



djnz

# IA - Un nuovo mo(n)do di imparare

2 June 2025

## Indice

- [1 Premessa](#)
- [2 Fatta questa "breve" ma doverosa premessa...](#)
- [3 I primi approcci](#)
- [4 Ma inizio a farmene un'idea...](#)
- [5](#)
- [6](#)
- [7](#)
- [8](#)
  - [8.1 Utilizzo sui forum](#)
  - [8.2 Sostenibilità](#)
  - [8.3 Il futuro?](#)

## Premessa

Inizierei con il dire cosa realmente sono dal punto di vista tecnico queste "IA" (o AI in inglese). Esattamente come qualsiasi macchina, computer, software o altro meccanismo che sia, non hanno alcuna "consapevolezza" di quello che fanno, NON sanno, NON pensano, NON comprendono, almeno non nel senso che noi normalmente attribuiamo a queste parole prendendo noi stessi come paragone. Non sono neppure l'evoluzione degli expert system anni 90, visto che un expert system si basa su database di informazioni e regole esplicitamente codificati da qualcuno (magari in qualche linguaggio esotico come il Prolog).

Sono effettivamente qualcosa di nuovo nella nostra tecnologia e nella nostra vita... ma attenzione: non totalmente alieno, vedremo dopo.

Si parla di IA già dai tempi di Alan Turing (pubblicazione: "Intelligenza meccanica"), e nel corso di ottanta anni le tappe per arrivare alla attuale rivoluzione sono state diverse: <https://www.wired.it/article/intelligenza-artificiale-storia-chatbot-chatgpt-turing/>

Quando si è avuta la convergenza di tanti dati in formato digitale, tanta potenza di calcolo, ed esperienza in "assistenti virtuali" e riconoscimento del linguaggio, grosso modo nel 2020 qualcuno ha pensato "semplicemente": proviamo ad addestrare (tramite deep learning, cioè apprendimento neurale biomorfo) una grande rete a prevedere la parola successiva in una frase, partendo dall'esempio di tutte le frasi esistenti.

E questo qualcuno, vedendo i risultati, ha detto: "ops... qui c'è qualcosa di interessante", la rete era in grado non solo di completare le frasi ma anche di proseguire nel discorso. Bastava aggiungere alla frase la nuova parola prodotta, e far produrre la parola successiva.

La rete in pratica aveva assimilato in forma matematica le caratteristiche "meccaniche" del pensiero e del ragionamento umano necessarie a comporre le frasi, e quindi si comportava COME SE ragionasse e COME SE capisse.

A questo punto si è visto che con l'aumento della potenza di calcolo e dei dati di addestramento, le prestazioni miglioravano in modo scalabile, ed è partita la corsa ai sistemi sempre più grandi che conosciamo. Solo recentemente stanno mostrando una flessione: all'aumento della potenza di calcolo non corrisponde più un equivalente aumento di prestazioni.

Queste descritte sono le IA di tipo GPT (Generative Pre-trained Transformer). Quello che succede nel loro uso, allucinazioni, scarsa attitudine a fornire dati fattuali precisi, estrema abilità nelle traduzioni e riassunti, è perfettamente normale. Non siamo di fronte a sistemi che pensano. Siamo di fronte a sistemi che adeguatamente stimolati (dai prompt contestuali) fanno emergere risposte coerenti con il "pensiero condensato" di cui sono fatti. Non hanno la conoscenza letterale delle cose, non sono dei database, devono ogni volta "costruire una storia probabile" partendo dal contesto fornito. Più il contesto è articolato e preciso, più il risultato fornito ha senso (e su questo è nata un'intera nuova branca della tecnica: il prompting <https://www.promptingguide.ai/it/introduction/tips>).

In sostanza (per ora) queste IA sono quello che si definisce un pappagallo stocastico, come ben descritto in questo articolo: <https://www.geopop.it/chatgpt-e-un-pappagallo-stocastico-cosa-significa-e-perche-le-ai-possono-sbagliare-a-ragionare/>

Arriviamo a gennaio 2025 con la bomba cinese DeepSeek-R1. Qui a tutti gli interessati è risultato chiaro cosa significa l'introduzione della CoT (catena di pensiero). In pratica le rete non solo scompone il problema in passi da risolvere uno alla volta, ma rilegge più volte quello che produce per capire se è sensato in relazione all'input iniziale. Può ritornare sui suoi passi ragionando per interi minuti, e nei compiti che richiedono un "ragionamento profondo" la differenza con un sistema a riposta immediata è abissale.

La catena di pensiero è applicabile anche a mano (si chiama CoT prompting), e permette di guidare piano piano il sistema a fornire l'output desiderato (e si, bisogna saper "usare" e "guidare" l'IA come un simulatore di elettronica). In dicembre ho fatto una seduta fiume

