridi, L., Cows, J., Beltrametti, M., Chatila, R., Chazerand, P., Dignum, V., Luetge, C., Madelin, R., Pagallo, U., Rossi, F., Schafer, B., Valcke, P., & Vayena, E. (2018). AI4People—an ethical framework for a good AI society: opportunities, risks, principles, and recommendations. *Minds and Machines*, 28(4), 689-707.
- Fornell, C., & Larcker, D. F. (1981). Evaluating structural equation models with unobservable variables and measurement error. *Journal of Marketing Research*, 18(1), 39–50.
- Fortunato, S. (2010). Community detection in graphs. *Physics Reports* 486, 75–174.
<https://doi.org/10.1016/j.physrep.2009.11.002>
- Fosnot, C. T. (2013). *Constructivism: Theory, perspectives, and practice*. Teachers College Press.
- Freire, P. (1970). *Pedagogy of the oppressed*. Continuum.
- Freire, P. (1973). *Education for critical consciousness*. Seabury Press.
- Fritz, C. O., Morris, P. E., & Richler, J. J. (2012). Effect size estimates: current use, calculations, and interpretation. *Journal of Experimental Psychology: General*, 141(1), 2-18.
- Gagné, R. M., Wager, W. W., Golas, K. C., & Keller, J. M. (2005). *Principles of instructional design* (5th ed.). Wadsworth/Thomson Learning.
- Gall, M. D., Gall, J. P., & Borg, W. R. (2007). *Educational research: An introduction* (8th ed.). Pearson.
- Gattico, E., & Storari, G. P. (2005). *Costruttivismo e scienze della formazione* (M. Ceruti, Pref.; R. Ortù & D. Seddone, Contrib.). Unicopli.
- GDPR (General Data Protection Regulation). (2016). Regulation (EU) 2016/679 of the European Parliament and of the Council of 27 April 2016 on the protection of natural persons with regard to the processing of personal data and on the free movement of such data. *Official Journal of the European Union*, 59, 1-88.
- Gee, J. (2003). *What video games have to teach us about learning and literacy*. New York: Palgrave, Macmillan.
- George, D., & Mallery, P. (2003). *SPSS for Windows step by step: A simple guide and reference*. 11.0 update (4th ed.). Boston: Allyn & Bacon.
- Ghallab, M. (2019). Responsible AI and AI literacy. *AI Magazine*, 40(2), 3–4.
<https://doi.org/10.1609/aimag.v40i2.2855>
- Ghasemi, A., & Zahediasl, S. (2012). Normality tests for statistical analysis: a guide for non-statisticians. *International Journal of Endocrinology and Metabolism*, 10(2), 486-489.
- Giddens, A. (1984). *The constitution of society: Outline of the theory of structuration*. University of California Press.
- Gilster, P. (1997). *Digital literacy*. New York: Wiley Computer Publishing.

- Giroux, H. A. (2011). *On critical pedagogy*. Continuum International Publishing Group.
- Glaserfeld, E. von (1995). *Radical constructivism: A way of knowing and learning*. Falmer Press.
- Goodfellow, I., Pouget-Abadie, J., Mirza, M., Xu, B., Warde-Farley, D., Ozair, S., ... & Bengio, Y. (2014). Generative adversarial nets. *Advances in Neural Information Processing Systems*, 27.
- Grant, M. J., & Booth, A. (2009). A typology of reviews: an analysis of 14 review types and associated methodologies. *Health information & libraries journal*, 26(2), 91-108.
- Greene, J. C., Caracelli, V. J., & Graham, W. F. (1989). Toward a conceptual framework for mixed-method evaluation designs. *Educational Evaluation and Policy Analysis*, 11(3), 255-274.
- Grimmer, J., & Stewart, B. M. (2013). Text as data: The promise and pitfalls of automatic content analysis methods for political texts. *Political Analysis*, 21(3), 267-297.
- Hair, J., Black, W. C., Babin, B. J., & Anderson, R. E. (2010). *Multivariate data analysis* (7th ed.). Upper Saddle River, NJ: Pearson Education International.
- Hamburg, I., O'brien, E., & Vladut, G. (2019). Entrepreneurial Learning and AI literacy to Support Digital Entrepreneurship. *Balkan Region Conference on Engineering and Business Education*, 3(1), 132–144. <https://doi.org/10.2478/cplbu-2020-0016>
- Hammersley, M., & Traianou, A. (2012). *Ethics in qualitative research: Controversies and contexts*. Sage Publications.
- Harari, Y. N. (2018). *21 Lessons for the 21st Century*. London: Jonathan Cape.
- Harper, D. (2002). Talking about pictures: A case for photo elicitation. *Visual Studies*, 17(1), 13-26.
- Hattie, J. (2009). *Visible learning: A synthesis of over 800 meta-analyses relating to achievement*. Routledge.
- Hermann, E. (2022). Artificial intelligence and mass personalization of communication content—An ethical and literacy perspective. *New Media & Society*, 24(5), 1258–1277. <https://doi.org/10.1177/14614448211022702>
- Hinkin, T. R. (1998). A brief tutorial on the development of measures for use in survey questionnaires. *Organizational Research Methods*, 1(1), 104-121.
- Hmelo-Silver, C. E. (2004). Problem-based learning: What and how do students learn? *Educational Psychology Review*, 16(3), 235-266.
