

Una prospettiva storico-tecnica sull'Intelligenza Artificiale dall'automazione del pensiero all'automazione dell'automazione

Stefano Coniglio



UNIVERSITÀ
DEGLI STUDI
DI BERGAMO

6 Novembre 2025

Chi sono?

Sono Stefano Coniglio, docente di informatica e intelligenza artificiale all'Università degli studi di Bergamo.

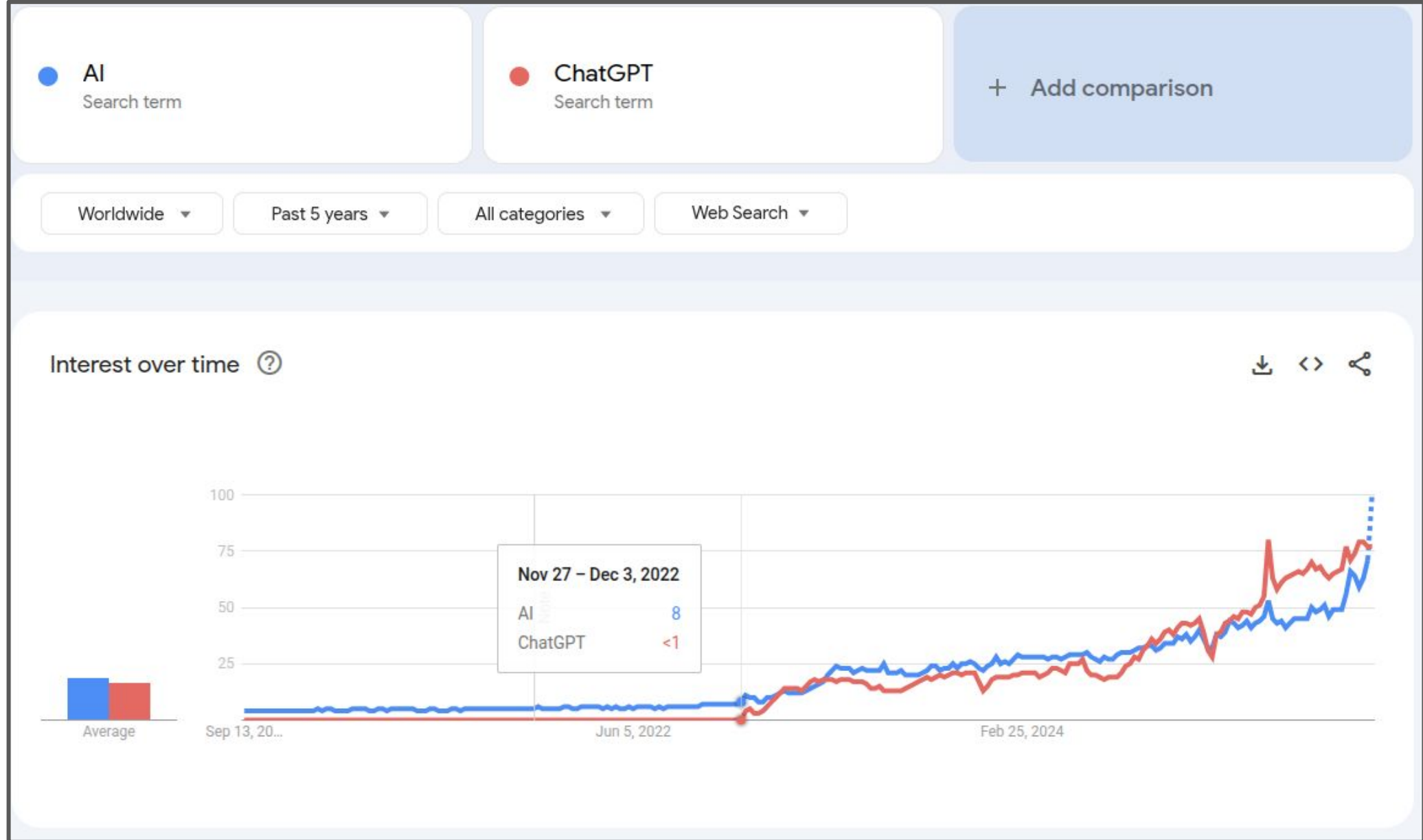
Ho una laurea in Ingegneria Informatica e un dottorato in Tecnologie dell'Informazione (ma ho studiato tutta la matematica che ho trovato).

Ho passato 10 anni all'esterno (Germania, UK).

Sono tornato in Italia da pochi anni principalmente per motivi familiari, ma devo dire che in Italia si vive meglio (e a UniBg si sta decisamente bene).



Lancio di ChatGPT 3.5



Una definizione di IA?

(def)

IA?

Anche se oggi tutti ne parlano, quando ho studiato io ci si vergonava quasi a farlo (e, non a caso, ho studiato altro): l'IA fino agli anni 2000 aveva ricevuto fondi ingentissimi senza aver ancora prodotto quanto promesso per quasi 50 anni.

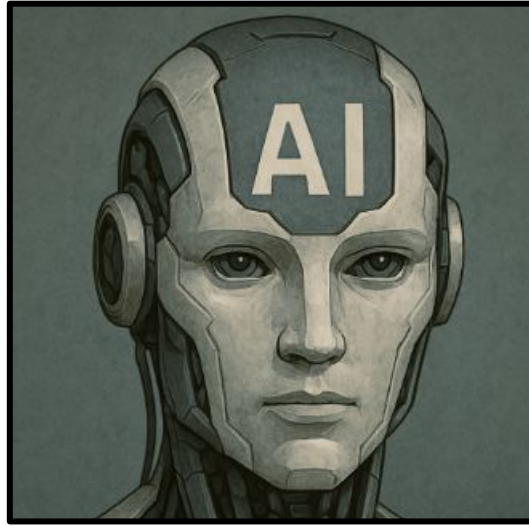
La si chiamava in altri modi:

- ~~Artificial intelligence~~ (1955/6)
- *Pattern recognition* (1960-70)
- *Data mining* (1990)
- *Soft Computing* (?)
- *Machine learning* (1980-2000)
- *Deep learning* (2000)
- *Artificial intelligence* (di nuovo dal 2022)



Un tentativo di definizione: cos'è l'IA?

È la disciplina (scientifica) preposta alla **progettazione** di **agenti intelligenti**.

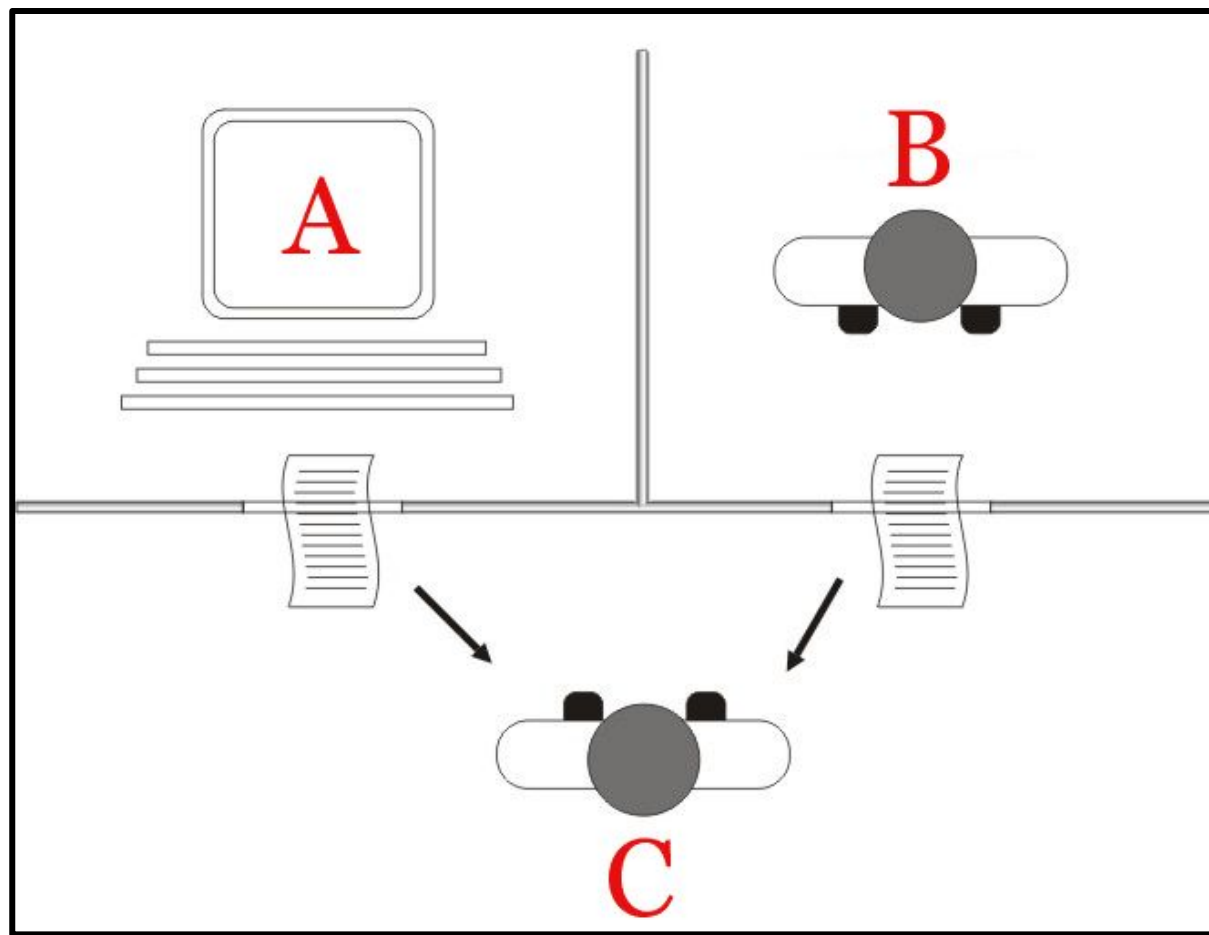


Cos'è un agente? Qualsiasi “cosa” (e.g., programma, robot, ...) capace di **prendere decisioni in autonomia**.

Cosa intendiamo per intelligenza? Dare una definizione è tutt'altro che semplice. Ci “accontentiamo” di un'**intelligenza funzionale**: ci basta che gli agenti che progettiamo si comportino (per un osservatore esterno) come se fossero intelligenti.



Il test di Turing



È un test di **intelligenza funzionale**: “accettiamo” che la macchina (A) sia funzionalmente intelligente se risulta indistinguibile per l'osservatore (C) rispetto a un umano (B).

MIND

A QUARTERLY REVIEW

OF

PSYCHOLOGY AND PHILOSOPHY



I.—COMPUTING MACHINERY AND INTELLIGENCE

BY A. M. TURING

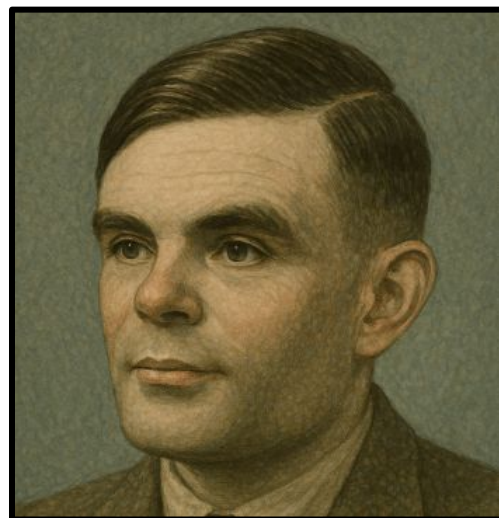
1. *The Imitation Game.*

I PROPOSE to consider the question, 'Can machines think?' This should begin with definitions of the meaning of the terms 'machine' and 'think'. The definitions might be framed so as to reflect so far as possible the normal use of the words, but this attitude is dangerous. If the meaning of the words 'machine' and 'think' are to be found by examining how they are commonly used it is difficult to escape the conclusion that the meaning and the answer to the question, 'Can machines think?' is to be sought in a statistical survey such as a Gallup poll. But this is absurd. Instead of attempting such a definition I shall replace the question by another, which is closely related to it and is expressed in relatively unambiguous words.