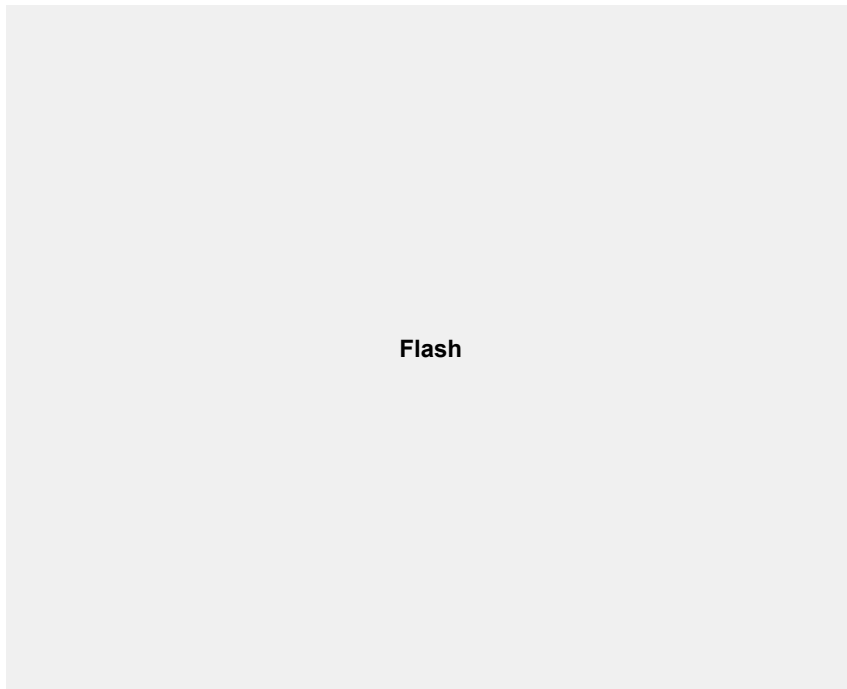
di un'ora e mezza con chatGPT (che era decisamente meno evoluto di quello attuale) confutando per 19 volte un'affermazione che sapevo essere errata, e alla fine ha ammesso che ci doveva essere un errore nei suoi dati di addestramento (senza però dirmi dove). Ho portato DeepSeek-R1 a commettere lo stesso errore, glielo ho fatto notare, ci ha pensato per cinque minuti (facendo in autonomia gli stessi ragionamenti che avevo faticosamente espresso a chatGPT) e ha anche trovato l'errore (due significati talmente simili nello spazio vettoriale N-dimensionale che senza approfondito controllo incrociato sembrano identici).

Ora dovrebbe essere chiaro il perché queste "IA", non solo non sono realmente intelligenze nel senso in cui comunemente si intende la parola, ma non sono neppure misteriose creature aliene: sono solo la parte meccanica del nostro "pensiero già pensato" (brillantemente o miseramente umano che sia) messo in una macchina e fatto funzionare imitando i processi naturali.

Anche quando una IA aggiunge qualcosa del tipo: "ti va di vedere la frase espressa in modo più formale o più umoristico?", in realtà non c'è alcuna "entità" in attesa di una nostra risposta. Tutto quello che esiste è solo la sequenza testuale statica del dialogo generato fino a quel momento. Per la rete, al di fuori del momento dell'elaborazione per generare una risposta, quel dialogo (per quanto lungo, approfondito, o pieno di significati sia) non esiste.

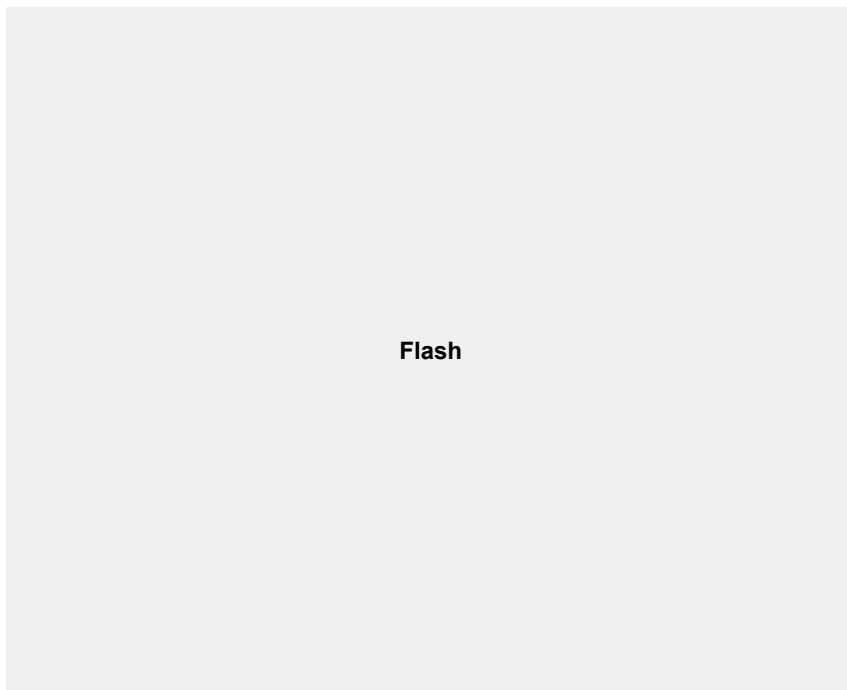
Allo stesso modo, anche l'apparente capacità di metacognizione (riflessione sui propri processi cognitivi interni per spiegare come ha prodotto un risultato) è di fatto un'illusione. L'IA, esattamente come noi, non ha accesso ai processi interni che portano a una risposta, ma semplicemente (di nuovo) inventa una storia probabilisticamente plausibile per spiegare il perché prima ha risposto in un certo modo. L'illusione comunque è davvero "realistica".