- Hwang, G. J., Xie, H., & Shin, D. (2020). AI-powered adaptive learning and formative assessment: Addressing new challenges in education. *Computers & Education*, 150(1), 103836. <https://doi.org/10.1016/j.compedu.2020.103836>
- Jacobson, N. S., & Truax, P. (1991). Clinical significance: a statistical approach to defining meaningful change in psychotherapy research. *Journal of Consulting and Clinical Psychology*, 59(1), 12-19.
- Johnson, D. W., & Johnson, R. T. (1999). *Learning together and alone: Cooperative, competitive, and individualistic learning* (5th ed.). Allyn & Bacon.
- Johnson, D. W., & Johnson, R. T. (2009). An educational psychology success story: Social interdependence theory and cooperative learning. *Educational Researcher*, 38(5), 365-379.

- Johnson, R. B., & Onwuegbuzie, A. J. (2004). Mixed methods research: A research paradigm whose time has come. *Educational Researcher*, 33(7), 14-26.
- Johnson, W. (1944). Studies in language behavior: I. A program of research. *Psychological Monographs*, 56(2), 1-15.
- Jonas, H. (1984). *The Imperative of Responsibility: In Search of an Ethics for the Technological Age*. University of Chicago Press.
- Jonassen, D. H. (1999). Designing constructivist learning environments. *Instructional Design Theories and Models: A New Paradigm of Instructional Theory*, 2, 215-239.
- Jonassen, D. H., & Grabowski, B. L. (1993). *Handbook of individual differences, learning, and instruction*. Lawrence Erlbaum Associates.
- Jordan, M. I., & Mitchell, T. M. (2015). Machine learning: Trends, perspectives, and prospects. *Science*, 349(6245), 255-260.
- Kandlhofer, M., Steinbauer, G., Hirschmugl-Gaisch, S., & Huber, P. (2016). Artificial intelligence and computer science in education: From kindergarten to university. 2016 IEEE Frontiers in Education Conference (FIE), 1–9. <https://doi.org/10.1109/FIE.2016.7757570>
- Kang, Y.-N., Chen, D.-R., & Chen, Y.-Y. (2023). Associations between literacy and attitudes toward artificial intelligence--assisted medical consultations: The mediating role of perceived distrust and efficiency of artificial intelligence. *Computers in Human Behavior*, 139, 107529. <https://doi.org/10.1016/j.chb.2022.107529>
- Kaplan, J. (2017). *Intelligenza Artificiale. Guida al futuro prossimo*. Roma: LUISS University Press.
- Kellner, D., & Share, J. (2007). Critical media literacy is not an option. *Learning Inquiry*, 1(1), 59-69.
- Kim, S., Jang, Y., Kim, W., Choi, S., Jung, H., Kim, S., & Kim, H. (2021). Why and What to Teach: AI Curriculum for Elementary School. *Proceedings of the AAAI Conference on Artificial Intelligence*, 35(17), 15569–15576. <https://doi.org/10.1609/aaai.v35i17.17833>
- Kohlberg, L. (1984). *The psychology of moral development: The nature and validity of moral stages (Essays on moral development, Volume 2)*. Harper & Row.
- Kolb, D. A. (1984). *Experiential learning: Experience as the source of learning and development*. Prentice-Hall.
- Kong, S.-C., & Zhang, G. (2021). A Conceptual Framework for Designing Artificial Intelligence Literacy Programmes for Educated Citizens.
- Kong, S.-C., Man-Yin Cheung, W., & Zhang, G. (2021). Evaluation of an artificial intelligence literacy course for university students with diverse study backgrounds. *Computers and Education: Artificial Intelligence*, 2, 100026. <https://doi.org/10.1016/j.caeai.2021.100026>
- Kress, G. (2000). Multimodality. In B. Cope & M. Kalantzis (Eds.), *Multiliteracies: Literacy learning and the design of social futures* (pp. 182–202). South Yarra, Victoria, Australia: Macmillan

- Kruger, J., & Dunning, D. (1999). Unskilled and unaware of it: How difficulties in recognizing one's own incompetence lead to inflated self-assessments. *Journal of Personality and Social Psychology*, 77(6), 1121.
- Kurzweil, R. (2014). *La Singolarità è vicina*. Santarcangelo di Romagna: Maggioli. (Opera originale pubblicata nel 2005 con il titolo *The Singularity Is Near*).
- Lakens, D. (2013). Calculating and reporting effect sizes to facilitate cumulative science: a practical primer for t-tests and ANOVAs. *Frontiers in Psychology*, 4, 863.
- Lakens, D. (2017). Equivalence tests: A practical primer for t tests, correlations, and meta-analyses. *Social Psychological and Personality Science*, 8(4), 355-362.
- Lakoff, G., & Johnson, M. (1980). *Metaphors we live by*. University of Chicago Press.
- Landis, J., & Koch, G. (1977). The measurement of observer agreement for categorical data. *Biometrics*, 33(1), 159–174. <https://doi.org/10.2307/2529310>
- Lankshear, C., & Knobel, M. (2003). *New literacies: Changing knowledge and classroom learning*. Philadelphia: Open University Press.
- Laupichler, M. C., Aster, A., & Raupach, T. (2023). Delphi study for the development and preliminary validation of an item set for the assessment of non-experts' AI literacy. *Computers and Education: Artificial Intelligence*, 4, 100126. <https://doi.org/10.1016/j.caeai.2023.100126>
- Laupichler, M. C., Aster, A., Schirch, J., & Raupach, T. (2022). Artificial intelligence literacy in higher and adult education: A scoping literature review. *Computers and Education: Artificial Intelligence*, 3, 100101. <https://doi.org/10.1016/j.caeai.2022.100101>
- Laurillard, D. (2012). *Teaching as a design science: Building pedagogical patterns for learning and technology*. Routledge.
