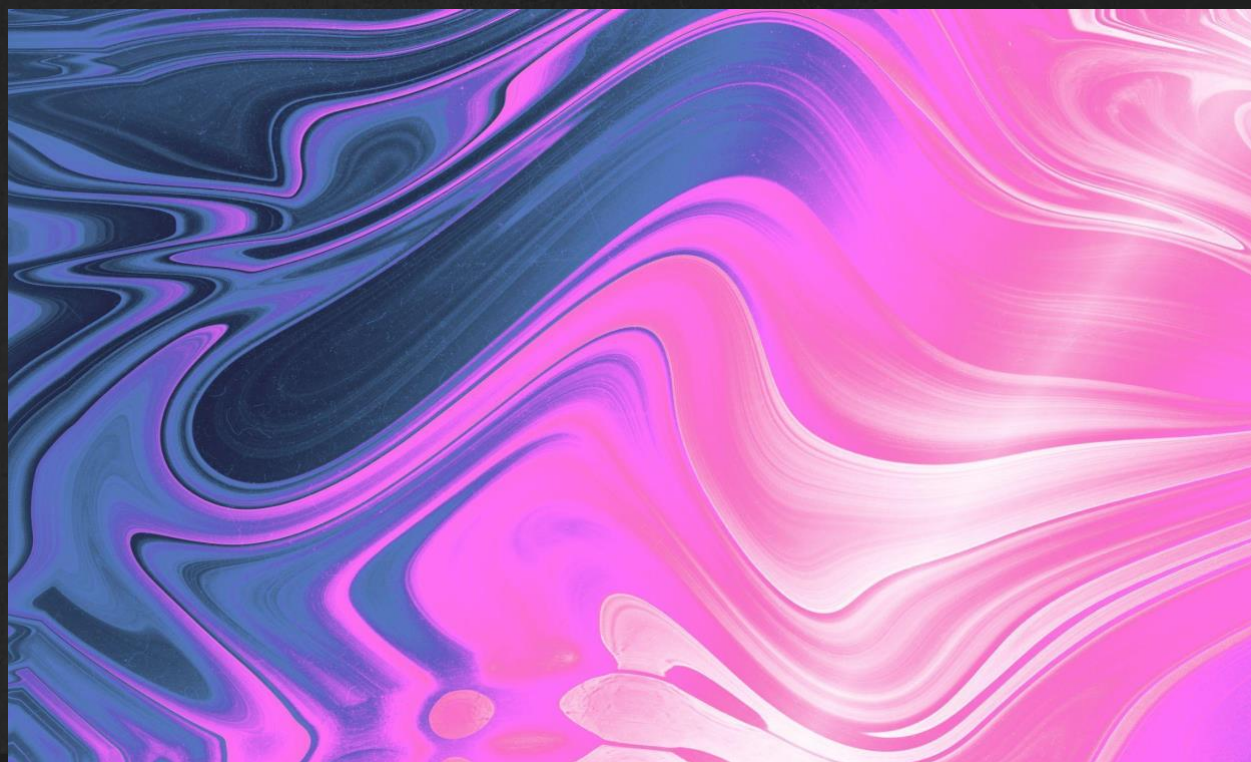




Funded by
the European Union



Valorizzare i Giovani Imprenditori con l'IA"

n° 2024-1-BG01-KA220-YOU-000249027



Funded by
the European Union





INDICE

INTRODUCTION

1: AI AND MACHINE LEARNING

2: IMPORTANCE OF TRAINING DATA

3: AI ARCHITECTURES

3.1: GENETIC ALGORITHMS

3.2: NEURAL NETWORKS

4: TRAINING AI

4.1: FITNESS FUNCTIONS AND REINFORCEMENT LEARNING

4.2: SUPERVISED AND UNSUPERVISED LEARNING

5: HOW TO USE AI EFFECTIVELY

6: RECAP OF KEY TAKEAWAYS



Funded by
the European Union



INTRODUZIONE

In questo modulo vengono trattati i fondamenti dell'Intelligenza Artificiale, fornendo informazioni sui pilastri necessari per creare un modello di IA. Vengono analizzati i colli di bottiglia nello sviluppo di modelli potenti, che comprendono decisioni di design a livello architetturale, nonché le sfumature da considerare durante l'addestramento dei modelli.

Dopo la lettura di questo modulo si dovrebbe avere una comprensione più approfondita dei temi legati all'IA. Questo modulo si concentra principalmente sul lato tecnologico dell'IA, ma suggerisce anche come utilizzare l'IA in modo più efficace.

- Imparerai l'importanza dei dati di addestramento e a cosa prestare attenzione.
- Imparerai le diverse architetture attualmente disponibili, con i loro punti di forza e debolezza.
- Imparerai i modi più produttivi per utilizzare i modelli di IA oggi disponibili.