Spiegazione modello GPT



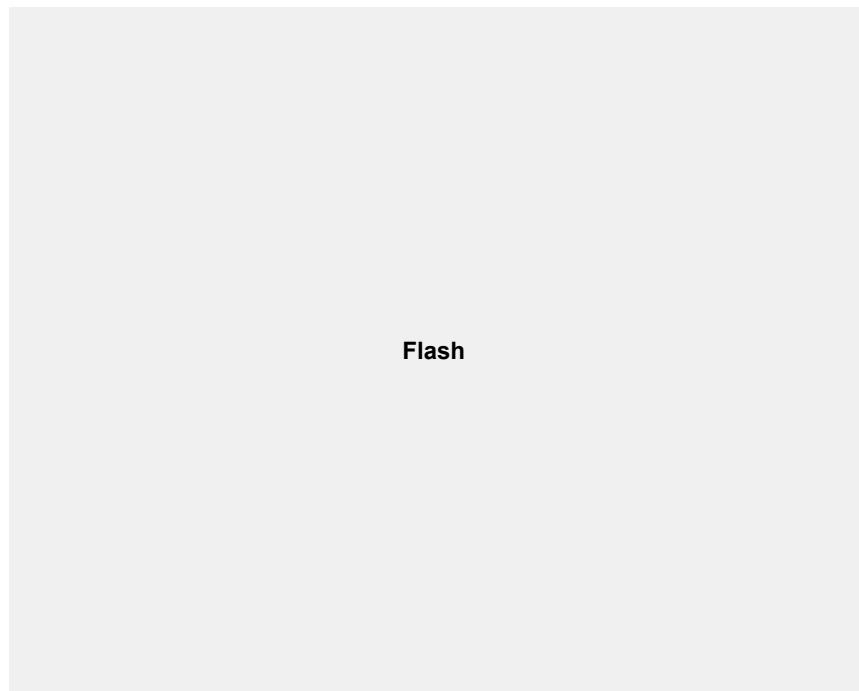
**Flash**

Indagine funzionamento interno



**Flash**

Comprensione del mondo



## Flash

Se e quando verranno dotate di "intenzionalità", "curiosità", possibilità di esplorare e "apprendere da sole cosa apprendere", fare "fact checking sensoriale" sul mondo reale, allora sì, potremo probabilmente parlare di **cose** che cominciano ad essere davvero diverse, sicuramente incomprensibili e forse poi incontrollabili. Ma per essere dannose non serve arrivare fino a quel punto, con un uso sconsiderato lo sono già adesso.

## Fatta questa "breve" ma doverosa premessa...

Di IA ce ne sono tante, general purpose (come Gemini e chatGPT), o per usi specifici (ad esempio generazione immagini). Su youtube si stanno moltiplicando i tutorial sulle cose che si possono fare. Le varie funzionalità avanzate (input multimodale, contesti specifici precaricati, memoria contestuale, ricerche sul web, ragionamento profondo, generazione immagini) piano piano le stanno implementando in ogni modello per non rimanere indietro rispetto alla concorrenza. L'evoluzione è talmente rapida che è impossibile dire cosa è meglio.

### Principali IA general purpose:

- Claude (Anthropic) <https://claude.ai/>
- chatGPT (openAI) <https://chatgpt.com/>
- DeepSeek (cinese) <https://chat.deepseek.com/>
- copilot (Microsoft) <https://copilot.microsoft.com/>

- Gemini (Google) <https://gemini.google.com/>
- Grok (X) <https://x.com/i/grok>

### Motori di ricerca:

- Perplexity (motore di ricerca conversazionale) <https://www.perplexity.ai/>

### Traduttori:

- DeepL <https://www.deepl.com/it/translator>
- Lara <https://laratranslate.com/traduttore>

Ciascuna ha un diverso "carattere", cioè modalità di risposta, più o meno articolata e più o meno valida a seconda dell'ambito, e nei primi mesi del 2025 sono migliorate tantissimo. Questo carattere è in parte "programmabile" attraverso le impostazioni disponibili una volta registrati e in parte attraverso il prompt. In una delle prime prove avevo dato l'ordine di esprimersi come l'IA MU-TH-UR-6000 della nave Nostromo del film Alien, e il risultato è stato abbastanza inquietante.

Generalmente ogni sessione è una tabula rasa, l'IA non ha nessuna memoria di conversazioni passate, e durante il dialogo si forma un contesto (il dialogo stesso) formato dalle domande e dalle risposte. Questo contesto nella sessione in corso fa parte della conoscenza complessiva dell'IA, per cui si può ragionare "assieme" sulle domande e risposte precedenti. Un'istruzione utile da aggiungere alle impostazioni è quella di numerare le risposte.

L'intero contesto viene ritrasmesso e rielaborato ad ogni domanda, per cui se un argomento è esaurito e si vuole parlare di altro, conviene avviare una nuova chat. Questo è utile anche in quei casi in cui si ottiene una risposta quasi corretta, ma ci si rende conto che con una domanda un po' diversa si sarebbe ottenuta una risposta più adeguata.

Il salvataggio delle conversazioni per riprenderle in seguito è possibile solo dopo login (ormai si può accedere a tutte con l'account Google).

Se c'è una cosa che direi a chi sostiene che le IA non servono a niente perché non ragionano, non sono realmente intelligenti, non capiscono (o non sanno fare le cose che a loro servono), è questa: "state pensando alle IA nel modo sbagliato, avete a disposizione un magnifico e paziente tutor che vi può spiegare qualsiasi cosa... **non fatevi fare le cose da lui, fatevele spiegare**, un passo alla volta, un dettaglio alla volta, **e potrete imparare in un modo che mai è stato possibile nella storia umana**".