- Lave, J., & Wenger, E. (1991). *Situated learning: Legitimate peripheral participation*. Cambridge University Press.
- Lee, A. (2021). The Effect of ARTIFICIAL INTELLIGENCE Literacy Education on University Students Ethical Consciousness of Artificial Intelligence. *J-Institute*, 6(3), 52–61. <https://doi.org/10.22471/ai.2021.6.3.52>
- Lee, I., Ali, S., Zhang, H., DiPaola, D., & Breazeal, C. (2021). Developing Middle School Students' AI literacy. *Proceedings of the 52nd ACM Technical Symposium on Computer Science Education*, 191–197. <https://doi.org/10.1145/3408877.3432513>
- Legrenzi, P. (2024). *Pensare con l'IA*. Bologna: Il Mulino.
- Likert, R. (1932). A technique for the measurement of attitudes. *Archives of Psychology*, 22(140), 5-55.
- Lincoln, Y. S., & Guba, E. G. (1985). *Naturalistic inquiry*. Sage Publications.
- Linn, R. L., & Slinde, J. A. (1977). The determination of the significance of change between pre- and posttesting periods. *Review of Educational Research*, 47(1), 121-150.
- Lipman, M. (1976). Philosophy for children. *Metaphilosophy*, 7(1), 17-39.
- Lipman, M. (2003). *Thinking in education* (2nd ed.). Cambridge University Press.

- Livingstone, S. (2004). What is media literacy? *InterMedia*, 32(3), 18–20.
- Long, D., & Magerko, B. (2020). What is AI literacy? Competencies and design considerations. In *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems* (pp. 1-16). ACM.
- Long, D., Roberts, J., Magerko, B., Holstein, K., DiPaola, D., & Martin, F. (2023). AI literacy: Finding Common Threads between Education, Design, Policy, and Explainability. Extended Abstracts of the 2023 CHI Conference on Human Factors in Computing Systems, 1–6.
<https://doi.org/10.1145/3544549.3573808>
- Lovelock, J. (1991). *Le nuove età di Gaia. Una biografia del nostro mondo vivente*. Torino: Bollati Boringhieri. (Opera originale pubblicata nel 1988 con il titolo *The Ages of Gaia. A Biography of Our Living Earth*).
- Lumley, T., Diehr, P., Emerson, S., & Chen, L. (2002). The importance of the normality assumption in large public health data sets. *Annual Review of Public Health*, 23(1), 151-169.
- Manning, C. D., Raghavan, P., & Schütze, H. (2008). *Introduction to information retrieval*. Cambridge University Press.
- Martin, A. (2005). DigEuLit—a European framework for digital literacy: a progress report. *Journal of eLiteracy*, 2(2), 130-136.
- Martinez, M. A., Sauleda, N., & Huber, G. L. (2001). Metaphors as blueprints of thinking about teaching and learning. *Teaching and Teacher Education*, 17(8), 965-977.
- Maxwell, J. A. (2013). *Qualitative research design: An interactive approach* (3rd ed.). Sage Publications.
- Mayer-Schönberger, V., & Cukier, K. (2013). *Big Data: A Revolution That Will Transform How We Live, Work, and Think*. London: John Murray. (Trad. it. *Big Data. Una rivoluzione che trasformerà il nostro modo di vivere e già minaccia la nostra libertà*, Milano: Garzanti, 2013).
- McCarthy, J. (2007). From here to human-level AI. *Artificial Intelligence*, 171(18), 1174-1182.
- McDonald, J. H. (2014). *Handbook of biological statistics* (3rd ed.). Baltimore, MD: Sparky House Publishing.
- McDonald, R. P. (1999). *Test theory: A unified treatment*. Mahwah, N.J.: L. Erlbaum Associates.
- Mello-Thoms, C. (2023). Teaching Artificial Intelligence Literacy: A Challenge in the Education of Radiology Residents. *Academic Radiology*, 30(7), 1488–1490.
<https://doi.org/10.1016/j.acra.2023.04.035>
- Merrill, M. D. (2002). First principles of instruction. *Educational Technology Research and Development*, 50(3), 43-59.
- Mertala, P., Fagerlund, J., & Calderon, O. (2022). Finnish 5th and 6th grade students' pre-instructional conceptions of artificial intelligence (AI) and their implications for AI literacy education. *Computers and Education: Artificial Intelligence*, 3, 100095.
<https://doi.org/10.1016/j.caeai.2022.100095>
- Mezirow, J. (1991). *Transformative dimensions of adult learning*. Jossey-Bass.

- Miao, F., Holmes, W., Huang, R., & Zhang, H. (2021). *AI and education: Guidance for policy-makers*. Paris: UNESCO.
- Miles, M. B., & Huberman, A. M. (1994). *Qualitative data analysis: An expanded sourcebook* (2nd ed.). Sage Publications.
- Mills, K. A. (2010). A Review of the "Digital Turn" in the New Literacy Studies. *Review of Educational Research*, 80(2), 246–271.
- Mishra, P., & Koehler, M. J. (2006). Technological pedagogical content knowledge: A framework for teacher knowledge. *Teachers College Record*, 108(6), 1017-1054.
- Mitchell, M. (2022). *L'intelligenza artificiale. Una guida per esseri umani pensanti*. Torino: Einaudi.