1: AI E APPRENDIMENTO AUTOMATICO

Con l'evoluzione del settore, anche la percezione del significato di IA è cambiata. Oggi, quando si parla di IA, ci si riferisce principalmente a grandi modelli di apprendimento automatico (machine learning), che è ciò di cui parleremo in questo modulo.

Il concetto di machine learning si riferisce a un algoritmo in grado di individuare schemi nei dati di input e migliorarsi man mano che viene addestrato più estensivamente. La chiave è che non gli viene detto "come pensare". Alcuni modelli funzionano imitando il cervello biologico.

In alcuni casi, IA è solo una parola di moda usata per attirare attenzione e vendere prodotti, quindi è fondamentale comprendere il funzionamento interno di questa tecnologia misteriosa e spesso fraintesa. Comprendendo i concetti chiave, avrai un'idea più chiara di ciò che è possibile e non sarai ingannato dall'hype. Inoltre, sarai meglio attrezzato per sfruttare la potenza che l'IA ha da offrire.

In questo modulo tratteremo la comprensione dei temi chiave del machine learning. Evidenzieremo l'importanza di dati di addestramento di alta qualità. Successivamente analizzeremo esempi selezionati di architetture di IA – dato che ne esistono moltissime – spiegando come funzionano e quali punti di forza e debolezza presentano. In seguito vedremo come i diversi modelli di IA possono essere addestrati e i tipi di addestramento disponibili, quindi combinando architettura e dati di training. Infine analizzeremo lo stato attuale dell'IA e commenteremo come ottenere il massimo dai prodotti disponibili sul mercato.

2: IMPORTANZA DEI DATI DI APPRENDIMENTO

Come nella maggior parte dei sistemi, l'IA è forte solo quanto il suo anello più debole. Se hai buoni dati ma nessun modo di elaborarli, non ottieni nulla; allo stesso modo, anche con l'architettura più avanzata, se la alimenti con dati di bassa qualità, i risultati saranno deludenti.

In questo capitolo discuteremo l'importanza di fornire dati di alta qualità per l'addestramento di un modello di IA e i criteri che i dati devono rispettare per essere considerati di qualità.

Esistono sistemi che utilizzano l'apprendimento simulato, dove un'IA interagisce con un ambiente virtuale e apprende per tentativi ed errori tramite il rinforzo (questo verrà spiegato nel capitolo 4). Per altri modelli di IA, i dati di addestramento rappresentano l'unica fonte di comprensione. Tutti gli schemi osservati dal modello derivano dai dati forniti. Questo perché l'IA, a differenza degli esseri umani, non ha un corpo o esperienze personali da cui trarre conoscenza, ma dipende esclusivamente dagli input per definire la sua "verità". Non ha altri modi per testare le proprie conclusioni, a meno che non le venga successivamente comunicato che l'informazione è errata. Per questo motivo, i dati di alta qualità sono necessari al corretto funzionamento dei modelli di IA.

Ma cosa significa dati di alta qualità?

La prima considerazione è **l'accuratezza**. Esistono dataset etichettati e non etichettati, e l'accuratezza è fondamentale in entrambi i casi. Nei dati etichettati usati con apprendimento supervisionato (si veda il capitolo 4), l'accuratezza dipende da quanto dettagliate e corrette siano le etichette rispetto ai dati.

Esempio: puoi addestrare un'IA a riconoscere oggetti in un'immagine.



Example input image <https://www.rspb.org.uk/birds-and-wildlife/robin>

La prima considerazione è l'accuratezza. Esistono dataset etichettati e non etichettati, e l'accuratezza è fondamentale in entrambi i casi. Nei dati etichettati usati con apprendimento supervisionato (si veda il capitolo 4), l'accuratezza dipende da quanto dettagliate e corrette siano le etichette rispetto ai dati.

Esempio: puoi addestrare un'IA a riconoscere oggetti in un'immagine.

[Immagine di esempio: pettirosso](#)

Questa immagine potrebbe essere usata come input per addestrare un'IA. L'immagine è la stessa, ma le etichette possono variare a seconda della complessità desiderata. Una semplice etichetta potrebbe essere "uccello": questo limiterà la comprensione dell'IA. Un dataset più



avanzato potrebbe etichettare con “uccello, pettirosso, appollaiato, su un ramo, petto arancione, sfondo sfocato” e così via. È facile immaginare la differenza tra i modelli risultanti. Se usi un dataset inaffidabile, questa immagine potrebbe addirittura essere etichettata “gatto”. Naturalmente questo è un esempio estremo, ma se insegni all’IA che si tratta di un gatto, potrebbe confondere gli uccelli con i gatti. Per questo motivo è fondamentale utilizzare dati etichettati accurati e dettagliati.