The new form of the problem can be described in terms of a game which we call the 'imitation game'. It is played with three people, a man (A), a woman (B), and an interrogator (C) who may be of either sex. The interrogator stays in a room apart from the other two. The object of the game for the interrogator is to determine which of the other two is the man and which is the woman. He knows them by labels X and Y, and at the end of the game he says either 'X is A and Y is B' or 'X is B and Y is A'. The interrogator is allowed to put questions to A and B thus:

C: Will X please tell me the length of his or her hair?

Now suppose X is actually A, then A must answer. It is A's



object in the game to try and cause C to make the wrong identification. His answer might therefore be

'My hair is shingled, and the longest strands are about nine inches long.'

In order that tones of voice may not help the interrogator the answers should be written, or better still, typewritten. The ideal arrangement is to have a teleprinter communicating between the two rooms. Alternatively the question and answers can be repeated by an intermediary. The object of the game for the third player (B) is to help the interrogator. The best strategy for her is probably to give truthful answers. She can add such things as 'I am the woman, don't listen to him!' to her answers, but it will avail nothing as the man can make similar remarks.

We now ask the question, 'What will happen when a machine takes the part of A in this game?' Will the interrogator decide wrongly as often when the game is played like this as he does when the game is played between a man and a woman? These questions replace our original, 'Can machines think?'

2. *Critique of the New Problem.*

As well as asking, 'What is the answer to this new form of the question', one may ask, 'Is this new question a worthy one to investigate?' This latter question we investigate without further ado, thereby cutting short an infinite regress.

The new problem has the advantage of drawing a fairly sharp line between the physical and the intellectual capacities of a man. No engineer or chemist claims to be able to produce a material which is indistinguishable from the human skin. It is possible that at some time this might be done, but even supposing this invention available we should feel there was little point in trying to make a 'thinking machine' more human by dressing it up in such artificial flesh. The form in which we have set the problem reflects this fact in the condition which prevents the interrogator from seeing or touching the other competitors, or hearing their voices. Some other advantages of the proposed criterion may be shown up by specimen questions and answers. Thus:

Q: Please write me a sonnet on the subject of the Forth Bridge.

A: Count me out on this one. I never could write poetry.

Q: Add 34957 to 70764

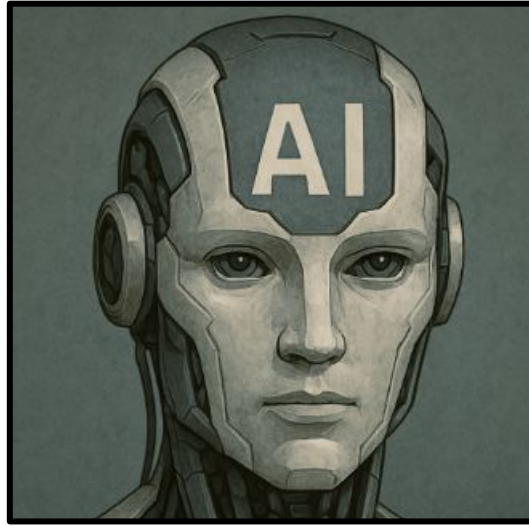
A: (Pause about 30 seconds and then give as answer) 105621.

Q: Do you play chess?

A: Yes.

Un secondo tentativo di definizione

L'agente intelligente è l'agente che **fa la cosa giusta** (secondo il progettista)



Sceglie quali azioni compiere così da massimizzare le proprie prestazioni (per un compito dato).

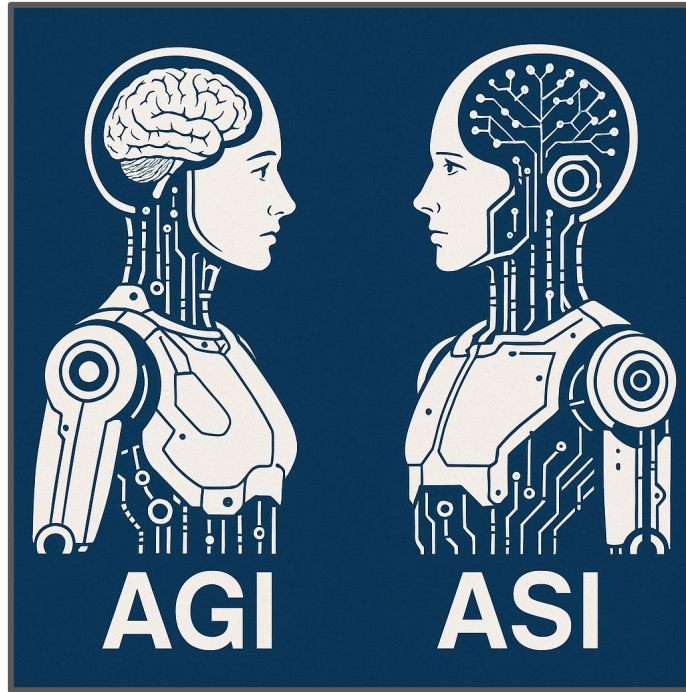
Esempio: minimizzare il tempo di percorrenza di un'auto a guida automatica.



Artificial General/Super Intelligence?

AGI: un'ipotetica IA capace di raggiungere l'intelligenza umana in tutti gli ambiti.

ASI: un'ipotetica IA capace di superare l'intelligenza umana in tutti gli ambiti.



Agenti con capacità super-umane sono ovunque ma, per un AGI/ASI, ne vorremmo uno superiore all'uomo da ogni punto di vista.

Infatti, di “IA” con competenze sovra-umane in attività specifiche (cosiddette **narrow AIs** o **IA strette**) ne conosciamo tantissime. Un esempio su tutti:

Project 1 - Excel

File Home Insert Page Layout Formulas Data Review View WPS PDF Tell me what you want to do... Share

Clipboard Font Alignment Number Styles Cells Editing

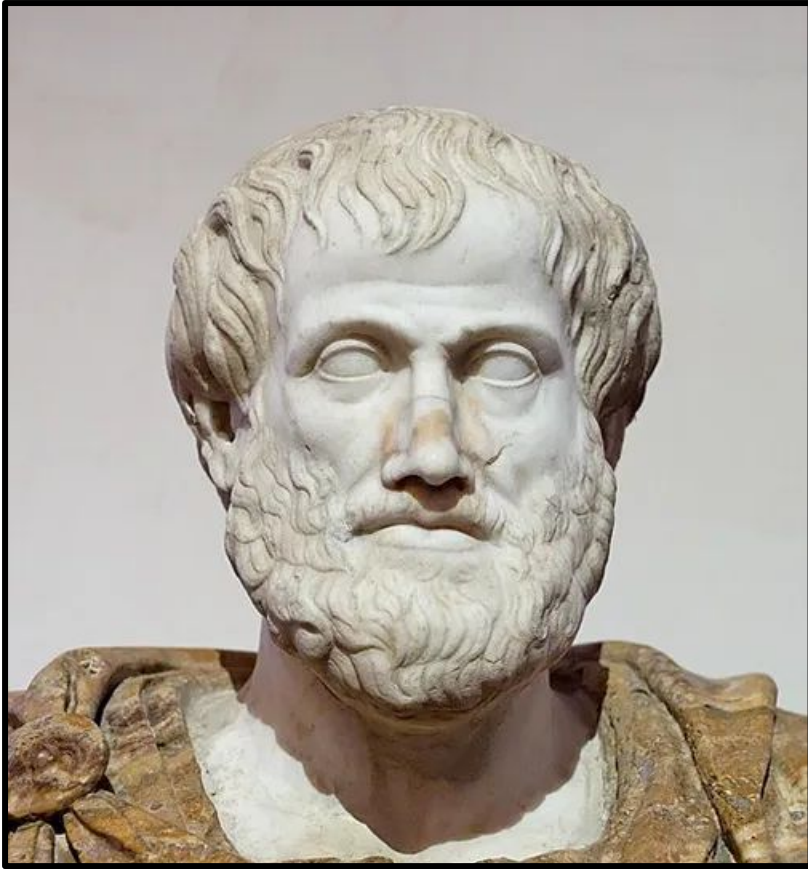
	G	H	I	J	K	L	M	N	O	P	Q	R	S
1	Sal-Mar	Total Sal	Avg Sal	Full Name	Numbers	Round	Round up	Round down	Round up	Round down	Months	Week	
2	1500	4200	1400	RNM KUMAR	1.03333	1	2	1	1.1	1.03	jan	Mon	
3	2000	5500	1833	GOPAL VERMA	2.0555	2	3	2	2.1	2.05	feb	Tue	
4	1900	5200	1733	JOSEPH PAUL	2.999	3	3	2	3	2.99	mar	Wed	
5	1800	4500	1500.00	HARI SINGH	8.96	9	9	8	9	8.95	apr	Thu	
6	2900	7400	2467	RAJA RAM	1.333	1	2	1	1.4	1.33	may	Fri	
7													
8													
9													
10													
11													
12													
13													
14													

project 1 Sheet1

Gli albori della disciplina

(h1.pio)

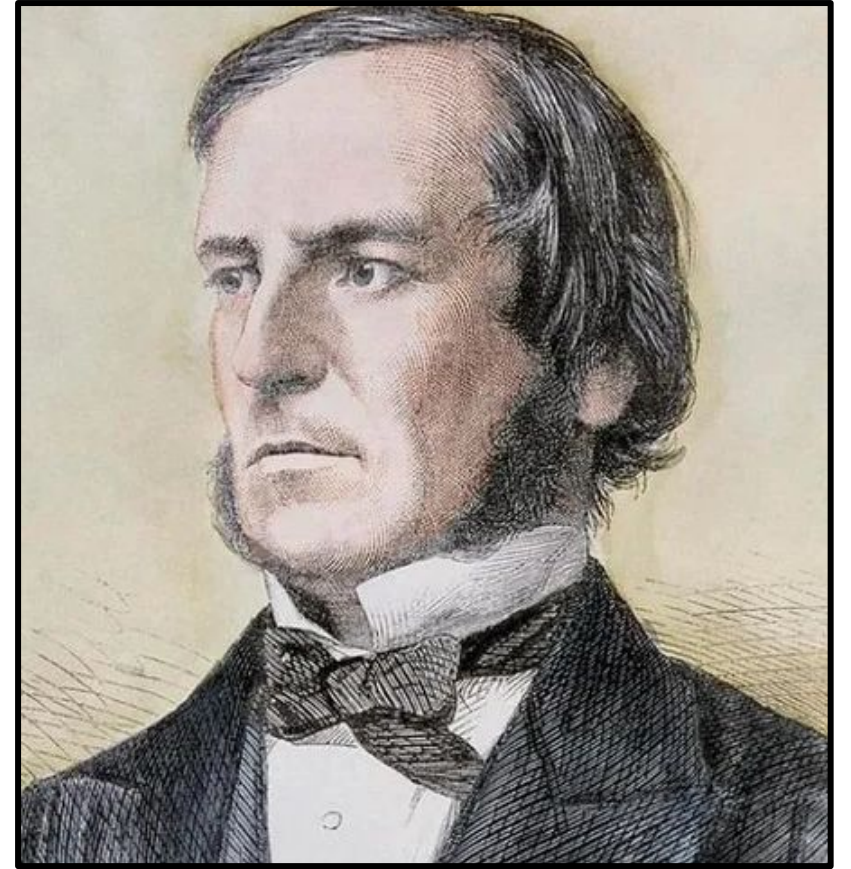
Le radici dell'IA sono profonde



Aristotele (384 – 322 a.C.)



Gottfried Wilhelm Leibniz (1646–1716)



George Boole (1815 – 1864)

Potremmo dire che le radici dell'IA si trovano nell'idea (curiosamente, ma innegabilmente, molto umana) di **automatizzare il pensiero**.