- Moher, D., Liberati, A., Tetzlaff, J., Altman, D. G., & The PRISMA Group. (2009). Preferred reporting items for systematic reviews and meta-analyses: The PRISMA statement. *PLoS Medicine*, 6(7), e1000097.
- Mollick, E. (2024). *Co-intelligence: living and working with AI*. Portfolio/Penguin.
- Moon, J. A. (2006). *Learning journals: A handbook for reflective practice and professional development* (2nd ed.). Routledge.
- National AI Initiative Act of 2020, H.R. 6216, 116th Cong. (2020).
<https://www.congress.gov/bill/116th-congress/house-bill/6216>
- Nelson, L. K. (2020). Computational grounded theory: A methodological framework. *Sociological Methods & Research*, 49(1), 3-42.
- Ng, D. T. K., Leung, J. K. L., Chu, S. K. W., & Qiao, M. S. (2021). Conceptualizing AI literacy: An exploratory review. *Computers and Education: Artificial Intelligence*, 2, 100041.
<https://doi.org/10.1016/j.caeai.2021.100041>
- Ng, D. T. K., Wu, W., Leung, J. K. L., & Chu, S. K. W. (2023). Artificial Intelligence (AI) literacy questionnaire with confirmatory factor analysis. In *2023 IEEE International Conference on Advanced Learning Technologies (ICALT)* (pp. 233–235). IEEE.
<https://doi.org/10.1109/ICALT58122.2023.00074>
- Ng, J. (2021). Artificial intelligence literacy for everyone. In *Handbook of Digital Higher Education* (pp. 321-334). Edward Elgar Publishing.
- Noddings, N. (2013). *Caring: A relational approach to ethics and moral education* (2nd ed.). University of California Press.
- Nowotny, H. (2022). *In AI We Trust: Power, Illusion and Control of Predictive Algorithms*. Cambridge: Polity Press.
- Nunnally, J. (1978). *Psychometric theory*. New York: McGraw-Hill.
- OECD. (2018). *OECD Learning Framework 2030*. Paris: OECD Publishing.
- O'Neil, C. (2016). *Weapons of math destruction: How big data increases inequality and threatens democracy*. Crown.
- Ong, W. (1982). *Oralità e scrittura. Le tecnologie della parola*. Bologna: il Mulino.

- Organisation for Economic Co-operation and Development (OECD). (2018a). *Bridging the digital gender divide: Include, upskill, innovate*. OECD Publishing. <http://www.oecd.org/digital/bridging-the-digital-gender-divide.pdf>
- Organisation for Economic Co-operation and Development (OECD). (2018b). *Future of education and skills 2030: Conceptual learning framework*. OECD Publishing. <https://www.oecd.org/education/2030/Education-and-AI-preparing-for-the-future-AI-Attitudes-and-Values.pdf>
- Organisation for Economic Co-operation and Development (OECD). (2019). *Recommendation of the Council on Artificial Intelligence*. OECD Publishing. <https://legalinstruments.oecd.org/en/instruments/OECD-LEGAL-0449>
- Orlandi, T. (1999). *Informatica umanistica*. Roma: NIS.
- Orlandi, T. (2010). *Informatica testuale. Teoria e prassi*. Roma-Bari: Laterza.
- Ouzzani, M., Hammady, H., Fedorowicz, Z., & Elmagarmid, A. (2016). Rayyan—a web and mobile app for systematic reviews. *Systematic Reviews*, 5(1), 210. <https://doi.org/10.1186/s13643-016-0384-4>
- P21 (Partnership for 21st Century Learning). (2019). Framework for 21st century learning. Battelle for Kids.
- Pacchioni, G. (2021). *L'ultimo sapiens. Viaggio al termine della nostra specie*. Bologna: Il Mulino.
- Pallant, J. (2020). *SPSS survival manual: A step by step guide to data analysis using IBM SPSS* (7th ed.). Routledge.
- Pancioli, C., & Rivoltella, P. C. (2023). *Pedagogia algoritmica. Per una riflessione educativa sull'Intelligenza artificiale*. Brescia: Scholé.
- Pancioli, C., Allegra, M., Gentile, M., & Rivoltella, P. C. (2023). Towards AI literacy: A proposal of a framework based on the Episodes of Situated Learning.
- Papert, S. (1980). *Mindstorms: Children, Computers, and Powerful Ideas*. New York: Basic Books.
- Park, J., Teo, T. W., Teo, A., Chang, J., Huang, J. S., & Koo, S. (2023). Integrating artificial intelligence into science lessons: Teachers' experiences and views. *International Journal of STEM Education*, 10(1), 61. <https://doi.org/10.1186/s40594-023-00454-3>
- Paul, R., & Elder, L. (2020). *Critical thinking: Tools for taking charge of your learning and your life* (3rd ed.). Foundation for Critical Thinking Press.
- Payne, B. H. (2019). An ethics of artificial intelligence curriculum for middle school students. MIT Media Lab, available at: <https://www.media.mit.edu/projects/ai-ethics-for-middle-school/overview/>
- Pedró, F., Subosa, M., Rivas, A., & Valverde, P. (2019). Artificial intelligence in education: Challenges and opportunities for sustainable development. Paris: UNESCO.
- Pellegrino, J. W., DiBello, L. V., & Goldman, S. R. (2016). A framework for conceptualizing and evaluating the validity of instructionally relevant assessments. *Educational Psychologist*, 51(1), 59-81.