Nel caso di dati non etichettati usati per l’apprendimento non supervisionato, si può pensare ai Large Language Model (LLM). Modelli come ChatGPT possono essere addestrati su testi non etichettati: articoli, tweet, post su Reddit, riviste scientifiche, ecc. L’IA usa questi dati per identificare schemi nella struttura delle frasi e creare relazioni tra parole. Con sufficiente quantità di dati, l’IA inizia a produrre frasi simili a quelle di un essere umano. Se però l’addestri su compiti scolastici scritti da bambini di 6 anni, otterrai output indesiderati. Ecco cosa significa “qualità” nel contesto di dati non etichettati.

In breve: se insegni cose false, l’IA non ha modo di correggerle. Assicurati che i dati siano di qualità, altrimenti potrebbe scambiare gatti per cani. Se addestri un modello su testi, questi devono essere coerenti e grammaticalmente corretti, altrimenti gli output potrebbero risultare privi di senso. Anche una piccola parte di dati di bassa qualità può contaminare il modello, poiché l’IA lavora su schemi e connessioni. La qualità dell’input è direttamente correlata alla qualità degli output.

Un’altra considerazione riguarda i bias incorporati. Anche se i dati non presentano errori tecnici, possono comunque generare enormi problemi se si trascurano i bias. In uno scandalo di PR, un’IA di riconoscimento facciale, addestrata prevalentemente su volti caucasici, ha identificato persone di altre etnie come primati. Questo problema può essere risolto considerando due ulteriori criteri di qualità: volume e diversità.

- Volume: più grande è la “mente” dell’IA – misurata in numero di parametri – più dati sono necessari per codificare correttamente i



Funded by
the European Union





valori nei nodi. Non esiste una quantità fissa di dati da utilizzare, l'unico modo è testare l'IA e valutarne le prestazioni.

- Diversità: se i dati sono troppo simili, l'IA rischia di "overfittare" (specializzarsi eccessivamente) e non riconoscere schemi generali. Ad esempio, se in tutte le foto di uccelli c'è un cielo blu sullo sfondo, l'IA potrebbe associare il blu alla presenza di un uccello. Nei modelli di linguaggio, se addestri solo su articoli giornalistici, il modello produrrà sempre output "da giornalista", perdendo la varietà necessaria.



Funded by
the European Union



3: ARCHITETTURE DI AI

Sebbene i dati di addestramento siano fondamentali per il funzionamento finale di un modello di IA, non sono l'unico elemento da considerare. Perché i dati risultino utili, è cruciale un'architettura adeguata su cui basare il modello. Pensa all'architettura come a uno scheletro: da solo è solo un insieme di ossa, che ha bisogno di muscoli per funzionare.

Esistono molte tipologie di architettura per scopi diversi. Alcuni modelli sono "stretti" (ANI – *Artificial Narrow Intelligence*), ottimizzati per compiti specifici (ad esempio, le IA per giocare a scacchi). Altri modelli, detti AGI – *Artificial General Intelligence*, cercano di imitare la comprensione umana in più campi. L'AGI non è ancora stato raggiunto e resta oggetto di dibattito. I Large Language Models (LLM), come ChatGPT, sono un tentativo verso l'AGI e rappresentano un enorme progresso rispetto a pochi anni fa.

Ogni architettura presenta punti di forza, limiti e casi d'uso specifici. In questo modulo tratteremo due esempi fondamentali: **algoritmi genetici** e **reti neurali**.

3.1: Algoritmi Genetici

Gli **algoritmi genetici (GA)** sono un tipo di architettura di IA, un po' diversa da quella a cui probabilmente sei abituato nel mainstream. Per ampliare la comprensione delle varie tipologie di IA, vedremo a cosa servono e come funzionano. Si tratta di un'architettura molto specializzata, non ampiamente utile per molti compiti differenti, ma altamente efficiente per problemi specifici. Di solito i GA vengono

utilizzati in **problemi di ottimizzazione** in cui ci sono più variabili di quante se ne possano gestire con la forza bruta. È una soluzione elegante e leggera per testare le possibili soluzioni a problemi rigidamente definiti. Tutto ciò può sembrare un po' astratto, quindi analizziamo l'esempio più famoso: **il problema del commesso viaggiatore**.