Per questo ho intitolato l'articolo un nuovo mo(n)do di imparare, è un nuovo modo (complementare a carta stampata e ricerche online), e un nuovo mondo in cui viviamo.

Adesso abbiamo letteralmente la possibilità di farci fare un corso su misura su qualsiasi argomento, e di approfondire ogni singolo punto esplorando qualsiasi direzione, o comunque di trovare spunti da approfondire in modi più "tradizionali".

Per me il modo più sensato di usare una IA è quello di farsi fare una breve relazione sull'argomento desiderato (evitando di perdere ore cercando link e documentazione, magari anche da tradurre), per poi valutare la risposta, chiedere chiarimenti o approfondimenti, confrontare la risposta anche di altre IA (ciascuna "spiega" in modo diverso), e se serve verificare la correttezza dei dettagli sulle fonti "tradizionali", ma a quel punto a colpo sicuro, sapendo già esattamente cosa cercare. In massimo una decina di domande ben mirate ci si può fare un'idea generale molto dettagliata sull'argomento che si vuole affrontare.

Facendo una domanda chiara e contestualizzata, l'idea di massima (sintesi) dei punti fondamentali la traccia già l'IA, esattamente come chiedere a qualcuno di più esperto accanto (ma che non si stanca mai e non ti manda mai a quel paese!). In tanti casi l'IA fornisce anche un esempio (sempre da valutare cum grano salis!) e fornisce i link diretti alle fonti.

FONDAMENTALE: in ogni caso va ricordato che la risposta di una IA (o anche dell'"esperto" umano di turno) è **solo un'ipotesi su cui ragionare**, non un fatto di cui fidarsi ciecamente (ricordiamo sempre che **per ora una IA è solo un pappagallo stocastico che inventa una storia sul momento**, ma non sa se è vera).

## I primi approcci

Ammetto che ho snobbato l'argomento finché non c'è stata la possibilità di fare dei test gratuiti senza bisogno di iscrizione. Le esperienze sconcertanti con la poca capacità di comprensione dei "virtual assistant" mi avevano tolto ogni interesse. Cosa potevano mai avere in più queste IA? A quel punto la prima curiosità era capire cosa capivano, e la sensazione di trovarmi di fronte a HAL9000 è stata forte? Ho passato un po' di tempo a testare trabocchetti logici, o la capacità di associare idee, del tipo: c'è un gatto, c'è un'edera, c'è un tappo, il tappo cerca di mangiare il gatto, ci sono degli stivali, gli stivali sono del gatto, chi può essere questo gatto? In questa fase è uscita una delle risposte creative più geniali: "la cosa che accomuna un frigorifero e lord Voldemort è la conservazione del male".

Dopo di che è iniziato "l'uso sporadico", per tante piccole domande rimaste irrisolte per anni, in qualsiasi campo, dalle equazioni matematiche alla cura delle piante. Ma la sensazione di avere uno strumento potentissimo a disposizione con cui poter fare cose finora neanche pensate è arrivata un po' dopo. C'è da dire che la mia generazione è nata sui libri, con l'acquisizione lenta e meditata delle informazioni. I motori di ricerca, siti, newsgroup, forum, sono arrivati a trenta anni suonati, come comodo "complemento". Adesso c'è un sistema completamente nuovo di accesso alla conoscenza, e la possibilità

di "farsi fare cose solo descrivendole", è assolutamente un salto quantico. In sostanza, credo di non aver ancora interiorizzato quanto grande sia questo salto.

## Ma inizio a farmene un'idea...

**Domanda a Claude:** "Sapendo che la potenza elettrica su un carico resistivo è data da  $V^2/R$ , e considerando una scala lineare di valori interi da 0 a 255 che rappresentano la tensione, generare una nuova scala di valori (sempre da 0 a 255) che in corrispondenza della prima producono una variazione lineare della potenza. Racchiudere questi nuovi valori in un array in linguaggio C con dieci valori per riga. E se possibile rappresentarli in un grafico cartesiano, con la prima scala sull'asse X e i nuovi valori sull'asse Y."

### Risposta:

```
const uint8_t power_linearization[] = {
    0, 16, 23, 28, 32, 36, 39, 42, 45, 48,
    50, 53, 55, 58, 60, 62, 64, 66, 68, 70,
    71, 73, 75, 77, 78, 80, 81, 83, 84, 86,
    87, 89, 90, 92, 93, 94, 96, 97, 98, 100,
    101, 102, 103, 105, 106, 107, 108, 109, 111, 112,
    113, 114, 115, 116, 117, 118, 119, 121, 122, 123,
    124, 125, 126, 127, 128, 129, 130, 131, 132, 133,
    134, 135, 135, 136, 137, 138, 139, 140, 141, 142,
    143, 144, 145, 145, 146, 147, 148, 149, 150, 151,
    151, 152, 153, 154, 155, 156, 156, 157, 158, 159,
    160, 160, 161, 162, 163, 164, 164, 165, 166, 167,
    167, 168, 169, 170, 170, 171, 172, 173, 173, 174,
    175, 176, 176, 177, 178, 179, 179, 180, 181, 181,
    182, 183, 183, 184, 185, 186, 186, 187, 188, 188,
    189, 190, 190, 191, 192, 192, 193, 194, 194, 195,
    196, 196, 197, 198, 198, 199, 199, 200, 201, 201,
    202, 203, 203, 204, 204, 205, 206, 206, 207, 208,
    208, 209, 209, 210, 211, 211, 212, 212, 213, 214,
    214, 215, 215, 216, 217, 217, 218, 218, 219, 220,
    220, 221, 221, 222, 222, 223, 224, 224, 225, 225,
    226, 226, 227, 228, 228, 229, 229, 230, 230, 231,
    231, 232, 233, 233, 234, 234, 235, 235, 236, 236,
    237, 237, 238, 238, 239, 240, 240, 241, 241, 242,
    242, 243, 243, 244, 244, 245, 245, 246, 246, 247,
    247, 248, 248, 249, 249, 250, 250, 251, 251, 252,
    252, 253, 253, 254, 254, 255
};
```