Calcolo automatico $\Rightarrow$ Pensiero automatico

Il nesso tra calcolo e IA è fortissimo.

Attorno agli anni '50, il problema dell'automatizzazione del calcolo viene risolto (col diffondersi del calcolatore digitale moderno).

L'asticella si sposta dal tentativo (andato a buon fine) di automatizzare il calcolo a quello di automatizzare il pensiero.

L'IA nasce qui.



ENIAC

L'**ENIAC** (Electronic Numerical Integrator and Computer) fu il primo computer elettronico ad essere programmabile e general-purpose.

Fu progettato da **John Mauchly** e **J. Presper Eckert** (University of Pennsylvania, U.S.) e completato nel 1945 per fare calcoli balistici nella Seconda Guerra Mondiale per lo United States Army's Ballistic Research Laboratory.

Il primo problema risolto da ENIAC fu uno studio di fattibilità per armi termonucleari. Ci mise un mese di lavoro. Il problema proveniva dal Los Alamos National Laboratory (e, nello specifico, dai militari e scienziati del Manhattan Project), dove il matematico ungherese (nazionalizzato americano) **John von Neumann** (uno dei genii, seppur guerrafondaio, del secolo scorso) stava calcolando come sganciare l'atomica (voleva colpire Tokio in origine) per massimizzarne i danni.

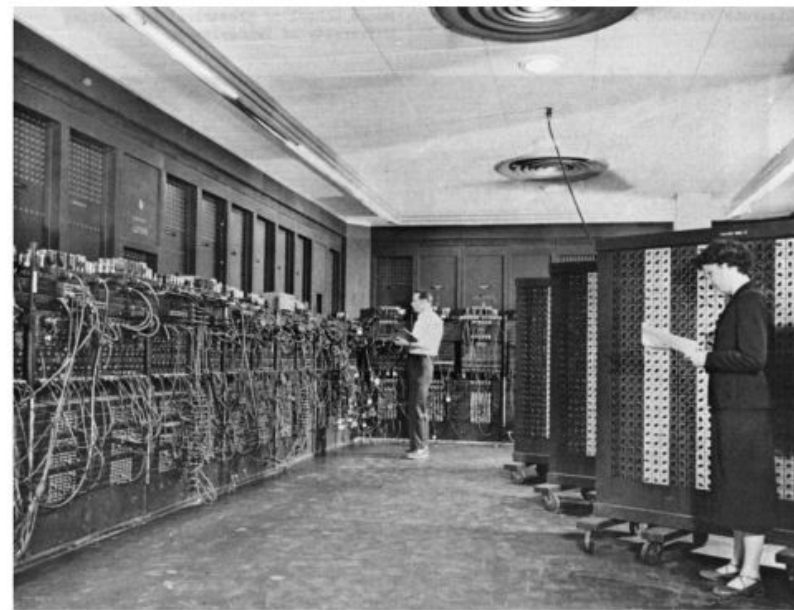
L'ENIAC poteva calcolare una traiettoria in 30 secondi (a un umano sarebbero servite 20 ore).

Costò 7.237 ore di lavoro e \$486.804,22 dollari (l'equivalente di circa 5 milioni dei giorni nostri) - la spesa prevista iniziale era di soli 61.700 dollari (costò 8 volte il costo preventivato...).

Occupava 180 metri quadrati (13,5 metri x 13,5 metri) per un peso di 30 tonnellate e conteneva 18.000 valvole termoioniche. Una volta acceso, riscaldava (a causa delle valvole) l'ambiente fino 50 °C.

Alla sua prima accensione, causò un black-out nel quartiere ovest di Filadelfia (consumava davvero molto!).

Era a **programma cablato** (cioè lo si programmava collegando cavi) (vedi foto).



La nascita dell'IA come disciplina scientifica

(h2.dart)

Dartmouth, 1956

A Proposal for the

DARTMOUTH SUMMER RESEARCH PROJECT ON ARTIFICIAL INTELLIGENCE

June 17 - Aug. 16

We propose that a 2 month, 10 man study of artificial intelligence be carried out during the summer of 1956 at Dartmouth College in Hanover, New Hampshire. The study is to proceed on the basis of the conjecture that every aspect of learning or any other feature of intelligence can in principle be so precisely described that a machine can be made to simulate it. An attempt will be made to find how to make machines use language, form abstractions and concepts, solve kinds of problems now reserved for humans, and improve themselves. We think that a significant advance can be made in one or more of these problems if a carefully selected group of scientists work on it together for a summer.

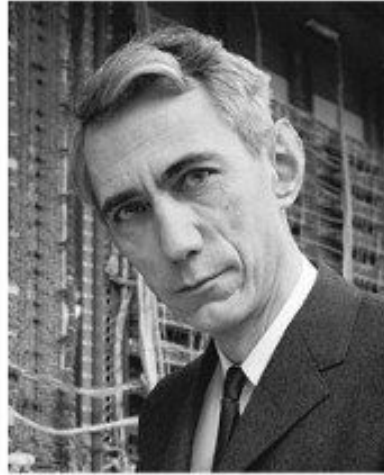
1956 Dartmouth Conference: The Founding Fathers of AI



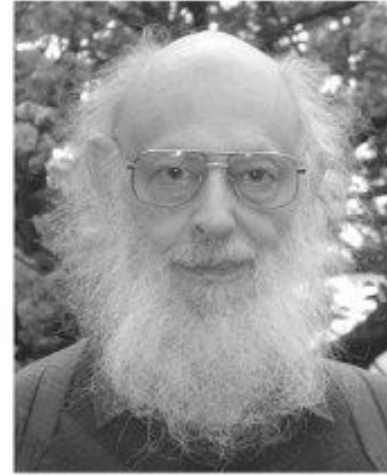
John MacCarthy



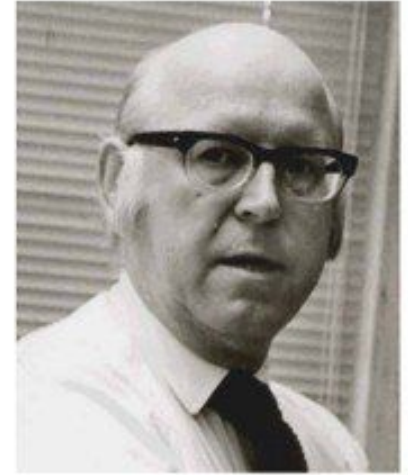
Marvin Minsky



Claude Shannon



Ray Solomonoff



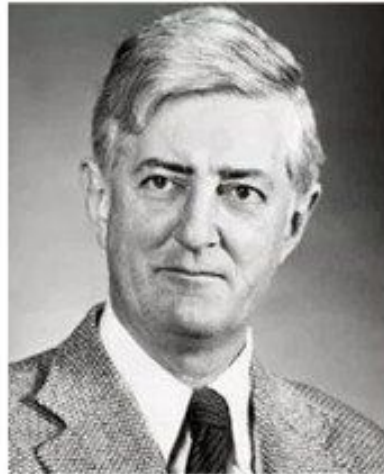
Alan Newell



Herbert Simon



Arthur Samuel



Oliver Selfridge



Nathaniel Rochester



Trenchard More

The following are some aspects of the artificial intelligence problem:

1) Automatic Computers

If a machine can do a job, then an automatic calculator can
be programmed to simulate the machine. The speeds and
memory capacities of present computers may be insufficient
to simulate many of the higher functions of the human brain,
but the major obstacle is not lack of machine capacity, but
our inability to write programs taking full advantage of what
we have.

2) How Can a Computer be Programmed to Use a Language

It may be speculated that a large part of human thought consists of manipulating words according to rules of reasoning

3. Neuron Nets

How can a set of (hypothetical) neurons be arranged so as to form concepts. Considerable theoretical and experimental work has been done on this problem by Uttley, Rashevsky and his group, Farley and Clark, Pitts and McCulloch, Minsky, Rochester and Holland, and others. Partial results have been obtained but the problem needs more theoretical work.

4. Theory of the Size of a Calculation

If we are given a well-defined problem (one for which it is possible to test mechanically whether or not a proposed answer is a valid answer) one way of solving it is to try all possible answers in order. This method is inefficient, and to exclude it one must have some criterion for efficiency of calculation. Some consideration will show that to get a measure of the efficiency of a calculation it is necessary to have on hand a method of measuring the complexity of calculating devices which in turn can be done if one has a theory of the complexity of functions. Some partial results on this problem have been obtained by Shannon, and also by McCarthy.

5. Self-Improvement

Probably a truly intelligent machine will carry out activities which may best be described as self-improvement. Some schemes for doing this have been proposed and are worth further study. It seems likely that this question can be studied abstractly as well.

6. Abstractions

A number of types of "abstraction" can be distinctly defined and several others less distinctly. A direct attempt to classify these and to describe machine methods of forming abstractions from sensory and other data would seem worthwhile.

7. Randomness and Creativity

A fairly attractive and yet clearly incomplete conjecture is that the difference between creative thinking and unimaginative competent thinking lies in the injection of a some randomness. The randomness must be guided by intuition to be efficient. In other words, the educated guess or the hunch include controlled randomness in otherwise orderly thinking.

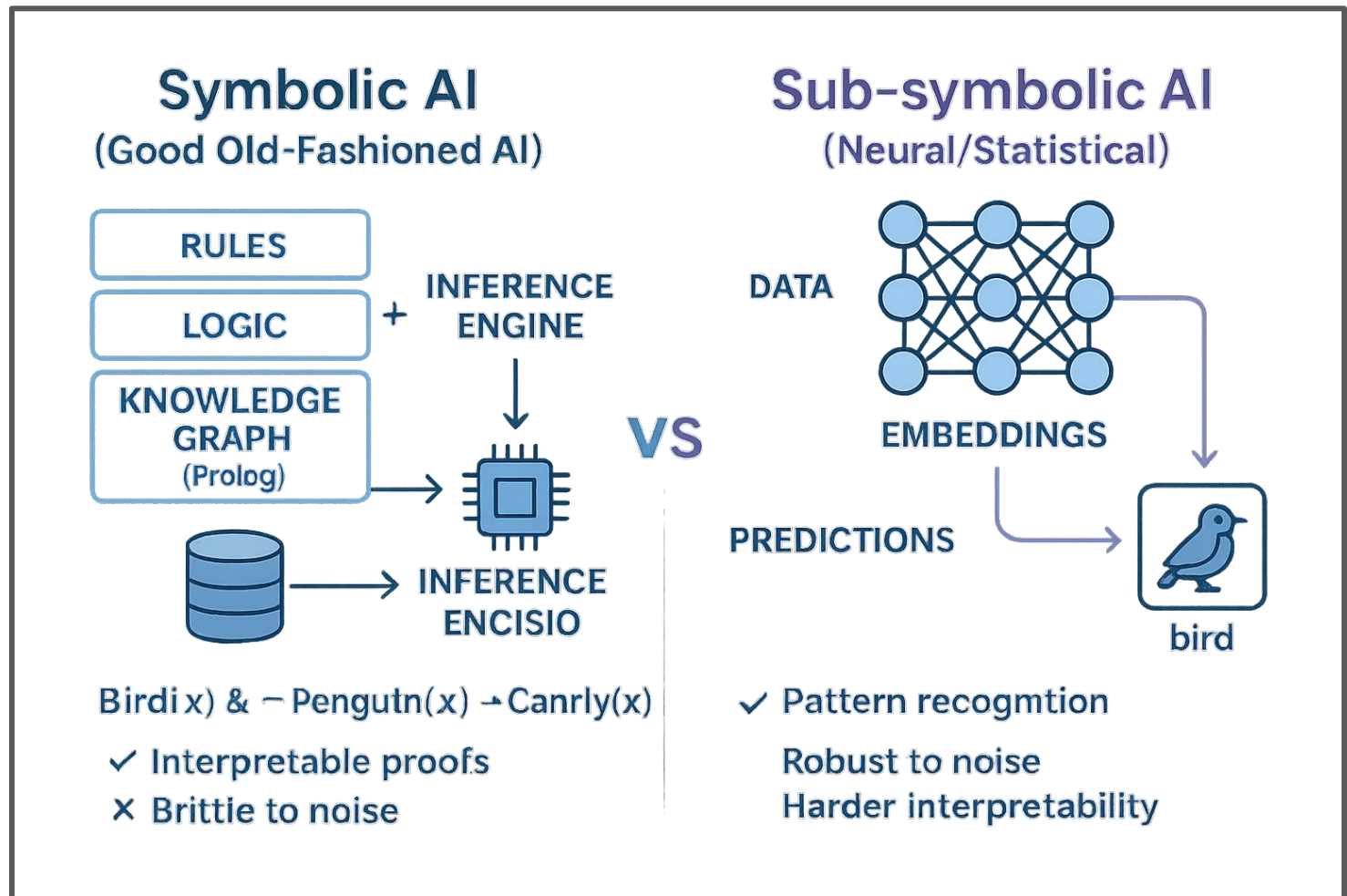
Due tipi di IA

Esistono due tipi (principali) di IA:

- **IA simbolica**
- **IA sotto-simbolica (o connessionista)** – l'IA moderna è quest'ultima

Vediamole in dettaglio.