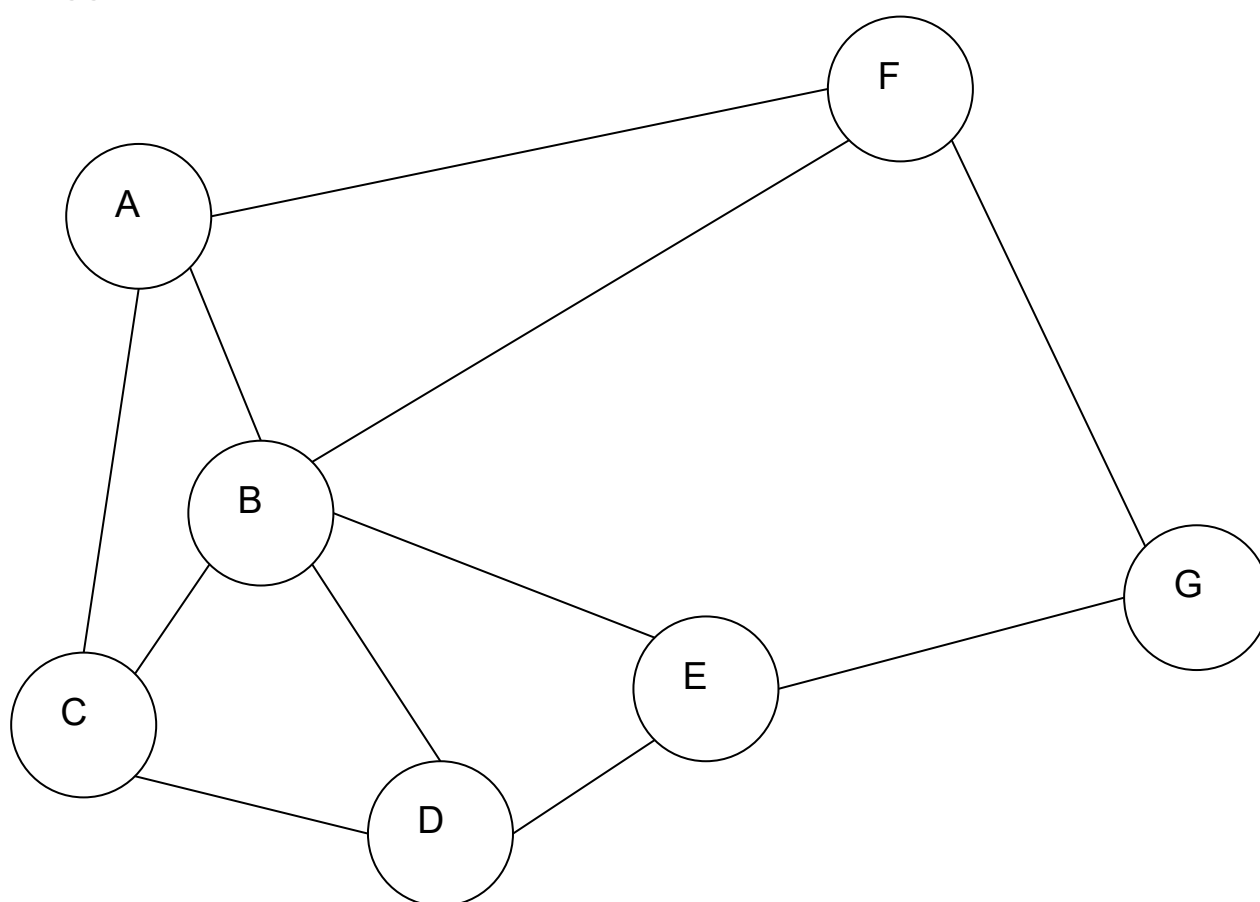


Diagramma di scenario potenziale

Nel disegno sopra si vedono 7 nodi. Il problema del commesso viaggiatore richiede una soluzione in cui ogni nodo o città sia visitato esattamente una volta, con l'obiettivo di trovare il percorso più breve. Per semplicità, in questo esempio non ci sono distanze tra i nodi. Una soluzione potenziale è **ABCDEFG**. Un'altra è **ABCDEGF**. Sebbene esistano altri modi per risolvere il problema, i GA funzionano molto bene in casi di questo tipo.



Come funzionano?

In questo esempio, ogni soluzione è lunga 7 lettere e include ciascuna lettera una sola volta. Pensa a questo come al **DNA** su cui si basano le “genetiche”: la posizione di ciascuna lettera è un attributo della soluzione. Ogni soluzione è un membro della **popolazione**.

Inizialmente viene generata una popolazione casuale con tutti i tipi di soluzioni, alcune inefficienti e altre più veloci. Una **funzione di fitness** è un sistema di valutazione usato per confrontare la qualità delle soluzioni. In questo esempio, la funzione di fitness è semplicemente la distanza del percorso: più è bassa, meglio è. Ciò significa che le soluzioni con distanza totale inferiore saranno classificate meglio e usate per produrre la generazione successiva. Con 3 nodi, se la soluzione è **ABC**, la distanza è **AB + AC**, dove AB è la distanza da A a B.