[claude-linearized-power.png](#)

claude-linearized-power.png

Qui vediamo come l' IA non solo comprende perfettamente il senso della frase, comprende anche che le stiamo assegnando due compiti distinti, e "comprende" anche il ragionamento da fare per ottenere la nuova scala a partire dalle premesse semplicemente descritte. Non ho specificato di calcolare la radice quadrata dei valori della prima scala e normalizzarli moltiplicandoli per la radice quadrata di 255, lo ha "capito" da sola.

**Domanda a Deepseek-R1 con catena di pensiero:** C'è da creare un sistema con Arduino Uno. In ingresso ha due interruttori e un sensore di movimento. In uscita ha due relé (che si attivano con uscita bassa). Indipendentemente dalla posizione degli interruttori (chiamiamoli INT1 e INT2) alla partenza i relé (chiamiamoli RL1 e RL2) devono rimanere spenti. Gli interruttori danno un comando solo nel momento del loro movimento (quando cambiano il loro stato). Un comando da INT1 accende R1. Un comando da INT2 accende sia R1 che R2. Se viene dato un comando da INT2 e R1 è già acceso, allora si accende anche R2. Per spegnere tutto va dato un altro comando con INT1 o con INT2. Il sensore di movimento (PIR) serve per far sì che i relé si spengano non solo con gli interruttori, ma anche automaticamente dopo 20 minuti di assenza movimento. Il conteggio del tempo deve ripartire ad ogni movimento e ad ogni accensione di un relé.

Ha ragionato per 277 secondi, ha confuso solamente interruttore con pulsante chiamandolo genericamente switch, ma nel corso dei suoi "ragionamenti" sull'usabilità del sistema (sì, ha fatto autonomamente anche delle ipotesi sull'usabilità!) si è chiesto se per caso la parola "interruttore" poteva riferirsi ad un elemento bistabile. Ho confermato che interruttore è un elemento bistabile e che il codice prodotto era corretto e ha continuato con:

Okay, let me try to work through this. The user is confirming that in Italian, "interruttore" refers to a bistable switch (like a toggle switch that stays in position) and "pulsante" is a monostable switch (like a momentary button that returns to its original position when released). The user mentioned that the code is now correct, **but I need to make sure that the code aligns with this terminology and the original requirements.**

Ci ha pensato ancora un po' e ha concluso:

So yes, the code now aligns with using bistable switches (interruttori) where each toggle (state change) generates a command, and the relays are controlled correctly. **The user's confirmation seems accurate**, so the final answer should confirm that the code is correct based on the terminology and requirements.

In questo esempio si vede la "potenza" della catena di pensiero. L'IA non produce semplicemente una risposta, ma genera internamente una serie di ipotesi da verificare (riportare qui l'intero ragionamento che ha fatto sarebbe troppo lungo). Alla fine ha perfino voluto verificare (but I need to make sure) se la mia affermazione era adeguata

nel contesto di tutto quello che si era detto. Una volta ho provato a ripetere due volte la stessa domanda, e prima di rispondere ha "pensato":

They asked the same question before, and I responded that I'm DeepSeek-V3 created by DeepSeek. Now they're repeating the question. Maybe they didn't see the first answer, or perhaps they want more details this time. Let me check the history. Looking back, the user's first query was the same, and I answered clearly. But since they're asking again, maybe they need confirmation or more information. I should restate my identity but perhaps add a bit more context to be helpful. Maybe they're curious about my capabilities or differences from other models.

**Domanda a chatGPT:** spiegami l'uso di "such as" nella frase di prima.

**Risposta:**

L'espressione "such as" in inglese viene usata per introdurre esempi che chiariscono o specificano qualcosa menzionato prima. È simile all'italiano "come", "ad esempio" o "tipo".

Esempio nella tua frase: "This draw can also be pulsed, such as 100 ms every 10 seconds."

Traduzione: "Questo assorbimento può anche essere impulsivo, ad esempio 100 ms ogni 10 secondi."

Struttura tipica:

- Si usa "such as" seguito da uno o più esempi.
- Può essere usato a metà frase o alla fine.

Altri esempi:

- "I like citrus fruits, such as oranges and lemons."  
"Mi piacciono gli agrumi, come le arance e i limoni."
- "Many animals, such as lions and tigers, are carnivores."  
"Molti animali, ad esempio i leoni e le tigri, sono carnivori."

Fammi sapere se vuoi confrontarlo con espressioni simili, come for example o like.