(A destra, un discutibile tentativo di ChatGPT di illustrare questi due tipi di IA)

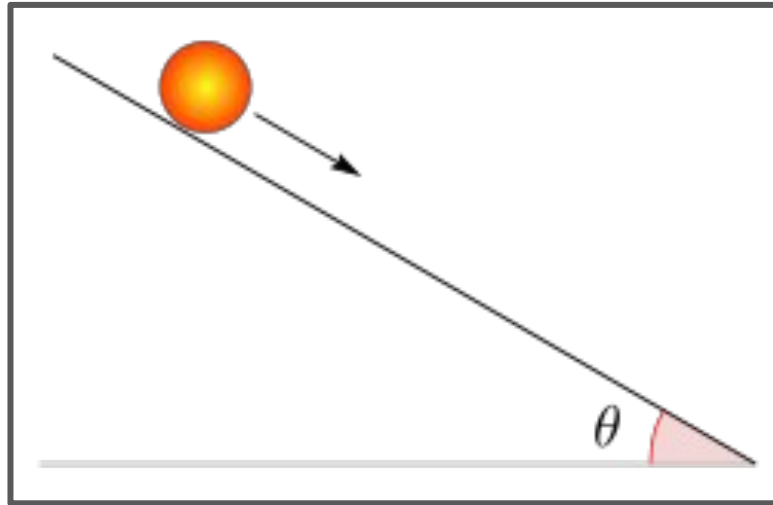


IA simbolica

(symb)

IA simbolica

Si basa essenzialmente sulla matematica “classica”, dove per “classica” intendiamo una matematica in cui i modelli (del mondo) sono scritti da un umano e, per tanto, sono interpretabili. Tra le matematiche, predilige (almeno fino agli anni 80 circa) la logica (di Boole e le sue estensioni) – pensiamo al costrutto *if-then-else* in un linguaggio di programmazione.



$$s = \frac{1}{2}at^2$$

E' il paradigma dominante fino agli anni '80.

Un primo esempio

$(SensorePolvere \wedge BatteriaCarica) \Rightarrow Aspira$

```
for iterazione in range(1,5):  
    sensorePosizione = rilevaPosizione()  
    sensoreSporco = rilevaSporcizia()  
    if sensoreSporco == "Sporco":  
        aspira()  
    if sensoreSporco == "Pulito":  
        if sensorePosizione == "A":  
            destra()  
        if sensorePosizione == "B":  
            destra()  
        if sensorePosizione == "C":  
            sinistra()
```

A, Pulito	B, Pulito	C, Sporco
-----------	-----------	-----------



Un secondo esempio: Google maps

The screenshot displays the Google Maps interface with a public transit route highlighted in blue. The route starts at the origin (red dot) and ends at the destination (red dot). The route is shown on a map of the Milan area, with labels for various locations including Merate, Seregno, Lissone, Vimercate, Trezzo sull'Adda, Cassano d'Adda, Treviglio, Pandino, Melegnano, and San Giuliano Milanese. The route is labeled with 'M' for Metro and '3' for Line 3. The total travel time is 1 hr 54 min. The route is shown with a red line segment indicating the current location.

Left Panel:

- Mode selection: Best, 1h 6m, 1h 53m, 12 hr, 3h 24m, —
- Origin: [Redacted]
- Destination: Via dei Caniana, 2, 24127 Bergamo BG
- Leave now ▾
- Options
- Send directions to iPhone
- Copy link
- Route summary: 2:42 PM—4:36 PM, 1 hr 54 min
- Route details: 2:48 PM from P.Za Ventiquattro Maggio, 25 min, every 30 min

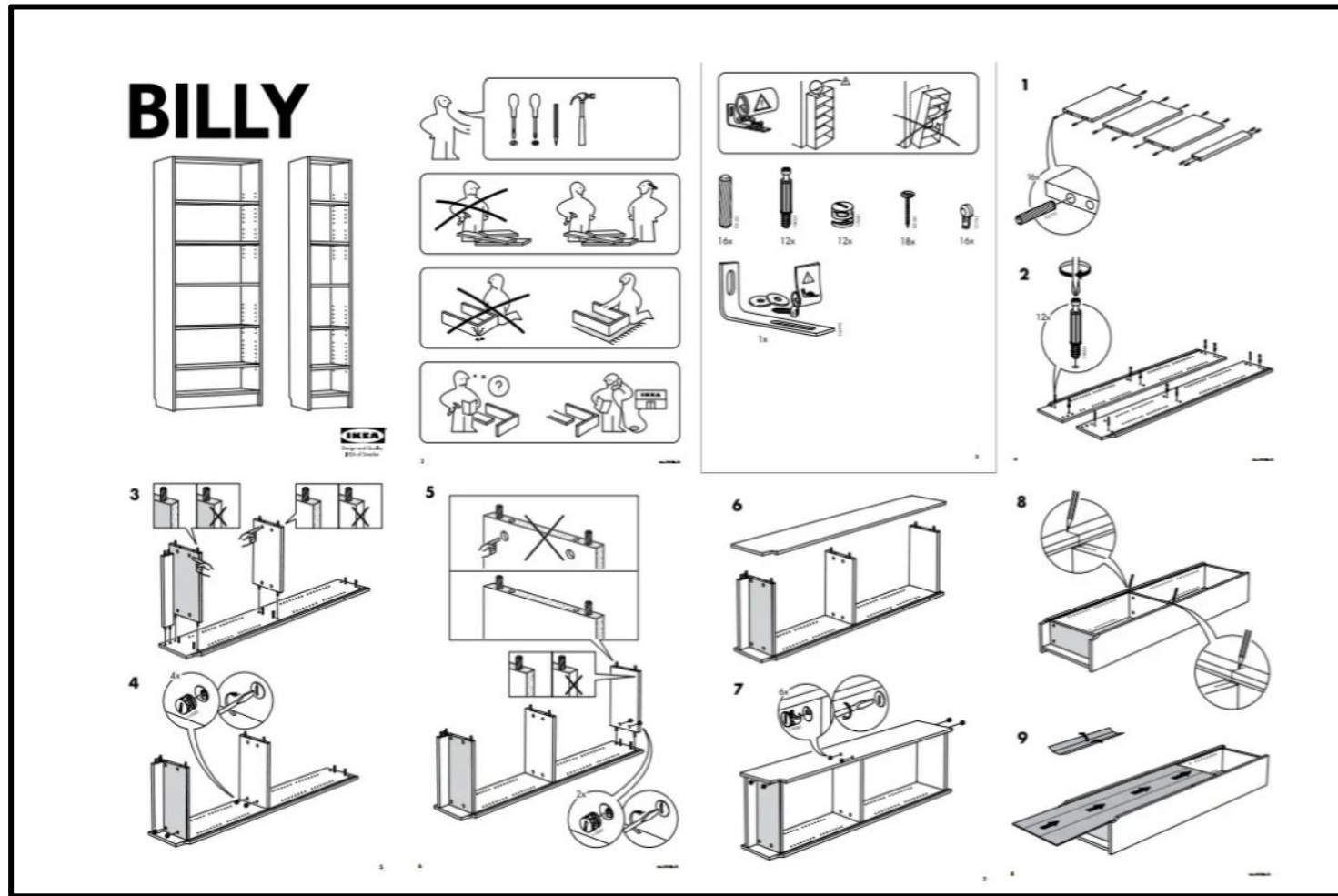
Map Panel:

- Route: M > 3 > M > 1 > Z301 > 25 min
- Travel time: 1 hr 54 min
- Intermediate stops: 1 hr 46 min, 1 hr 41 min
- Map data ©2025 Google, Italy, Terms, Privacy, Send Product Feedback, 10 km

Algoritmi

Nell'IA simbolica troviamo (modelli matematici e) algoritmi scritti da un umano:

Algoritmo = una sequenza (finita) di istruzioni che possono essere eseguite (anche da un umano) mediante un certo numero (finito) di operazioni.



Ricetta per la torta di mele



PREPARAZIONE

COME PREPARARE LA TORTA DI MELE



Per preparare la torta di mele come prima cosa sbucciate le mele **1**, tagliatele a spicchi ed eliminate il torsolo. Tagliatele poi a fettine sottili dal lato corto **2**. In una ciotola versate le uova e lo zucchero **3**.



Montate con le fruste elettriche e non appena il composto sarà chiaro e spumoso aggiungete un pizzico di sale **4**. Versate poi anche il burro tiepido **5** e la cannella **6**.

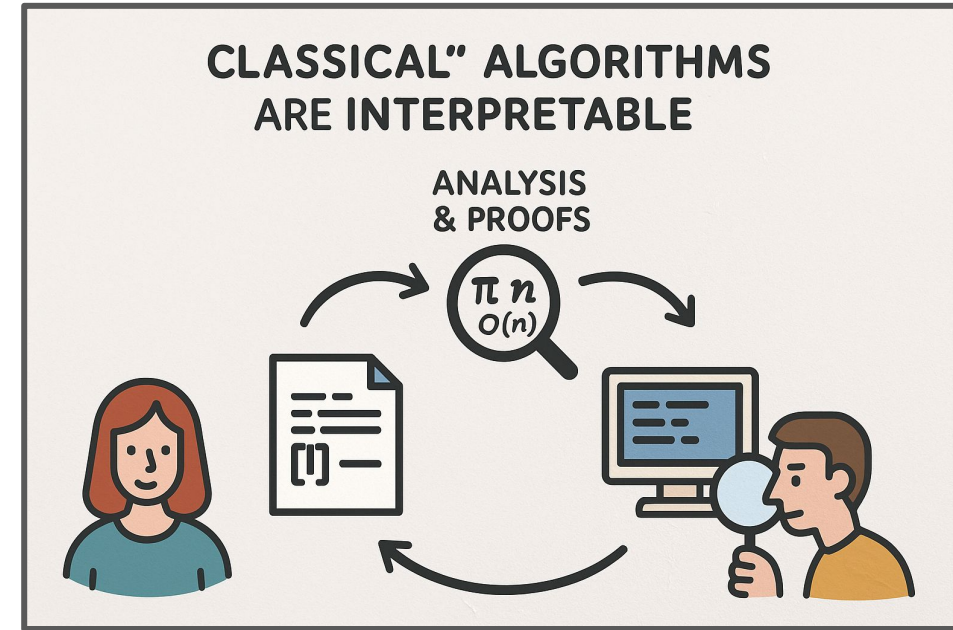


Tipicamente...

... gli algoritmi vengono scritti dagli umani. Pertanto:

- possono essere analizzati (spesso matematicamente)
- ci si possono dimostrare sopra dei teoremi (di correttezza, efficienza, robustezza,)
- se qualcosa va storto eseguendoli, possiamo ispezionarli e capire dove stia il problema.

In summa, gli algoritmi “classici” sono **interpretabili**.