Una volta generata la popolazione, essa deve essere ordinata in base alla funzione di fitness. Poi, la generazione successiva può essere prodotta dai genitori della generazione attuale. Anche se le soluzioni con ranking più alto hanno più probabilità di riprodursi, c'è un certo grado di **casualità**, per evitare una convergenza prematura e massimi locali. Questo argomento è troppo complesso per essere trattato in dettaglio all'interno di questo corso, ma per chi è curioso sono disponibili risorse aggiuntive.

La **riproduzione** avviene prendendo due genitori e combinandoli per variare la prole. Sebbene esistano molti modi diversi per combinare la “genetica”, l'idea generale è che la nuova popolazione sia in media migliore, poiché le soluzioni meno efficaci di ogni generazione vengono eliminate. Facendo funzionare l'algoritmo per più generazioni, si possono monitorare le migliori soluzioni fino a trovarne una sufficientemente valida per lo scopo, poiché non vi è mai garanzia che sia la soluzione ottimale assoluta.



Funded by
the European Union



3.2: Reti Neurali

Le **reti neurali (NNs)** sono la pietra angolare dell'IA moderna. L'idea è che esse imitino il funzionamento del cervello umano, con i neuroni (cellule cerebrali) e le connessioni tra di essi. Nel contesto dell'IA, ogni neurone è chiamato **nodo** e la sua rappresentazione informatica è una funzione per l'elaborazione dei dati. Un caso semplice è quando il nodo ha un valore compreso tra 0 e 1 e l'input viene moltiplicato per quel valore: questa è una funzione. Le connessioni hanno spesso anch'esse valori, detti **pesi**, che rappresentano l'influenza di un nodo su quello successivo.

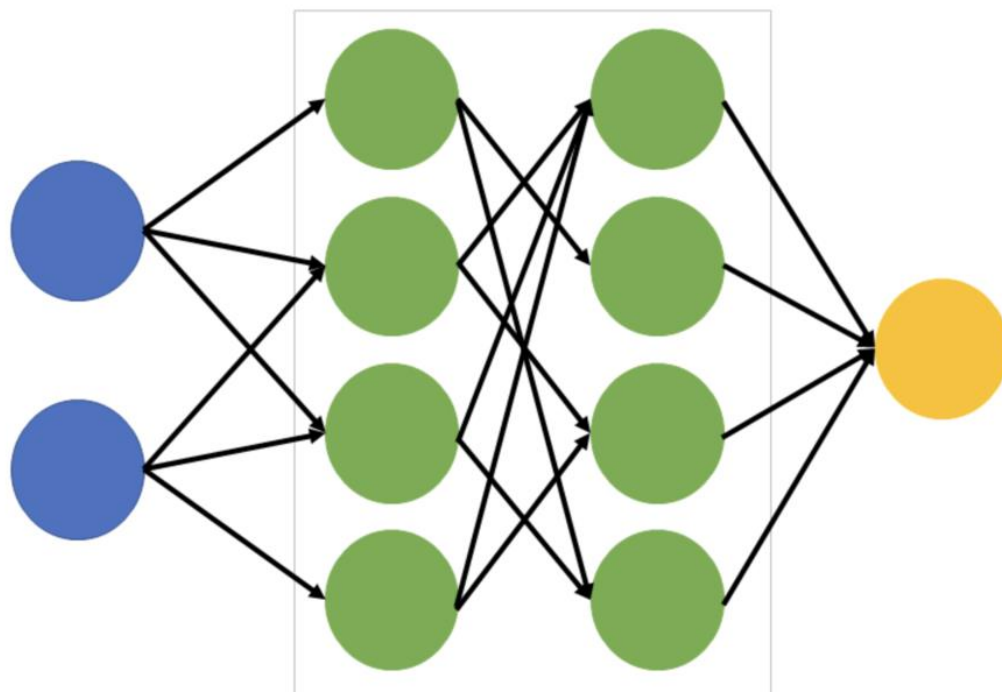
Esistono numerosissime strutture di reti neurali, tra cui l'**architettura transformer** – quella su cui si basano la maggior parte dei LLM (Large Language Models) come ChatGPT. In questo capitolo tratteremo l'idea generale di rete neurale piuttosto che i dettagli delle architetture più avanzate.

Per sviluppare un'intuizione su come funzionano le reti neurali, useremo un esempio visivo di un'IA per il riconoscimento di immagini.