- Perchik, J. D., Smith, A. D., Elkassem, A. A., Park, J. M., Rothenberg, S. A., Tanwar, M., Yi, P. H., Sturdivant, A., Tridandapani, S., & Sotoudeh, H. (2023). Artificial Intelligence Literacy: Developing a Multi-institutional Infrastructure for AI Education. *Academic Radiology*, 30(7), 1472–1480. <https://doi.org/10.1016/j.acra.2022.10.002>
- Perkins, D. N., & Salomon, G. (1992). Transfer of learning. *International Encyclopedia of Education*, 2, 6452-6457.
- Perneger, T. V. (1998). What's wrong with Bonferroni adjustments. *BMJ*, 316(7139), 1236-1238.
- Phillips, P. J., Hahn, C. A., Fontana, P. C., Broniatowski, D. A., & Przybocki, M. A. (2022). Four principles of explainable artificial intelligence. NIST Interagency/Internal Report (NISTIR), 8312.
- Piaget, J. (1936). *La naissance de l'intelligence chez l'enfant*. Neuchâtel: Delachaux et Niestlé.
- Piaget, J. (1964). *Cognitive development in children*. *Journal of Research in Science Teaching*, 2, 176-186.
- Pian, A. (2000). *L'ora di Internet. Manuale critico di pedagogia informatica*. Firenze: La Nuova Italia.
- Pinsky, M., & Benlian, A. (2023). AI literacy – towards measuring human competency in artificial intelligence. In *Proceedings of the 56th Hawaii international conference on system sciences* (pp. 165–174). Maui, USA: IEEE Computer Society Press.
- Prencipe, A. (2024). *Università generativa. Internazionale, interdisciplinare, innovativa*. Bologna: Il Mulino
- Prensky, M. (2001). Digital natives, digital immigrants part 1. *On the Horizon*, 9(5), 1-6.
- Quintarelli, E. (2020). *Intelligenza Artificiale: cos'è davvero, come funziona, che effetti avrà*. Torino: Bollati Boringhieri.
- R Core Team. (2021). *R: A language and environment for statistical computing*. R Foundation for Statistical Computing. <https://www.R-project.org/>
- Rachels, J., & Rachels, S. (2019). *The elements of moral philosophy* (9th ed.). McGraw-Hill.
- Rainie, L., & Anderson, J. (2017). Code-dependent: Pros and cons of the algorithm age. Pew Research Center.
- Ranieri, M. (2011). *Le insidie dell'ovvio. Tecnologie educative e critica della retorica tecnocentrica*. Pisa: ETS.
- Ranieri, M. (2019). Literacy, Technology, and Media. In: Renee Hobbs and Paul Mihailidis (Editors-in-Chief), Gianna Cappello, Maria Ranieri, and Benjamin Thevenin (Associate Editors). *The International Encyclopedia of Media Literacy*, pp. 1-12, Hoboken, NJ: JohnWiley & Sons, Inc.
- Ranieri, M. (2019). *Media Literacy Education between old and new challenges in a hyper-connected world*. In: M. Gui (ed.) *Digital well-being at school and at home. A path to media literacy education in the permanent connection*, Mondadori Università, Florence, pp. 9-35.
- Ranieri, M. (2022). *Competenze digitali per insegnare. Modelli e proposte operative*. Roma: Carocci.

- Ranieri, M. (Ed.). (2018). *Teoria e pratica delle New Media Literacies nella scuola*. Roma: Th. Factory.
- Ranieri, M., Cuomo, S., & Biagini, G. (2024). *Scuola e Intelligenza Artificiale. Percorsi di alfabetizzazione critica*. Roma: Carocci.
- Razali, N. M., & Wah, Y. B. (2011). Power comparisons of Shapiro-Wilk, Kolmogorov-Smirnov, Lilliefors and Anderson-Darling tests. *Journal of Statistical Modeling and Analytics*, 2(1), 21-33.
- Ricciardi, M. (2019). *Artificial Intelligence: ultima chiamata. Il sistema Europa sfida la supremazia tecnologica di Cina e Stati Uniti*. Roma: Università Luiss.
- Riguzzi, F. (2006). *Fondamenti di Intelligenza Artificiale*. Bologna: Esculapio.
- Rivoltella, P. C. (2020). *Nuovi alfabeti. Educazione e culture post-mediale*. Trento: Scholé.
- Rivoltella, P. C., & Panciroli, C. (2023). *Pedagogia algoritmica: Per una riflessione educativa sull'intelligenza artificiale*. Scholé.
- Roll, I., & Wylie, R. (2016). Evolution and revolution in artificial intelligence in education. *International Journal of Artificial Intelligence in Education*, 26(2), 582-599.
- Romano, L. (2024). *Artificio e intelligenza. Prospettive digitali e pedagogiche*. Milano: Mondadori Education.
- Roncaglia, G. (2018). *L'età della frammentazione. Cultura del libro e scuola digitale*. Roma-Bari: Laterza.
- Roncaglia, G. (2023). *L'architetto e l'oracolo. Forme digitali del sapere da Wikipedia a ChatGPT*. Roma-Bari: Laterza.
- Rosenbaum, P. R. (2002). *Observational studies* (2nd ed.). New York: Springer.
- Rosenthal, R. (1991). *Meta-analytic procedures for social research*. Sage.
- Rosenthal, R., & DiMatteo, M. R. (2001). Meta-analysis: Recent developments in quantitative methods for literature reviews. *Annual Review of Psychology*, 52(1), 59-82.
- Rudin, C. (2019). Stop explaining black box machine learning models for high stakes decisions and use interpretable models instead. *Nature Machine Intelligence*, 1(5), 206-215.