Esempio: se qualcosa va storto in Google maps, o l'algoritmo è buggato (scritto male) o concepito male (errore matematico), o usa dati sbagliati – in ogni caso, possiamo capire dove sta l'errore e sistemarlo.

Declino

L'IA simbolica è completamente basata su di un modello matematico preciso (“simbolico” come $s = \frac{1}{2} a t^2$) del compito che l'IA deve risolvere – costruire questo modello è tutt'altro che facile.

In termini moderni diremmo che l'IA simbolica non è capace di generalizzare out of sample

L'idea che i modelli dell'IA (in particolare dopo la crisi (anche di mercato) dei sistemi esperti)

(I sistemi esperti sono sistemi di supporto alle decisioni molto complessi e basati su regole; si rivelano fragili, con scarse proprietà di *scaling* e incapaci di adattarsi a contesti nuovi o incerti.)



A Symbolics 3640 Lisp machine: an early (1984) platform for expert systems

Il disastro del Columbia (2003)

Il 1° febbraio 2003 lo Shuttle **Columbia** si disintegrò durante il rientro nell'atmosfera, provocando la morte dei sette astronauti a bordo. La causa fu un danno al rivestimento termico dell'ala sinistra, colpita da un pezzo di materiale isolante staccatosi dal serbatoio esterno al momento del lancio.



È un altro esempio di modello (simbolico) usato fuori dal dominio di validità: il modello era stato tarato solo per danni piccoli (non si pensava che si potesse staccare un frammento così grosso dal rivestimento termico) e classificò il danno come *non critico*.



IA sotto-simbolica
(connessionista)
e apprendimento automatico

(nn.int)

Simboli = etere luminifero?

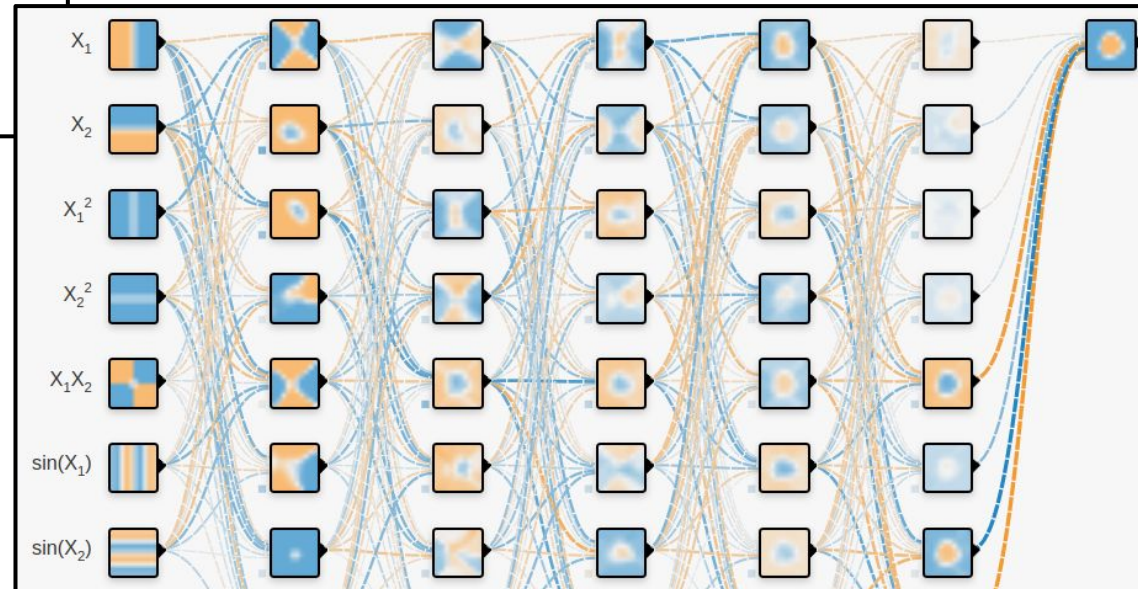
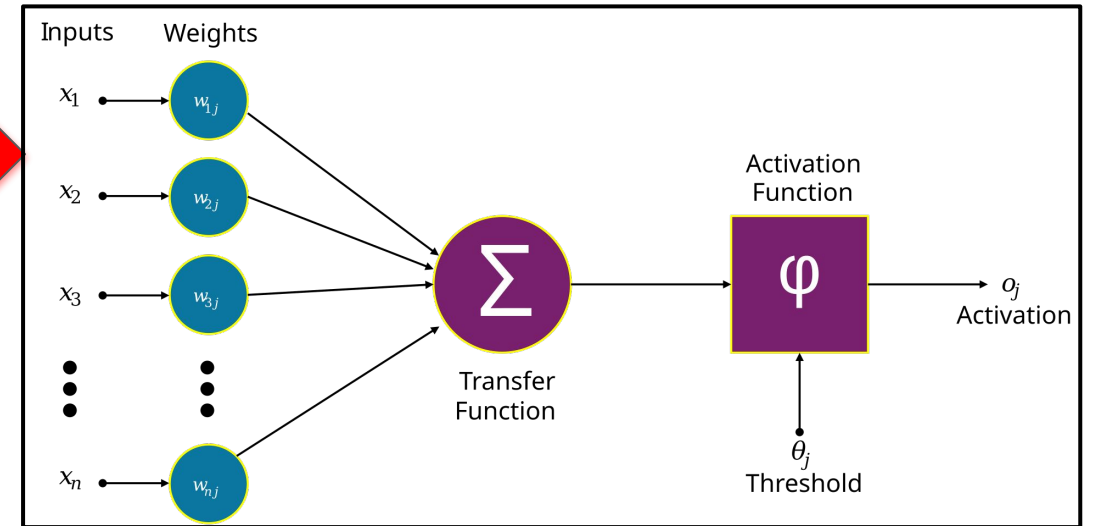
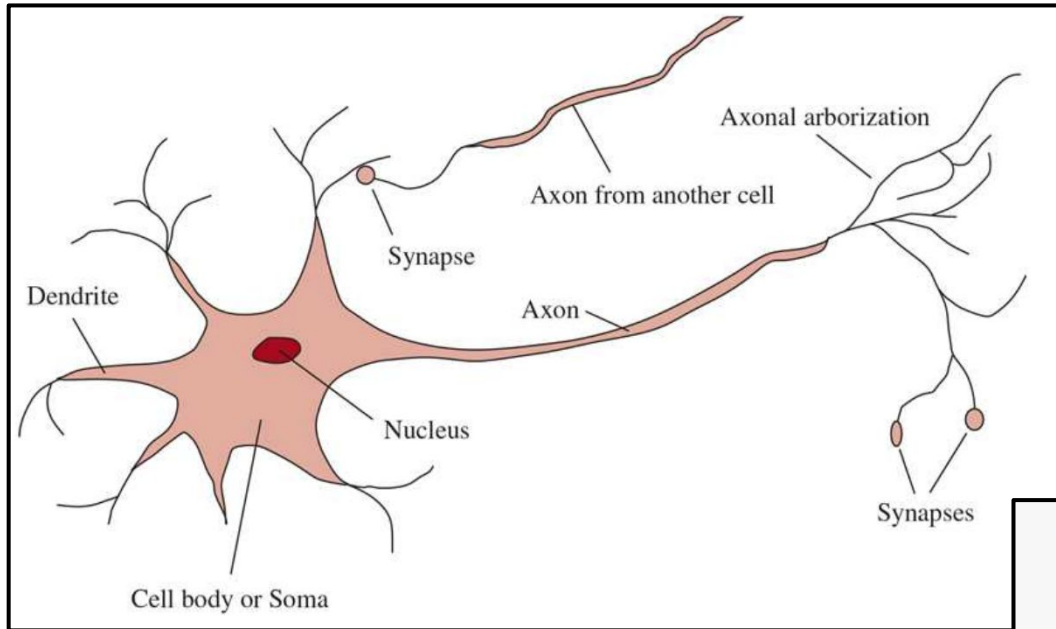


Geoffrey Hinton (Nobel per la fisica nel 2024), descrive i simboli come l'**etere luminifero** dell'IA, riferendosi alla sostanza fisica (inesistente) attraverso cui le onde elettromagnetiche si propagano che i fisici del 1800 cercavano.

Il suo punto è: **il cervello manipola numeri** (attivazioni neurali) **e non simboli**. I simboli entrano nel cervello (immagini, suono) e ne escono (parole, gesti, azioni), ma l'elaborazione interna che avviene nel cervello si basa su qualcosa di più basso livello (*subsimbolico*, appunto).

Motore dell'IA moderna sono le reti neurali

Le **reti neurali** sono modelli matematici “grezzi” ispirati al funzionamento dei neuroni biologici.