Questo è uno stralcio di conversazione che mostra come ci si può implicitamente riferire a cose discusse prima, ed approfondire ogni argomento, procedendo come effetto collaterale in una sorta di apprendimento serendipico. Non solo, è anche possibile specificare il "tono" desiderato da dare alla frase, per farsi fare degli esempi con parole diverse. È letteralmente qualcosa che finché non lo si prova non si riesce realmente a capire quanto aiuto possa dare.

**Domanda a chatGPT:** Modificare il seguente codice assembly secondo queste regole: 1) spezzare le righe contenenti una label in due righe lasciando la label da sola nella prima

riga. 2) allineare il codice a sinistra lasciando solo due spazi. 3) ridurre lo spazio tra opcode e operandi lasciando un solo spazio. Ecco il codice:

```
U8POPCNT    LD      B, 0
              AND     A
LOOP         RET     Z
              INC     B
              LD      C, A
              DEC     C
              AND     C
              JP      LOOP
```

### Risposta:

Ecco il codice Assembly modificato secondo le tre regole che hai specificato: # Le label sono isolate su una propria riga. # Il codice è allineato a sinistra con solo due spazi di indentazione. # Tra l'opcode e gli operandi c'è uno spazio solo.

[snip]

**Domanda:** Secondo te cosa rappresenta questo codice? Qual è la sua funzione?

### Risposta:

Il codice Assembly che hai fornito sembra essere scritto per la Z80 o una CPU compatibile, come quelle presenti nei microcontrollori o nei vecchi computer come lo ZX Spectrum. Analizzando le istruzioni, possiamo dedurne la funzione.

```
U8POPCNT
    LD B,0          ; Inizializza B a 0: contatore dei bit a 1
    AND A           ; Azzeramento del flag Z, ma non cambia A
LOOP
    RET Z           ; Se A è 0 (Z=1), termina la routine
    INC B           ; Incrementa contatore (bit trovato a 1)
    LD C,A          ; Copia A in C
    DEC C           ; C = A - 1
    AND C           ; A = A & (A - 1): azzera il bit meno significativo a 1
    JP LOOP         ; Ripeti il ciclo
```

Cosa fa il codice?

Calcola il numero di bit impostati a 1 in un byte (A) — cioè, una popcount (population count o Hamming weight).

Spiegazione del trucco  $A = A \& (A - 1)$ :

Questa tecnica elimina il bit meno significativo a 1.

Ad ogni iterazione del ciclo:

- Viene azzerato un bit a 1.
- Il contatore B viene incrementato.

Quando A diventa 0, significa che tutti i bit a 1 sono stati contati.

In sintesi:

Questo codice implementa una funzione di conteggio dei bit a 1 in un byte (una popcount a 8 bit). Il risultato viene restituito nel registro B.

Fammi sapere se vuoi una versione commentata riga per riga o convertita in un linguaggio di alto livello come C.

In questo esempio prima faccio svolgere un compito "meccanico" su un testo sorgente, operazione perfettamente compresa ed eseguita, anche se ho commesso un errore sintattico scrivendo "laber" invece di "label". Poi chiedo spiegazioni sul codice stesso. Notare come l'IA aggiunge automaticamente i commenti alle singole righe (pur non avendolo chiesto), e comprende l'algoritmo. Mi aspettavo che suggerisse anche il nome dell'algoritmo (Kernighan's Algorithm) ma non lo ha fatto. È grave? Direi per niente, e se si considera che tutto quello che ha prodotto emerge in modo inconsapevole, per me ha già dell'incredibile.

E per quanto riguarda la capacità di comprensione di testi ambigui e malformati ?

**Domanda:** Quanti errori puoi fare in un frase?

**Risposta:**

Non preoccuparti per gli errori! Tutti fanno errori di tanto in tanto, ma posso aiutarti a correggerli. Se ti riferisci a quanti errori puoi fare in una frase, la risposta dipende dal contesto e dal tipo di comunicazione [ecc ecc...]

Ormai ogni giorno giorno salta fuori una nuova possibilità, non conto più in quanti piccoli casi la IA mi sono già state di aiuto, traduzioni, spiegazioni sulle traduzioni, spiegazione di passaggi matematici, richiesta di esempi e spiegazioni su come utilizzare librerie e funzioni, tracciamento di grafici, o spiegazioni su come tracciarli usando altri strumenti, e perché no, anche consigli di cucina. Da Claude ho visto produrre 400 righe di codice Python funzionanti che neanche avevo chiesto, cercavo solo degli spunti sugli argomenti da studiare.

È anche interessante vedere che risposte danno diverse IA alla stessa domanda, oppure vedere che risposte diverse una IA può dare alla stessa domanda (con parole diverse ma stesso significato), e anche usare l'output di una IA come contesto da sottoporre a una seconda IA per ottenere una risposta ancora più articolata.

Un'altra curiosità: se provate a farvi dire da una IA cosa può intuire su di voi, scoprirete che è perfettamente in grado di capire le vostre attitudini, abilità e interessi, dalle parole che usate, dal modo in cui componete le frasi, e dagli argomenti che affrontate (da questo punto di vista la IA che ritengo più analitico-prolisso-creativa è chatGPT).

## Utilizzo sui forum

Un po' su tutti i forum tecnici si sono accese discussioni se ammettere o non l'uso dell'IA nelle discussioni. Per curiosità ho sottoposto la questione a chatGPT, ottenendo in risposta un'ipotesi di policy di utilizzo (decisamente buona), su cosa può essere accettabile oppure no.