- Russell, S. J., & Norvig, P. (2020). *Artificial intelligence: A modern approach* (4th ed.). Pearson.
- Saban, A., Kocbeker, B. N., & Saban, A. (2007). Prospective teachers' conceptions of teaching and learning revealed through metaphor analysis. *Learning and Instruction*, 17(2), 123-139.
- Sakai, T. (2023). Literacy for Appropriate Use of Medical Big Data and Artificial Intelligence. *YAKUGAKU ZASSHI*, 143(6), 491–495. <https://doi.org/10.1248/yakushi.22-00179-2>
- Savery, J. R. (2006). Overview of problem-based learning: Definitions and distinctions. *Interdisciplinary Journal of Problem-Based Learning*, 1(1), 9-20.
- Savery, J. R., & Duffy, T. M. (1995). Problem based learning: An instructional model and its constructivist framework. *Educational Technology*, 35(5), 31-38.
- Schank, R. C., & Abelson, R. P. (1977). *Scripts, Plans, Goals, and Understanding: An Inquiry into Human Knowledge Structures*. Hillsdale: Lawrence Erlbaum.

- Schepman, A., & Rodway, P. (2020). Initial validation of the general attitudes towards Artificial Intelligence Scale. *Computers in Human Behavior Reports*, 1, 100014. <https://doi.org/10.1016/j.chbr.2020.100014>
- Schiff, D. (2021). Out of the laboratory and into the classroom: The future of artificial intelligence in education. *AI & Society*, 36(1), 331-348.
- Schön, D. A. (1983). *The reflective practitioner: How professionals think in action*. Basic Books.
- Schön, D. A. (1987). *Educating the reflective practitioner: Toward a new design for teaching and learning in the professions*. Jossey-Bass.
- Schriesheim, C. A., Powers, K. J., Scandura, T. A., Gardiner, C. C., & Lankau, M. J. (1993). Improving construct measurement in management research: Comments and a quantitative approach for assessing the theoretical content adequacy of paper-and-pencil survey-type instruments. *Journal of Management*, 19(2), 385-417.
- Searle, J. R. (1980). Minds, brains, and programs. *Behavioral and Brain Sciences*, 3(3), 417-424.
- Searle, J. R. (1990). Is the Brain a Digital Computer? *Proceedings of the American Philosophical Association*, 64(3), 21-37.
- Sefton-Green, J. (2007). Literature review of informal learning with technology outside school. Bristol, UK: Futurelab.
- Selwyn, N. (2019). *Should Robots Replace Teachers? AI and the future of education*. Cambridge: Polity Press.
- Selwyn, N. (2022). The future of AI and education: Some cautionary notes. *European Journal of Education*, 57(4), 620-631. <https://doi.org/10.1111/ejed.12532>
- Sfard, A. (1998). On two metaphors for learning and the dangers of choosing just one. *Educational Researcher*, 27(2), 4-13.
- Shadish, W. R., Cook, T. D., & Campbell, D. T. (2002). *Experimental and quasi-experimental designs for generalized causal inference*. Houghton Mifflin.
- Shih, P.-K., Lin, C.-H., Wu, L. Y., & Yu, C.-C. (2021). Learning Ethics in AI—Teaching Non-Engineering Undergraduates through Situated Learning. *Sustainability*, 13(7), 3718. <https://doi.org/10.3390/su13073718>
- Shute, V. J., Sun, C., & Asbell-Clarke, J. (2017). Demystifying computational thinking. *Educational Research Review*, 22, 142-158.
- Siemens, G. (2005). Connectivism: A learning theory for the digital age. *International Journal of Instructional Technology and Distance Learning*, 2(1), 3-10.
- Sindermann, C., Sha, P., Zhou, M., Wernicke, J., Schmitt, H. S., Li, M., & Montag, C. (2021). Assessing the attitude towards artificial intelligence: Introduction of a short measure in German, Chinese, and English Language. *KI – Künstliche Intelligenz*, 35(1), 109–118.
- Singer, J. D., & Willett, J. B. (2003). *Applied longitudinal data analysis: Modeling change and event occurrence*. Oxford University Press.

- Skinner, B. F. (1970). *La tecnologia dell'insegnamento* (L. Magliano, Trad.; C. Scurati, Pref.). La Scuola. (Opera originale pubblicata 1968)
- Skinner, B. F. (1975). *Walden due: Utopia per una nuova società* (E. M. Peron, Trad.; C. Metelli di Lallo, Pres.). Le Lettere. (Opera originale pubblicata 1948)
- Snider, A., & Schnurer, M. (2006). *Many sides: Debate across the curriculum*. IDEBATE Press.
- Somalvico, M. (1992). Intelligenza Artificiale. *Enciclopedia Italiana*, V appendice.
- Steinbauer, G., Kandlhofer, M., Chklovski, T., Heintz, F., & Koenig, S. (2021). A Differentiated Discussion About AI Education K-12. *KI - Künstliche Intelligenz*, 35(2), 131–137. <https://doi.org/10.1007/s13218-021-00724-8>
- Straub, D. W. (1989). Validating instruments in MIS research. *MIS Quarterly*, 13(2), 147–169.
- Strauß, S. (2021). "Don't let me be misunderstood": Critical AI literacy for the constructive use of AI technology. *TATuP - Zeitschrift Für Technikfolgenabschätzung in Theorie Und Praxis*, 30(3), 44–49. <https://doi.org/10.14512/tatup.30.3.44>
- Street, B. (2003). What's "new" in New Literacy Studies? Critical approaches to literacy in theory and practice. *Current issues in comparative education*, 5(2), 77-91.
