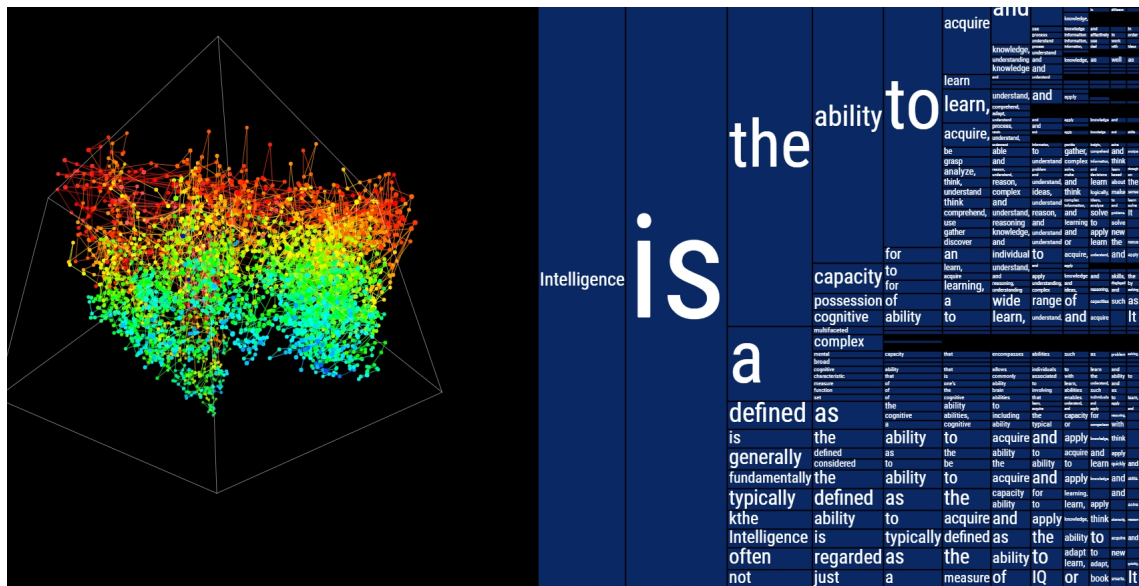




Per spiegare a vostri ragazzi come funziona ChatGPT e quali strutture di tipo statistico e probabilistico caratterizzano i modelli di AI generativa per replicare alle nostre richieste con parole dotate di significato in base alle frasi con cui è stato addestrato, può tornare molto utile il sito che vado a presentarvi.

Moebio.com/mind è un esperimento in cui il creatore mostra l'importanza delle distribuzioni di probabilità sulle parole (o token) per generare testo

In sostanza, ChatGPT viene addestrato su un enorme corpus di dati di testo. Questo processo di addestramento prevede l'apprendimento delle probabilità di sequenze di parole in base alle loro occorrenze e contesti nei dati di addestramento. Il modello non comprende il linguaggio nel senso umano, ma ha una comprensione matematica dei modelli e delle sequenze di parole.

Quando il modello genera testo, lo fa prevedendo la parola successiva (token) in una sequenza basata sulle parole precedenti. Per ogni sequenza di parole, il modello calcola una distribuzione di probabilità su tutte le parole del vocabolario su quale potrebbe essere la parola successiva. Questa probabilità si basa sugli schemi appresi durante l'allenamento.

Il processo di generazione delle risposte è il seguente:

**Inizializzazione:** il processo inizia con un input iniziale (che può essere semplicemente un messaggio effettivo fornito dall'utente).

**Previsione dei token:** ad ogni passaggio, il modello esamina la sequenza generata fino a quel momento e prevede una distribuzione di probabilità per la parola successiva.

**Campionamento:** da questa distribuzione viene selezionata la parola successiva. La selezione può essere deterministica (scegliere la parola più probabile) o comportare una qualche forma di casualità (per generare risultati più vari o creativi).

**Ripetizione:** questo processo si ripete e il modello prende la sequenza estesa (inclusa la parola appena aggiunta) come nuovo input per ogni passaggio successivo, finché non viene soddisfatta una condizione di arresto (come il raggiungimento di una lunghezza massima o la generazione di un token di arresto).

**Perfezionamento e addestramento:** l'addestramento di modelli come ChatGPT prevede la regolazione dei parametri del modello (rilevanza nella rete neurale) in modo che le probabilità assegnate alla parola successiva in una sequenza siano il più accurate possibile, dato il contesto fornito dalle parole precedenti.

La pagina del sito parte dal prompt "l'intelligenza è" e consente di visualizzare le centinaia di possibili risultati, semplicemente spostando il mouse sulla pagina.

[Moebio.com/mind](https://moebio.com/mind)

{jcomments on}