Input layer

Hidden layer

Output layer



Artificial neural networks

Rappresentazione potenziale di una rete neurale

<https://www.sciencelearn.org.nz/images/5156-neural-network-diagram>

In questo esempio, i nodi sono i cerchi e le frecce rappresentano le connessioni. I pesi delle connessioni e i valori dei nodi influenzano entrambi i risultati. Vedremo da dove provengono questi valori nel capitolo 4 – addestrare l'IA.



Immagine con griglia <https://www.photoshopessentials.com/photo-effects/photo-strips/>

Ora immagina una griglia come questa sovrapposta a un'immagine. Invece dei quadrati, tuttavia, ogni **pixel** è un cerchio blu di input nella rete neurale. Questi input vengono quindi elaborati dalla rete: ogni



connessione analizza schemi tra i pixel vicini. Quando l'input attraversa un numero sufficiente di **strati nascosti**, la rete diventa capace di identificare strutture basate sui pixel di input, arrivando quindi a descrivere cosa rappresentano tali strutture.

Le capacità dipendono dalla dimensione della rete neurale, poiché essa può gestire solo tanti input quanti sono i suoi nodi. Puoi pensarla come un puzzle: se ogni nodo sapesse dirti dove va un pezzo, avresti bisogno di tanti nodi quanti sono i pezzi. Questa è una semplificazione: nella realtà, avere più nodi è utile per ottenere maggiore accuratezza e sfumature. Come abbiamo visto nel capitolo precedente, i dati di addestramento sono importanti, ma anche la dimensione della rete neurale rappresenta un collo di bottiglia per le prestazioni. Se avessi un solo nodo, per quanto buoni e numerosi siano i dati, tutto ciò che potresti fare è aggiornare un singolo valore – e alla fine la tua IA non sarebbe utile.

Da quanto visto in questo capitolo dovrebbe essere chiaro che **l'architettura di un modello di IA incide direttamente sulle sue capacità**. Le reti neurali sono più versatili, poiché i dati di input possono essere flessibili se la dimensione della "mente" (numero di parametri) è sufficiente. Gli algoritmi genetici, invece, sono specifici per singoli casi d'uso, a causa della rigidità degli input e degli output, ma hanno comunque grande efficacia se impiegati nel contesto giusto.

Esistono molte altre architetture, alcune molto avanzate e altre più basilari. In generale, i modelli semplici sono più veloci da eseguire e preferiti quando le aspettative sono ben definite. I modelli avanzati, invece, sono costosi da addestrare e richiedono hardware all'avanguardia, ma producono risultati straordinari se progettati correttamente. È sempre fondamentale analizzare quale architettura sia adatta, perché essa determinerà i risultati. Ad esempio, **Stockfish** è una potente IA per gli scacchi, migliore di qualunque essere umano nella storia, mentre **ChatGPT** commette ancora mosse illegali nelle partite.



Funded by
the European Union





Entrambi sono sistemi di IA, ma sono stati addestrati con intenzioni e obiettivi diversi.



Funded by
the European Union



4: ADDESTRARE L'IA

Ora che abbiamo trattato l'importanza di scegliere i dati di addestramento e l'architettura corretti, è il momento di vedere come tutto si combina. L'ultimo passo è il processo di **training dell'IA**. Hai costruito l'edificio, hai comprato i mobili, e ora è il momento di sistemarli. In questo capitolo discuteremo diversi aspetti rilevanti riguardanti l'addestramento dell'IA. Nel contesto degli algoritmi genetici parleremo dell'importanza delle **funzioni di fitness** e collegheremo questo più in generale al **reinforcement learning**. Successivamente analizzeremo la differenza tra **apprendimento supervisionato** e **non supervisionato**.

4.1: Funzioni di Fitness e Reinforcement Learning

Sebbene vi siano differenze tra questi due concetti, vengono trattati insieme perché entrambi prevedono che l'IA impari interagendo con un ambiente. Le funzioni di fitness sono centrali nei GA, poiché le soluzioni devono essere classificate in qualche modo per selezionare i "genitori". È quindi importante stabilire quale sia la funzione di fitness più efficace.

Immaginiamo per un momento che Google Maps utilizzasse un'IA per trovare il percorso più breve. Quale sarebbe la funzione di fitness? La distanza più corta? Il minor consumo di carburante? Il minor tempo per raggiungere la destinazione? Ci sono sempre compromessi: calcolare la distanza è relativamente semplice, ma calcolare il tempo richiede di considerare la velocità massima delle strade, i semafori, il traffico, eventuali chiusure, ecc. Questo potrebbe rallentare la valutazione delle soluzioni e rallentare il processo evolutivo. È comunque fondamentale avere una funzione di fitness appropriata, perché determina i risultati.



Le risposte per “percorso più veloce” e “più efficiente in termini di carburante” potrebbero perfino essere diverse.