- Su, J., & Zhong, Y. (2022). Artificial Intelligence (AI) in early childhood education: Curriculum design and future directions. *Computers and Education: Artificial Intelligence*, 3, 100072. <https://doi.org/10.1016/j.caeai.2022.100072>
- Su, J., Ng, D. T. K., & Chu, S. K. W. (2023). Artificial Intelligence (AI) Literacy in Early Childhood Education: The Challenges and Opportunities. *Computers and Education: Artificial Intelligence*, 4, 100124. <https://doi.org/10.1016/j.caeai.2023.100124>
- Su, J., Zhong, Y., & Ng, D. T. K. (2022). A meta-review of literature on educational approaches for teaching AI at the K-12 levels in the Asia-Pacific region. *Computers and Education: Artificial Intelligence*, 3, 100065. <https://doi.org/10.1016/j.caeai.2022.100065>
- Sullivan, G. M., & Feinn, R. (2012). Using effect size—or why the P value is not enough. *Journal of Graduate Medical Education*, 4(3), 279-282.
- Taddia, F., & Mazzolari, B. (2021). *Perché i robot sono stupidi. E tante altre domande sulla robotica?*. Firenze: Editoriale Scienza.
- Tanica, R. (2022). *Non siamo mai stati sulla Terra*. Milano: Mondadori.
- Tavakol, M., & Dennick, R. (2011). Making sense of Cronbach's alpha. *International Journal of Medical Education*, 2, 53-55.
- Teddlie, C., & Tashakkori, A. (2009). *Foundations of mixed methods research: Integrating quantitative and qualitative approaches in the social and behavioral sciences*. Sage Publications.
- Tegmark, M. (2017). *Life 3.0: Being Human in the Age of Artificial Intelligence*. New York: Knopf. (Trad. it. Vita 3.0. Essere umani nell'era dell'intelligenza artificiale, Milano: Raffaello Cortina, 2018).
- Templin, M. C. (1957). Certain language skills in children: Their development and interrelationships. University of Minnesota Press.

- Tomlinson, C. A. (2014). *The differentiated classroom: Responding to the needs of all learners* (2nd ed.). ASCD.
- Tonfoni, G. (1984). *Artificial Intelligence and Literary Creativity: Inside the Mind of Brutus, a Storytelling Machine*. Cambridge: The MIT Press.
- Touretzky, D., Gardner-McCune, C., Martin, F., & Seehorn, D. (2019). Envisioning AI for K-12: What should every child know about AI? *Proceedings of the AAAI Conference on Artificial Intelligence*, 33(1), 9795-9799.
- Turing, A. M. (1950). Computing machinery and intelligence. *Mind*, 59(236), 433-460.
- Turkle, S. (1985). *The Second Self: Computers and the Human Spirit*. New York: Simon & Schuster.
- Turkle, S. (2011). *Alone Together: Why We Expect More from Technology and Less from Each Other*. New York: Basic Books.
- UNESCO (United Nations Educational, Scientific and Cultural Organization). (2022). K-12 AI curricula: A mapping of government-endorsed AI curricula. UNESCO. <https://unesdoc.unesco.org/ark:/48223/pf0000380602>
- UNESCO. (2020). Artificial intelligence and gender equality: Key findings of UNESCO's Global Dialogue. Paris: UNESCO.
- UNESCO. (2021). Recommendation on the Ethics of Artificial Intelligence. Paris: UNESCO.
- United Nations Children's Fund (UNICEF). (2020). *Policy guidance on AI for children Draft 1.0*. <https://www.unicef.org/globalinsight/media/1171/file/UNICEF-Global-Insight-policy-guidance-AI-children-draft-1.0-2020.pdf>
- United Nations Children's Fund (UNICEF). (2021a). *AI policy guidance: How the world responded*. <https://www.unicef.org/globalinsight/stories/ai-policy-guidance-how-world-responded>
- United Nations Children's Fund (UNICEF). (2021b). *Policy guidance on AI for children 2.0*. UNICEF. <https://www.unicef.org/globalinsight/media/2356/file/UNICEF-Global-Insight-policy-guidance-AI-children-2.0-2021.pdf>
- United Nations Educational, Scientific and Cultural Organization (UNESCO). (2019a). *Beijing consensus on artificial intelligence and education*. <https://unesdoc.unesco.org/ark:/48223/pf0000368303>
- United Nations Educational, Scientific and Cultural Organization (UNESCO). (2019b). *Stepping up AI for social good*. United Nations Educational, Scientific and Cultural Organization.
- United Nations Educational, Scientific and Cultural Organization (UNESCO). (2021). *AI and education: Guidance for policy makers*. <https://unesdoc.unesco.org/ark:/48223/pf0000376709>
- Urwin, R. (2019). *Artificial Intelligence: The Quest for the Ultimate Thinking Machine*. London: Arcturus Publishing.
- Van Brummelen, J., Heng, T., & Tabunshchyk, V. (2021). Teaching Tech to Talk: K-12 Conversational Artificial Intelligence Literacy Curriculum and Development Tools. *Proceedings of the AAAI Conference on Artificial Intelligence*, 35(17), 15655–15663. <https://doi.org/10.1609/aaai.v35i17.17844>

- Van Eck, N.J., & Waltman, L. (2007). VOS: a new method for visualizing similarities between objects. In H.-J. Lenz, & R. Decker (Eds.), *Advances in Data Analysis: Proceedings of the 30th Annual Conference of the German Classification Society* (pp. 299-306). Springer
- Van Eck, N.J., & Waltman, L. (2010). Software survey: VOSviewer, a computer program for bibliometric mapping. *Scientometrics*, 84(2), 523-538.