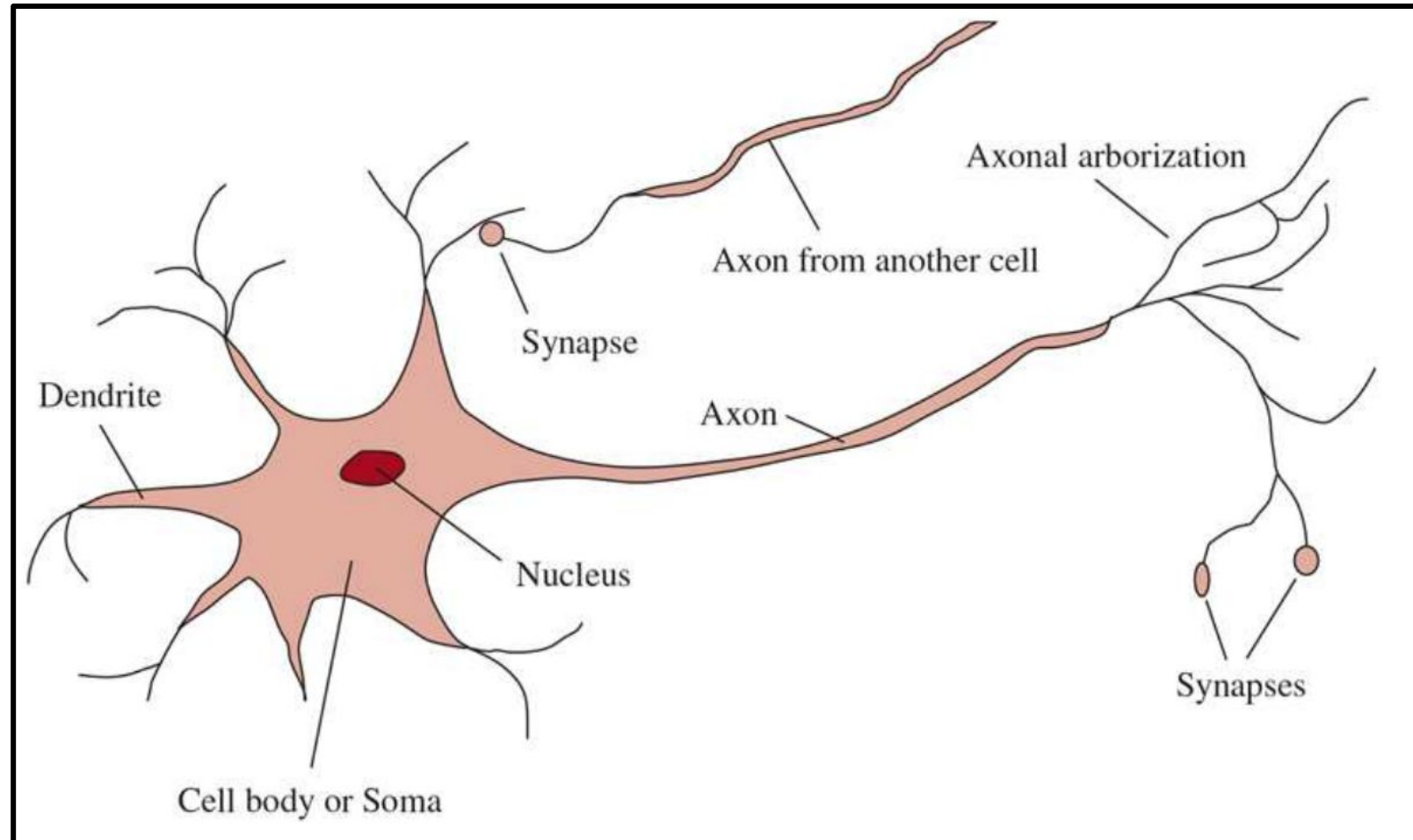


I neuroni biologici

Camillo Golgi (1843–1926) li vede per primo al microscopio

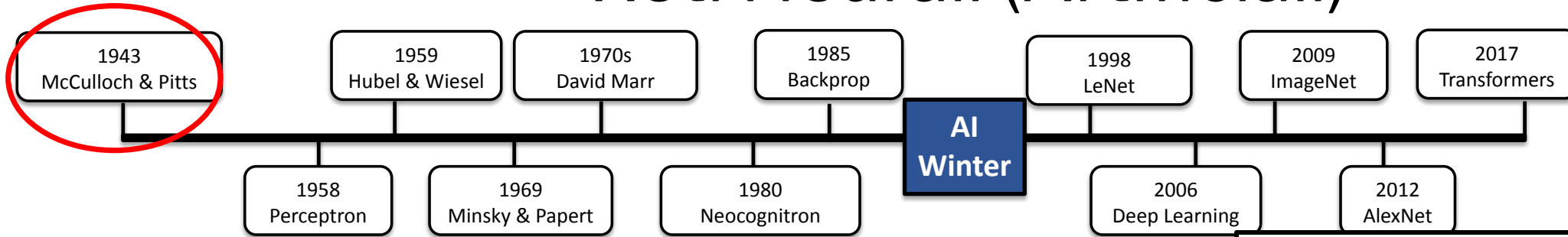
Santiago Ramon y Cajal (1852–1934) ne studia l'organizzazione funzionale

John Searle: “*neurons cause minds*” (distruggendoli perdiamo facoltà cognitive – WW2)



I neuroni interagiscono mediante reazioni elettrochimiche attraverso le **sinapsi**.

Reti Neurali (Artificiali)



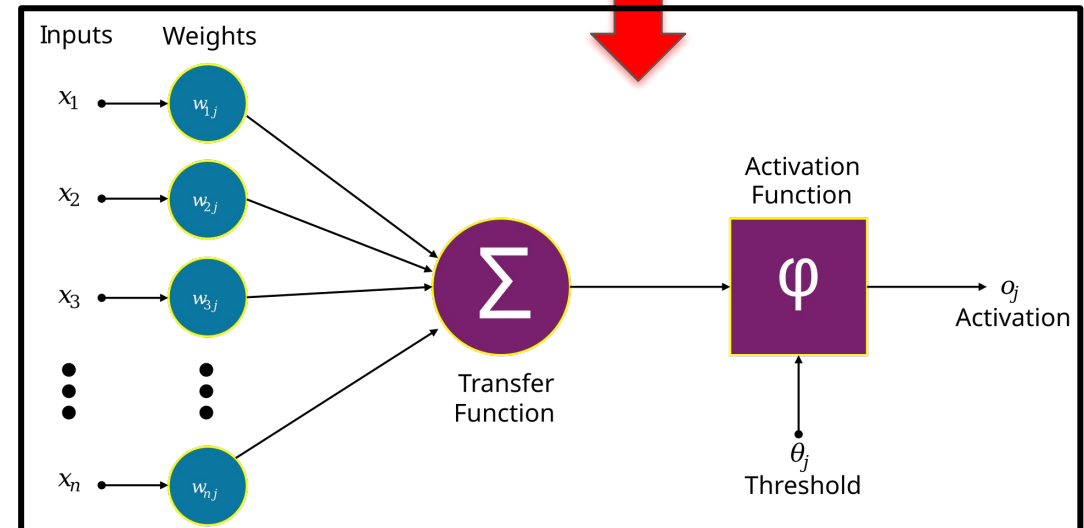
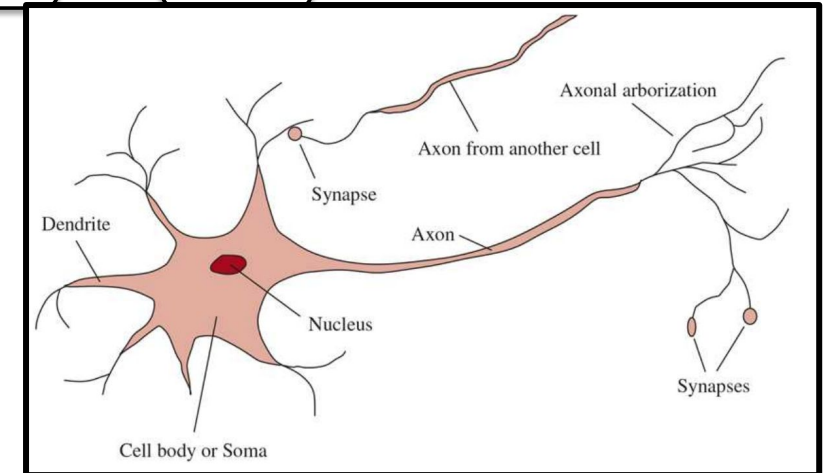
BULLETIN OF
MATHEMATICAL BIOPHYSICS
VOLUME 5, 1943

A LOGICAL CALCULUS OF THE IDEAS IMMANENT IN NERVOUS ACTIVITY

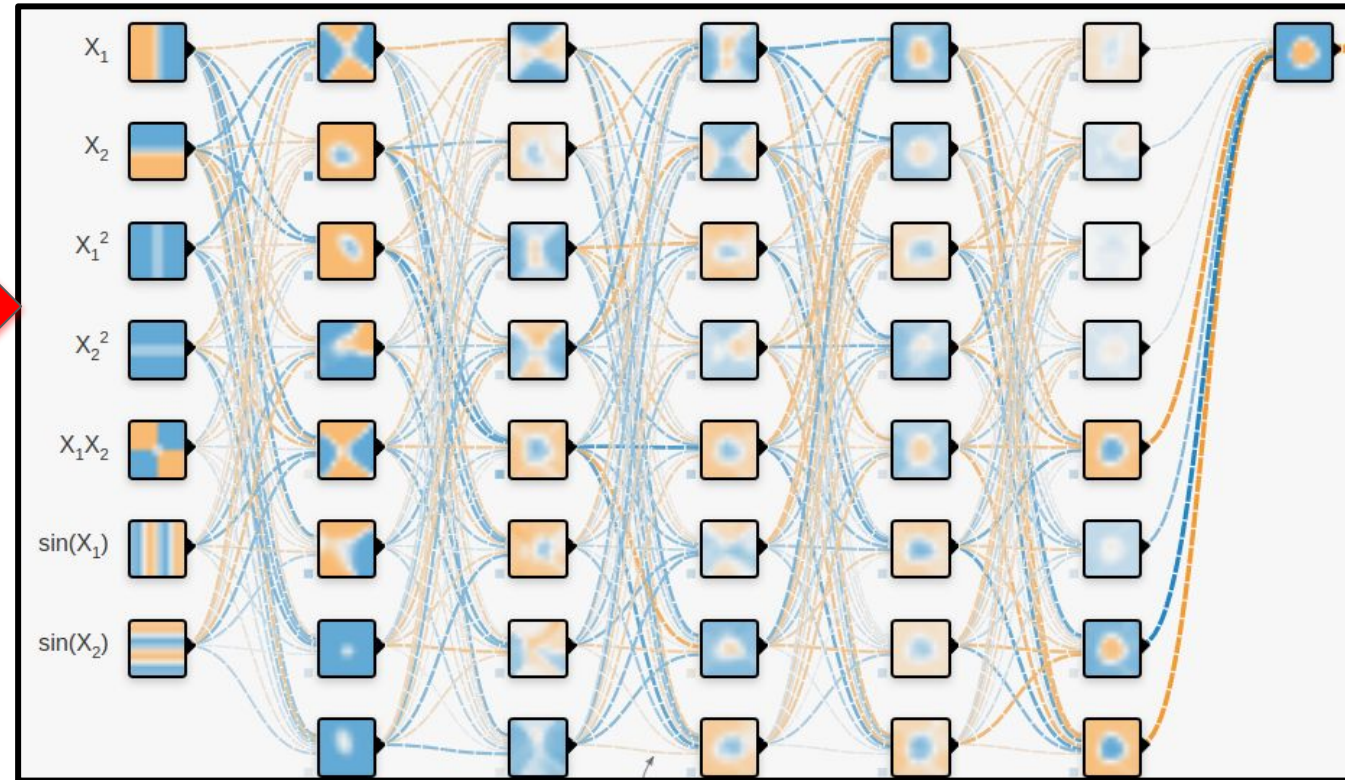
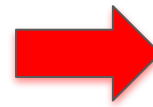
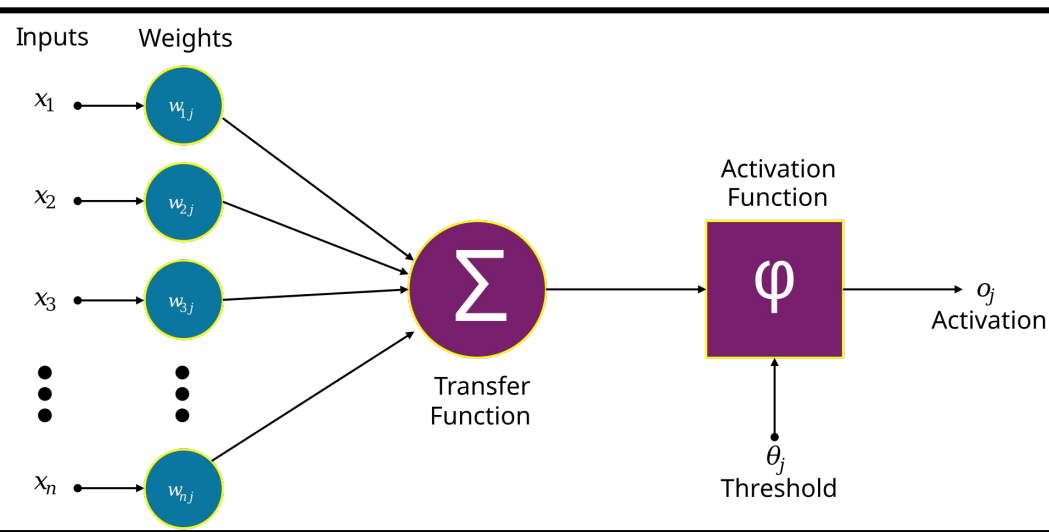
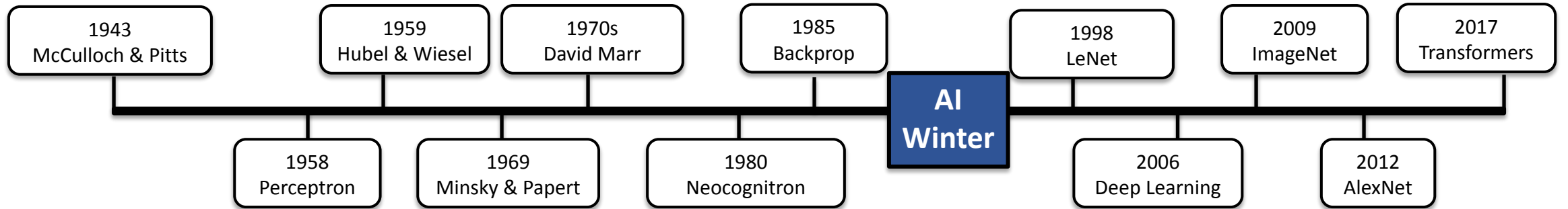
WARREN S. MCCULLOCH AND WALTER PITTS

FROM THE UNIVERSITY OF ILLINOIS, COLLEGE OF MEDICINE,
DEPARTMENT OF PSYCHIATRY AT THE ILLINOIS NEUROPSYCHIATRIC INSTITUTE,
AND THE UNIVERSITY OF CHICAGO

Because of the "all-or-none" character of nervous activity, neural events and the relations among them can be treated by means of propositional logic. It is found that the behavior of every net can be described in these terms, with the addition of more complicated logical means for nets containing circles; and that for any logical expression satisfying certain conditions, one can find a net behaving in the fashion it describes. It is shown that many particular choices among possible neurophysiological assumptions are equivalent, in the sense that for every net behaving under one assumption, there exists another net which behaves under the other and gives the same results, although perhaps not in the same time. Various applications of the calculus are discussed.