Un altro esempio: se stai addestrando un’IA per gli scacchi e definisci la funzione di fitness come “numero di pezzi sulla scacchiera”, potresti comunque venire messo sotto scacco matto pur senza perdere alcun pezzo. Ecco perché è necessario considerare attentamente quale funzione di fitness utilizzare.

Il **reinforcement learning** è collegato a questo concetto, in quanto include funzioni di fitness. È particolarmente diffuso nell’apprendimento simulato – quando l’IA si trova in un ambiente virtuale e tenta di raggiungere un obiettivo. Il principio è che l’IA venga **ricompensata** per azioni corrette e **punita** per azioni sbagliate. Ad esempio, se l’IA è un’auto in un videogioco di corse, riceve più punti se taglia il traguardo più velocemente e perde punti se va nella direzione sbagliata. In questo modo, generazione dopo generazione, vengono mantenute le soluzioni migliori mentre quelle peggiori non si riproducono. Anche in questo caso, la funzione usata per valutare l’IA deve essere attentamente progettata.

L’IA deve avere abbastanza spazio per sperimentare, perché potrebbe trovare soluzioni migliori di quelle che ti aspetti. Se sei troppo rigido con la funzione di fitness, l’IA troverà soltanto la soluzione che desideri, che potresti trovare anche da solo. Nell’esempio della corsa, se dai punti per “stare stretti sulle curve”, sei troppo restrittivo e impedisce di esplorare strategie alternative. Potrebbe infatti risultare più efficace, nell’economia dell’intera gara, affrontare una curva larga mantenendo la velocità.

4.2: Apprendimento Supervisionato e Non Supervisionato



Funded by
the European Union





Gli algoritmi genetici, le funzioni di fitness e il reinforcement learning sono tutti collegati, poiché utilizzano l'interazione con un ambiente per addestrare l'IA tramite una valutazione dei risultati. L'altro paradigma consiste nel **non dire all'IA se sta facendo bene o male**, ma addestrarla su grandi quantità di dati e vedere cosa succede.

Nell'**apprendimento supervisionato**, nonostante il nome, non c'è una persona che supervisiona direttamente il processo. La differenza sta nel tipo di dati forniti all'IA. L'apprendimento supervisionato si basa su **dati etichettati** (solitamente da esseri umani). Significa che qualcuno ha fornito le "risposte corrette", come studiare per un esame utilizzando le soluzioni degli esami precedenti. È utile quando c'è un risultato chiaro e atteso, e produce modelli di IA più prevedibili e meno volatili (con meno variazioni in accuratezza e coerenza).

Nell'**apprendimento non supervisionato**, i dati non sono etichettati e l'IA cerca schemi all'interno dei dati stessi, piuttosto che collegare i dati alle etichette. È utile quando le aspettative sono soggettive, come nella lingua scritta: ci sono molti modi per dire la stessa cosa e nessuno è intrinsecamente sbagliato. In questi casi è poco pratico etichettare i dati ed è più efficace lasciare che l'IA scopra autonomamente schemi e relazioni.

Ora che sappiamo che la differenza tra apprendimento supervisionato e non supervisionato riguarda il tipo di dati forniti all'IA, vediamo le applicazioni di entrambi. Nei modelli avanzati come i LLM (ad esempio ChatGPT) si usano entrambi, per dare all'IA maggiore capacità. Ma a cosa serve ciascun approccio?

La distinzione è importante perché ciascun metodo ha caratteristiche specifiche. L'apprendimento supervisionato può essere costoso, poiché richiede dati etichettati – un processo intensivo considerando la quantità di dati necessaria per addestrare un modello complesso. Ha però il vantaggio di fornire output ben definiti e strutturati. L'apprendimento non supervisionato, invece, può usare qualsiasi dato disponibile e costruire schemi a partire da quello. È meno esplicito e



Funded by
the European Union





strutturato nei risultati, ma più semplice da addestrare, poiché non richiede dati speciali. È particolarmente efficace nell'individuare sfumature e schemi altrimenti poco evidenti.



Funded by
the European Union



5: COME USARE L'IA IN MODO INTELLIGENTE

Ora che hai sviluppato una comprensione di ciò che serve per creare un modello di IA, è il momento di analizzare il panorama attuale dei prodotti esistenti e come possano essere usati in modo più efficace. Non è un segreto che l'uso dell'IA possa potenziare enormemente la produttività: in questo capitolo discuteremo i migliori utilizzi dei LLM e vedremo in quali situazioni l'IA possa risultare controproducente.

I LLM come ChatGPT possono essere utili in molti contesti, ma poiché questo corso è rivolto a giovani imprenditori, ci concentreremo sugli utilizzi più orientati al business.