- Varanini, F. (2020). *Le cinque leggi bronzee dell'era digitale. E perché conviene trasgredirle*. Milano: Guerini e Associati.
- Varisco, B. M. (2002). *Socio-Costruttivismo. Genesi filosofiche, sviluppi psicopedagogici, applicazioni didattiche*. Roma: Carocci.
- Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., ... & Polosukhin, I. (2017). Attention is all you need. *Advances in Neural Information Processing Systems*, 30.
- Vazhayil, A., Shetty, R., Bhavani, R. R., & Akshay, N. (2019). Focusing on Teacher Education to Introduce AI in Schools: Perspectives and Illustrative Findings. In *IEEE Tenth International Conference on Technology for Education (T4E)* (pp. 71-77).
- Vygotsky, L. S. (1978). *Mind in society: The development of higher psychological processes*. Harvard University Press.
- Wang, B., Rau, P.-L. P., & Yuan, T. (2022). Measuring user competence in using artificial intelligence: Validity and reliability of artificial intelligence literacy scale. *Behaviour & Information Technology*, 1–14. <https://doi.org/10.1080/0144929X.2022.2072768>
- Wang, P. (2019). On defining artificial intelligence. *Journal of Artificial General Intelligence*, 10(2), 1–37. <https://doi.org/10.1515/jagi-2019-0001>
- Wang, Y. Y., & Wang, Y. S. (2022). Development and validation of an artificial intelligence anxiety scale: An initial application in predicting motivated learning behavior. *Interactive Learning Environments*, 30(4), 619–634. <https://doi.org/10.1080/10494820.2019.1674887>
- Wasserstein, R. L., & Lazar, N. A. (2016). The ASA statement on p-values: context, process, and purpose. *The American Statistician*, 70(2), 129-133.
- Weber, P., Pinski, M., & Baum, L. (2023). Toward an Objective Measurement of AI literacy.
- Wickham, H., & Grolemund, G. (2017). *R for data science: Import, tidy, transform, visualize, and model data*. O'Reilly Media.
- Wiedemann, G. (2022). Computational text analysis within the humanities: How to combine working practices from the humanities and computer science? *Journal of Contemporary Ethnography*, 51(1), 87-109.
- Wilcoxon, F. (1945). Individual comparisons by ranking methods. *Biometrics Bulletin*, 1(6), 80-83.
- William, D. (2011). What is assessment for learning? *Studies in Educational Evaluation*, 37(1), 3-14.
- Williams, R., Park, H. W., & Breazeal, C. (2018). A is for artificial intelligence: the impact of artificial intelligence activities on young children's perceptions of robots. In *Proceedings of the 2018 CHI Conference on Human Factors in Computing Systems* (pp. 1-11).
- Wing, J. M. (2006). Computational Thinking. *Communications of the ACM*, 49(3), 33-35.

- Wing, J. M. (2008). Computational thinking and thinking about computing. *Philosophical Transactions of the Royal Society A: Mathematical, Physical and Engineering Sciences*, 366(1881), 3717-3725.
- Wong, G. K. W., Ma, X., Dillenbourg, P., & Huan, J. (2020). Broadening artificial intelligence education in K-12: Where to start? *ACM Inroads*, 11(1), 20–29. <https://doi.org/10.1145/3381884>
- Wood, D., Bruner, J. S., & Ross, G. (1976). The role of tutoring in problem solving. *Journal of Child Psychology and Psychiatry*, 17(2), 89-100.
- World Economic Forum. (2020). *The Future of Jobs Report 2020*. Geneva: World Economic Forum.
- Wright, K. B. (2005). Researching Internet-based populations: Advantages and disadvantages of online survey research, online questionnaire authoring software packages, and web survey services. *Journal of Computer-Mediated Communication*, 10(3), JCMC1034.
- Yang, W. (2022). Artificial Intelligence education for young children: Why, what, and how in curriculum design and implementation. *Computers and Education: Artificial Intelligence*, 3, 100061. <https://doi.org/10.1016/j.caeai.2022.100061>
- Yi, Y. (2021). Establishing the concept of AI literacy: Focusing on competence and purpose. *JAH*, 12(2), 353–368. <https://doi.org/10.21860/j.12.2.8>
- Zaltman, G., & Coulter, R. H. (1995). Seeing the voice of the customer: Metaphor-based advertising research. *Journal of Advertising Research*, 35(4), 35-51.
- Zawacki-Richter, O., Marín, V. I., Bond, M., & Gouverneur, F. (2019). Systematic review of research on artificial intelligence applications in higher education – where are the educators? *International Journal of Educational Technology in Higher Education*, 16, 39. <https://doi.org/10.1186/s41239-019-0171-0>
- Zhao, L., Wu, X., & Luo, H. (2022). Developing AI literacy for Primary and Middle School Teachers in China: Based on a Structural Equation Modeling Analysis. *Sustainability*, 14(21), 14549. <https://doi.org/10.3390/su142114549>
- Zhou, J., Gandomi, A. H., Chen, F., & Holzinger, A. (2020). Evaluating the quality of machine learning explanations: A survey on methods and metrics. *Electronics*, 10(5), 593.
- Zimmerman, D. W. (1987). Comparative power of Student t test and Mann-Whitney U test for unequal sample sizes and variances. *The Journal of Experimental Education*, 55(3), 171-174.