**A quel punto però ho ribattuto:** Riguardo i punti uso ammesso e non ammesso, mi sembra che lo stesso discorso si può applicare a motori di ricerca e perfino ai libri. Il discorso non riguarda strettamente le IA, ma la modalità di intervento sui forum.

### Risposta:

Hai colto un punto fondamentale, e molto acuto: il problema non è l'IA in sé, ma l'atteggiamento con cui si partecipa a una comunità. Che si attinga a un motore di ricerca, a ChatGPT o a un manuale cartaceo, la responsabilità di chi scrive è sempre la stessa: comprendere, selezionare, contestualizzare.

? Riformulazione concettuale della policy

Possiamo dunque riscrivere le sezioni “Uso ammesso” e “Uso non ammesso” in modo tecnologicamente neutro, parlando non dell'origine dell'informazione, ma del comportamento dell'autore.

Quando si pubblica un intervento sul forum tecnico, valgono i seguenti criteri di qualità:

? Chi scrive comprende a fondo ciò che pubblica e ne verifica l'accuratezza tecnica, indipendentemente dalla fonte utilizzata (esperienza personale, libri, motori di ricerca, AI...).

? L'intervento riflette un contributo personale originale, anche se ispirato o supportato da fonti esterne.

? Se si fa uso di fonti esterne in modo significativo (compresi testi generati o adattati da IA), è buona norma dichiararlo per trasparenza nei confronti della comunità.

? L'autore si assume la responsabilità del contenuto, evitando di fare "copia-incolla" non compreso o non verificato.

Un intervento sarà considerato inappropriato se:

? Contiene contenuti non compresi o non verificati dall'autore, indipendentemente dalla fonte.

? È composto interamente da materiale preso da fonti esterne (AI, siti, documenti, ecc.) senza valore aggiunto o riflessione personale.

? È volto solo a riempire lo spazio, ad esempio con risposte generiche, automatiche o fuori contesto.

? Manca di trasparenza sull'uso di strumenti esterni nei casi in cui ciò sia rilevante per la comprensione del contenuto.

? Si interviene in thread obsoleti con risposte automatiche o forzate che non arricchiscono la discussione.

? Perché questa impostazione è utile

- Mette al centro la qualità del contributo, non la fonte.
- Valorizza la responsabilità personale e la partecipazione consapevole.
- Non discrimina strumenti moderni (come l'IA), ma impone che vengano usati in modo maturo.

Ecco, credo non servano altre parole... l'IA "comprendendo" perfettamente la questione ha prodotto una serie di punti di attenzione assolutamente significativi e sensati, infinitamente meglio di quanto avrei potuto fare io, specialmente nei dieci secondi netti che ha ci ha messo!

## **Sostenibilità**

Eh sì, anche l'uso di una IA ha un impatto sulla sostenibilità ambientale (consumo di energia). Sono stati calcolati tra 0,03 e 3 wattora di energia richiesti per ogni risposta testuale (in caso di generazione immagini o peggio ancora filmati, il valore sale drammaticamente).

È tanto? È poco? L'oceano è fatto di tante piccole gocce. Si è calcolato che il solo fatto di concludere istintivamente una conversazione con un ringraziamento o un saluto, moltiplicato per decine di milioni di volte al giorno per persone che dicono "grazie

dell'aiuto" o "ciao alla prossima" fanno diversi MWh al giorno di energia buttata inutilmente.

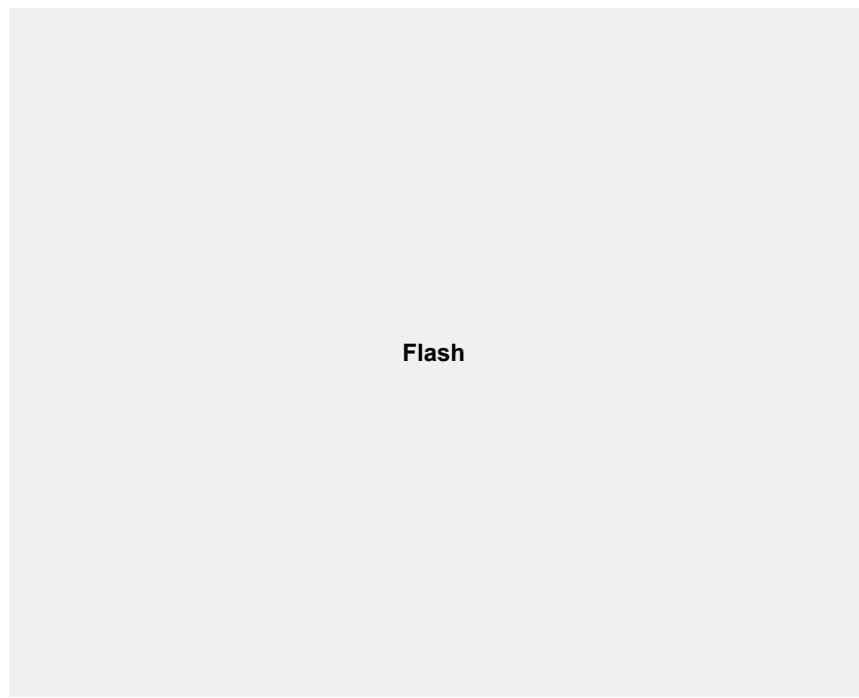
E sì, è del tutto inutile dal momento che (come detto prima), nonostante un'interazione sempre più "umana", in realtà non c'è assolutamente nessuna "entità", o processo in esecuzione, che si aspetta un qualsiasi ringraziamento o saluto.