Brainstorming:

Vantaggi

- Generazione rapida di argomenti correlati.
- Ampia familiarità con molti concetti e capacità di collegare punti affini.
- Permette di confrontarsi con le idee, in modo simile a una conversazione con una persona.

Limitazioni

- Sometimes does not understand what you are expecting so you need to be more specific
- Only trained on available data so if it is something proprietary or outside of the dataset you are out of luck



Scrittura di testi:

Benefici

- Generalmente scrive in modo più coerente della maggior parte delle persone, soprattutto se non madrelingua.
- Ampia familiarità con concetti e capacità di collegare punti affini.
- Può interagire sulle idee, come in un dialogo

Limitazioni

- A volte il tono non è adeguato al tipo di testo.
- Se lavori con qualcuno che non accetta contenuti generati da IA, potrebbe usare scanner di rilevamento.
- Potrebbe richiedere più tempo spiegare contesto ed esigenze che scrivere il testo direttamente.

Dare feedback:

Benefici

- Individua errori che potresti non notare.
- Più veloce rispetto a rileggere da solo, soprattutto testi lunghi.
- Può suggerire miglioramenti una volta individuate le criticità.

Limitazioni

- È programmato per dare una risposta comunque, anche forzata, generando feedback banali.
- A volte è eccessivamente critico rispetto al compito richiesto.
- Tende a usare più parole del necessario, a meno che non venga istruito diversamente.



Funded by
the European Union





Coding:

Benefici

- Generazione rapida di codice.
- Conoscenza di molti linguaggi e framework.
- Può correggere da solo il codice se inizialmente non funziona.
- È in grado di spiegare il codice riga per riga.

Limitazioni

- Non riesce a generare sistemi di file e integrazioni complesse.
- Può far perdere tempo producendo codice che non comprendi.

Considerazioni generali sull'uso dell'IA:

- A volte "allucina" e produce fatti falsi.
- Limitato ai dati di addestramento disponibili.
- Conosce solo le informazioni contestuali che tu fornisci.

I LLM funzionano sulla base della predizione della parola successiva. Anche se i modelli più recenti hanno sistemi più avanzati per citazioni e fornitura di informazioni, è sempre opportuno **verificare i fatti più importanti**.

Se gli argomenti che cerchi non sono prominenti nei dati di addestramento, l'IA non saprà gestirli bene.



Funded by
the European Union





Ad esempio, se devi scrivere un testo per il sito web della tua azienda, devi fornire all'IA molte informazioni sulla tua visione, operazioni e obiettivi. Se queste informazioni non sono già disponibili in un testo copiabile, potrebbe richiedere più tempo raccoglierle che scrivere direttamente il testo. Questo è particolarmente vero per testi brevi con significato molto preciso. In questi casi, può essere utile scrivere tu stesso la bozza e poi chiedere all'IA suggerimenti di miglioramento.



Funded by
the European Union



6: RIEPILOGO DEI PUNTI CHIAVE

Sei arrivato alla fine del modulo, congratulazioni! Rivediamo ora i concetti principali per non dimenticare gli aspetti più importanti.

Dal **capitolo 2** abbiamo appreso perché i dati di addestramento hanno un impatto significativo sugli output di un modello di IA. Abbiamo visto che dati di alta qualità devono:

- Essere accurati.
- Avere quantità sufficiente.
- Essere diversificati.
- Considerare i bias involontari.

Nel **capitolo 3** abbiamo visto come il tipo di architettura sia legato al problema da risolvere. Gli algoritmi genetici sono efficienti per problemi con formato di risultato ben definito, in particolare quelli di ottimizzazione. Abbiamo sviluppato una comprensione di base del funzionamento delle reti neurali (NN) e delle loro numerose varianti. Le NN sono flessibili per scopi e funzionalità.

Nel **capitolo 4** abbiamo compreso come l'apprendimento per rinforzo e le funzioni di fitness siano correlati, come funzionano e quale sia la loro finalità. Abbiamo inoltre distinto l'apprendimento supervisionato da quello non supervisionato.

Infine, nel **capitolo 5**, abbiamo discusso come usare al meglio l'IA e quali considerazioni possano migliorare il flusso di lavoro. Abbiamo visto come l'IA si comporta in attività come **coding, brainstorming, scrittura di testi e revisione**, con punti di forza e debolezza in ciascuna categoria.

Si spera che questa comprensione ti sia utile nel tuo percorso con l'IA. Buona fortuna!

1. BIBLIOGRAFIA

<https://www.oracle.com/nl/artificial-intelligence/ai-model-training/#role-of-data>

<https://www.nytimes.com/2021/09/03/technology/facebook-ai-race-primates.html>

<https://www.geeksforgeeks.org/genetic-algorithms/>