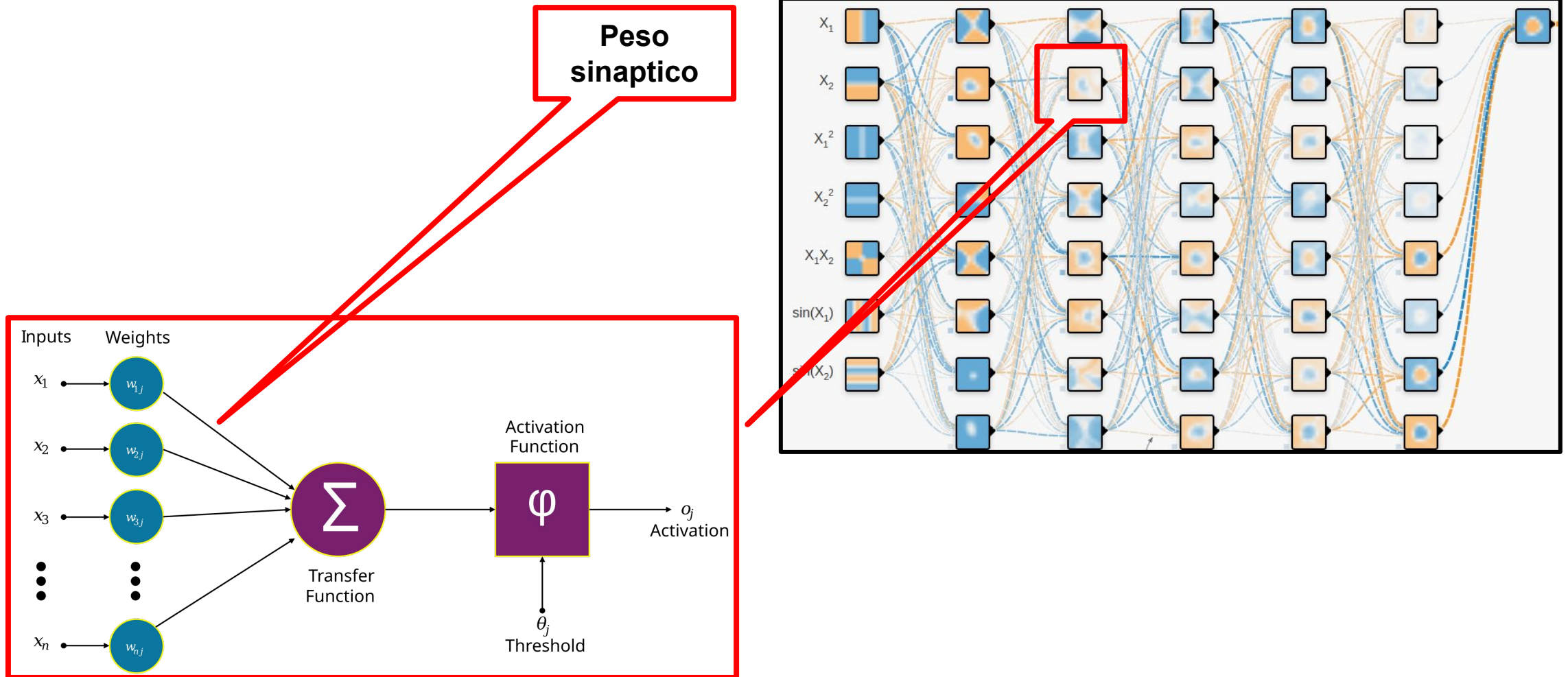


IA subsimbolica (o connessioneista) = reti neurali



Reti neurali: addestramento

Le reti neurali vanno **addestrate** attraverso un algoritmo che ne determina il valore dei pesi sinaptici così che (*apprendimento supervisionato*) la rete possa riprodurre la relazione ingresso-uscita che ha osservato in addestramento.



Addestramento

Donald Hebb: “*neurons that fire together wire together*”.

Esempio: Immaginiamo di voler imparare a suonare uno strumento. Abbiamo ben in mente il motivo che vogliamo suonare e proviamo a riprodurlo un po' a casaccio.

All'inizio, ne uscirà un suono molto cattivo, ma perseveriamo.



Dopo molti giorni di esercizio:

- il nostro orecchio funge da guida (“*loss function*”): ci dice quanto ci siamo avvicinati alla qualità d’esecuzione che cerchiamo
- se ci accorgiamo di aver fatto bene, tentiamo di ripetere
- se ci accorgiamo di aver fatto male, cerchiamo di evitare l’errore.

Neuralmente (*), stiamo rafforzando le *activation pathways* neuronali che ci portano a suonare bene e indebolendo quelle che ci portano a suonare male.

(*) Durante il giorno, stiamo solo “marcando” le pathways “giuste”; il loro rafforzamento avviene nel sonno.

Un esempio: l'interpretazione di immagini

Anche se tutti sappiamo distinguere un gatto da un cane a colpo d'occhio...



... sfido chiunque a descrivere la sequenza di passaggi logici che seguiamo per farlo.

L'IA connessionista funziona in modo analogo: non diamo alla macchina un insieme di regole per determinare il contenuto dell'immagine: le chiediamo di imparare a farlo da sé.

Image 1



Image 2



Image 3



Image 4



Image 5



Image 6



Image 7



Image 8



Image 9



Image 10



Image 11



Image 12



Image 13



Image 14



Image 15



Image 16



Image 17



Image 18



Image 19



Image 20



Image 21



Image 22



Image 23

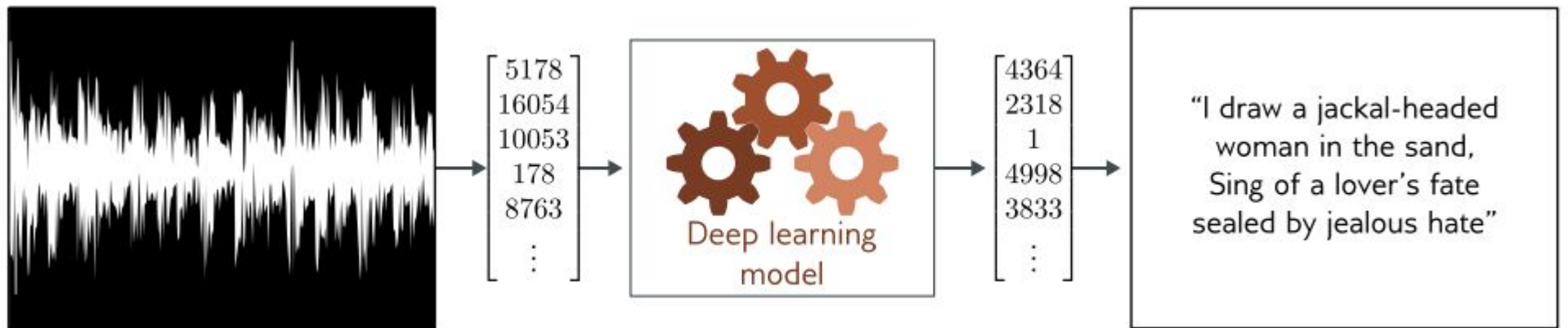
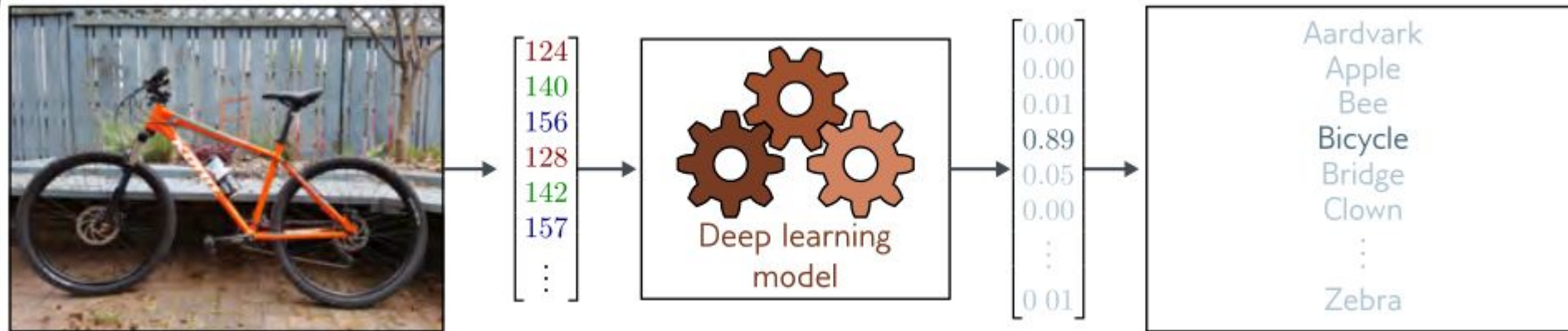


Image 24



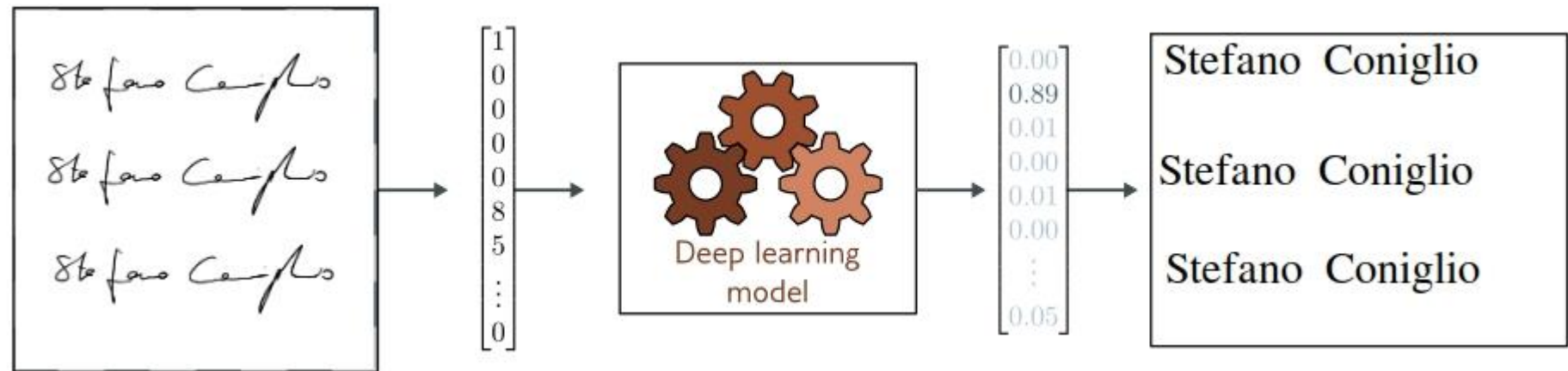
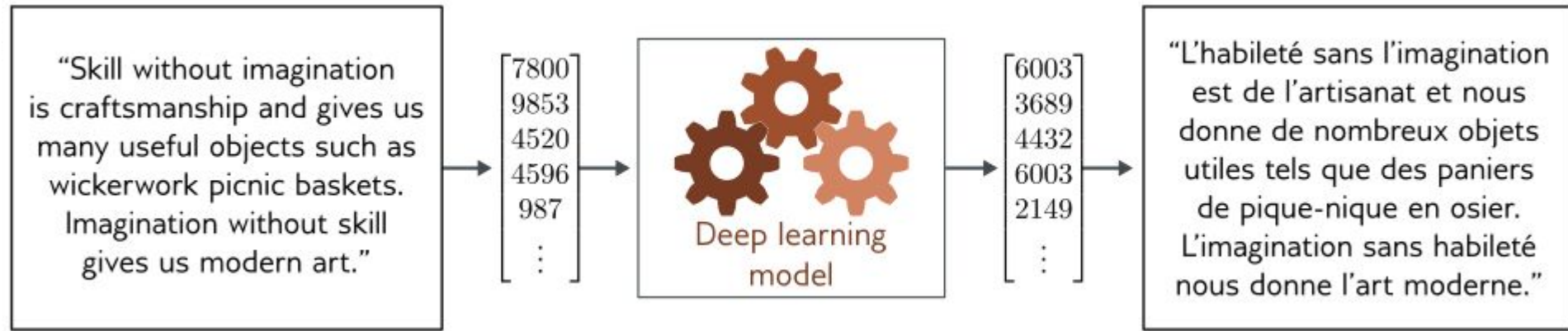
Generalizzazione