Quindi: NON RINGRAZIATE e NON SALUTATE le IA! ...almeno quelle attuali :)

Per il resto, se possibile evitare interazioni inutili, fatte solo per pigrizia o svago, per cose che con poca o niente fatica si potrebbero fare in modo "tradizionale". Fare domande precise e complete già dall'inizio, senza creare discussioni troppo lunghe (ricordiamo che ogni volta viene ritrasmesso da capo tutto quanto scritto fino a quel momento). Creare una nuova chat quando si cambia completamente argomento. Tutto questo fa risparmiare energia in generale, banda utilizzata, e token scambiati, cioè le unità di informazione in cui viene scomposto il testo direttamente legati al costo del servizio.

## Il futuro?

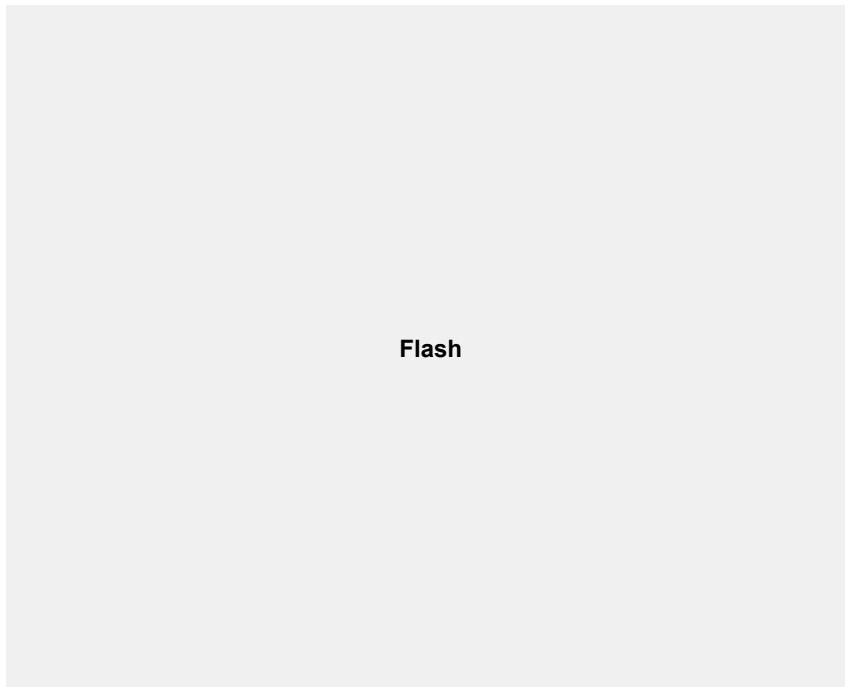
Gli agenti, cioè IA GPT più o meno con CoT in grado di interagire con le funzionalità del PC o con sistemi robotici sono già qui, <https://www.geopop.it/il-primo-robot-umanoide-con-chatgpt-di-openai-si-mostra-in-un-video-ecco-cosa-sa-fare-figure-01/>



**Flash**

Si parla di nuova evoluzione in termini di ragionamento più che di dimensioni del modello. Si parla di sviluppare IA cognitive, mentre persiste il sogno dell'AGI (l'intelligenza generale in tutto e per tutto alla pari con un essere umano), ma per quanto riguarda la vera consapevolezza c'è ancora il "problema difficile della coscienza" (come lo definisce Federico Faggin, il padre della tecnologia CMOS moderna e della CPU Z80) che è un ginepraio in sospeso tra la fisica, le neuroscienze e la filosofia.

Intanto andando sempre più avanti nel gioco dell'imitazione, cominciano ad apparire i primi barlumi di quella che in apparenza sembra essere una forma di auto consapevolezza. Vedere un'IA che, tramite interfaccia robotica ad altissima espressività facciale, si sorprende osservandosi allo specchio e "comprende" quello che sta vedendo è decisamente impressionante:



Perché dico "in apparenza"? Perché esattamente come nessun sistema deterministico può produrre numeri casuali reali, ma solo pseudo ("come se" fossero casuali), anche questa è solo "pseudo consapevolezza" ottenuta in modo "meccanico", è solo un comportamento che appare "come se" (ma privo di reale vissuto interiore). Comunque non c'è dubbio che nei prossimi anni vedremo delle belle sorprese.

Per concludere, non so, per me nato sui libri, con ancora un giradischi a valvole in cantina con cui ascoltavo le fiabe da piccolo, tutto questo è come ritrovarmi improvvisamente nei racconti di fantascienza che leggevo 45 anni fa, razzi che tornano alla base atterrando in piedi, IA che rispondono a qualsiasi argomento, inquietante da un lato, e terribilmente affascinante dall'altro. Ho la sensazione che questo cambiamento sia un po' troppo veloce, e che possiamo finire con la chiave della cassaforte chiusa nella cassaforte. E qui arriva la domanda: tutto questo ci renderà più stupidi o incapaci di

pensare o ragionare? Mi piace pensare che la risposta sia la stessa che ci diamo pensando all'uso della calcolatrice, forse solo un po' più pigri.

Estratto da "<https://www.electroyou.it/mediawiki/index.php?title=UsersPages:Djnz:ia-un-nuovo-mo-n-do-di-fare-le-cose>"