Esempi: interpretazione/trascrizione audio&video

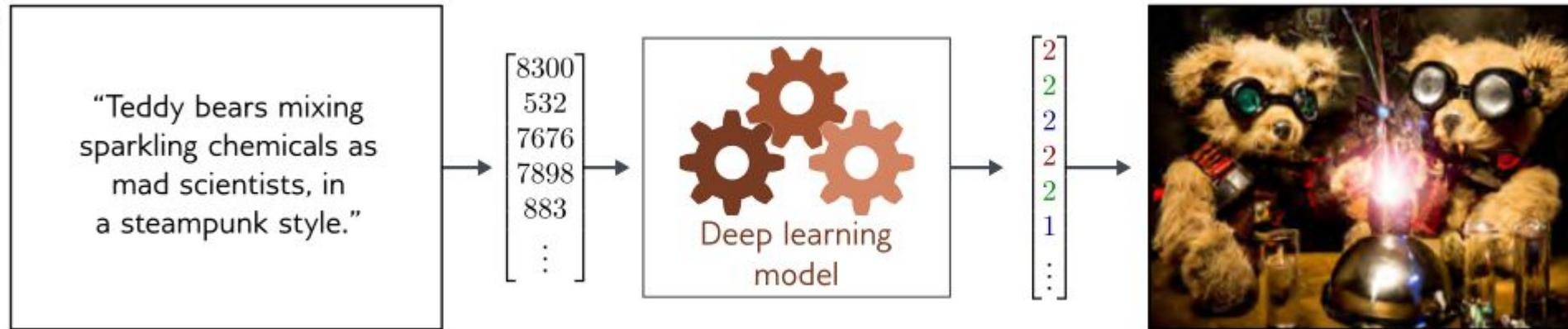


(e viceversa)

Esempi: traduzione, OCR



Esempi: generazione immagini e testo



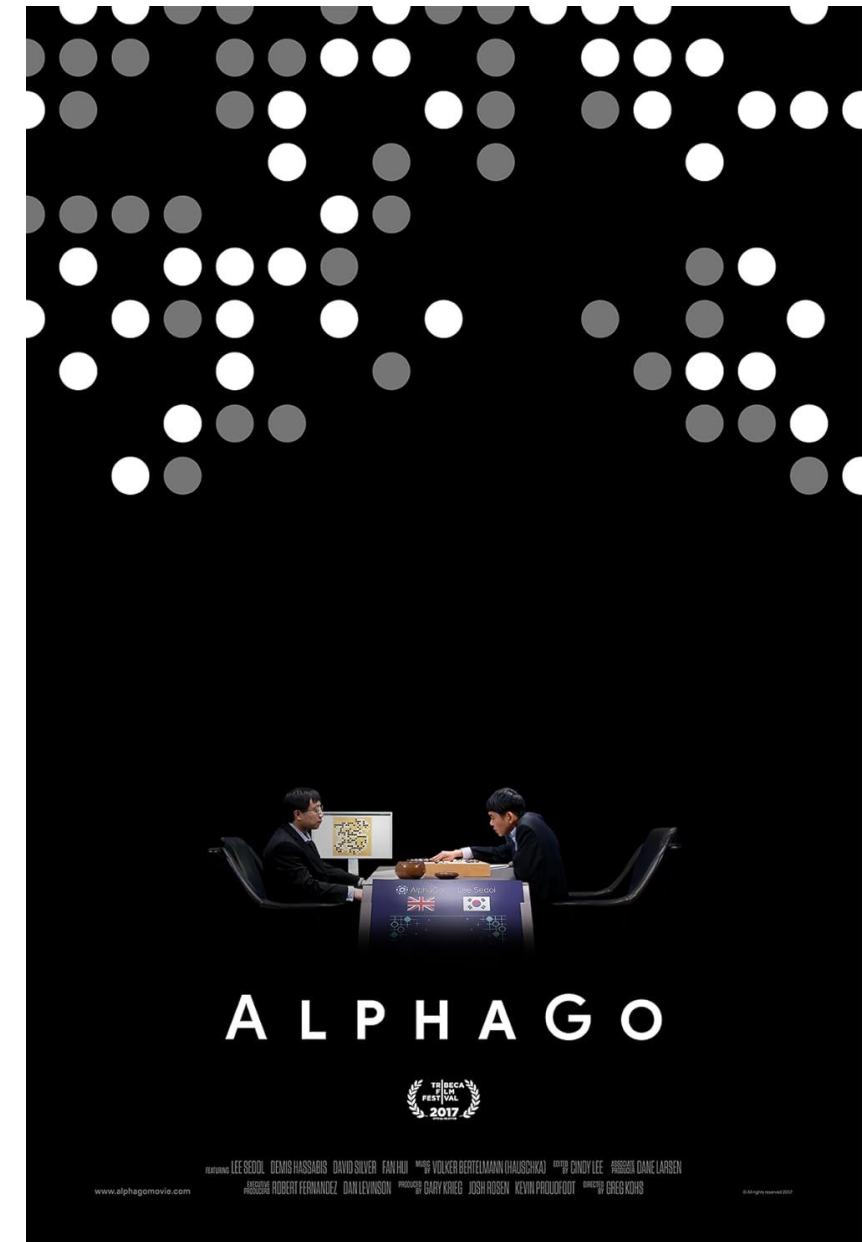
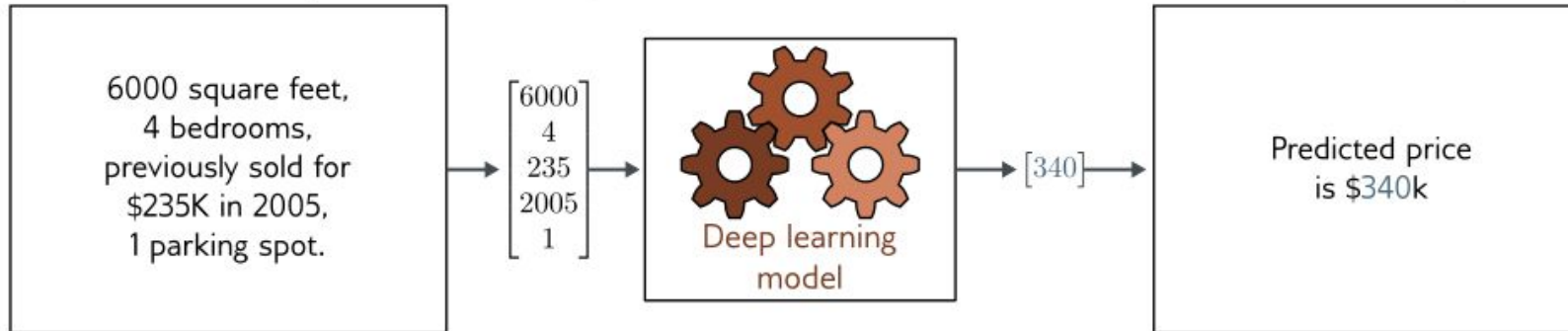
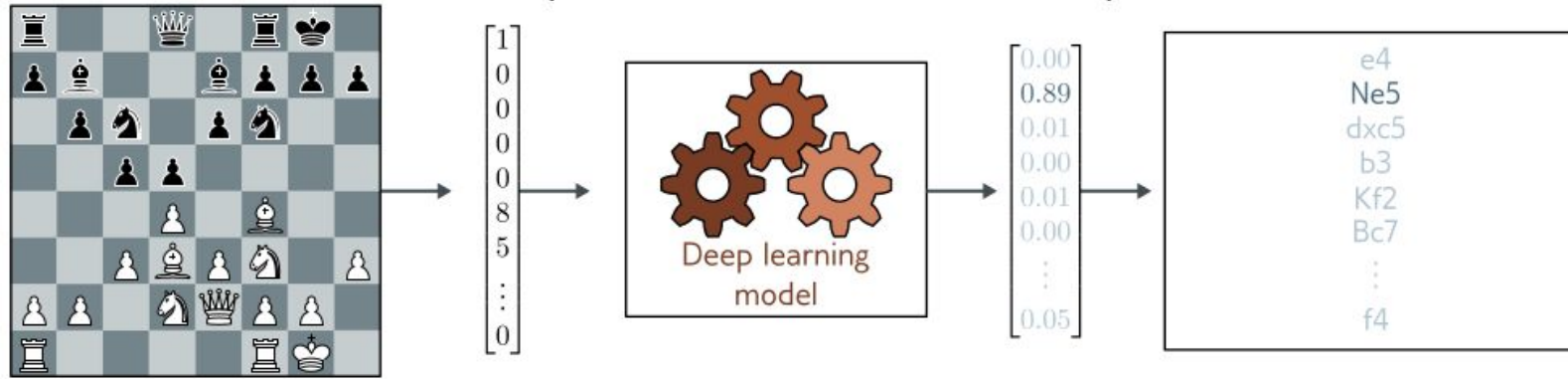
Suggerisci
d'arte
nord
da un
con v

Suggeriscimi una città d'arte lombarda con aeroporto low-cost facilmente raggiungibile, anche di piccole dimensioni

Una scelta molto interessante sarebbe **Bergamo** ✈️. È una città d'arte *di medie-piccole dimensioni*, con un centro storico magnifico — **Città Alta** — pieno di fascino rinascimentale e medievale, e soprattutto ha l'**aeroporto di Orio al Serio (BGY)**, uno dei principali hub europei per compagnie low-cost come Ryanair. È collegata benissimo sia con Milano (45 min di bus o treno) sia con molte città europee.

visitare

Esempi: presa di decisioni e predizione





Lee Sedol



L'importanza delle reti neurali

(nn.rel)

Reti neurali: una nuova concezione di *software*

L'IA sottosimbolica ci porta ad una nuova concezione del *computer* e della *computazione*. Passiamo da

- **software 1.0**: il programmatore dice alla macchina cosa fare, istruzione per istruzione
a
- **software 2.0**: la macchina impara da sé ad assolvere a un compito preposto (a patto di averla addestrata su di un numero sufficiente di esempi).

L'idea ha almeno '70 anni: l'ha avuta Turing nel '50, quando ha è detto:

<<Invece di elaborare un programma per la simulazione di una mente adulta, perché non proviamo piuttosto a realizzarne uno che simuli quella di un bambino? Se la macchina fosse poi sottoposta ad un appropriato corso d'istruzione, si otterrebbe un cervello adulto.>>



Software 2.0

I sometimes see people refer to neural networks as just "another tool in your machine learning toolbox". They have some pros and cons, they...



Andrej Karpathy

Follow

androidstudio · November 11, 2017 (Updated: December 7, 2022) · Free: Yes

I sometimes see people refer to neural networks as just "another tool in your machine learning toolbox". They have some pros and cons, they work here or there, and sometimes you can use them to win Kaggle competitions.

Unfortunately, this interpretation completely misses the forest for the trees.

Neural networks are not just another classifier, they represent the beginning of a fundamental shift in how we develop software. They are Software 2.0.

Large Language Models (LLM)

(genia.llm)

Large Language Models (LLMs)

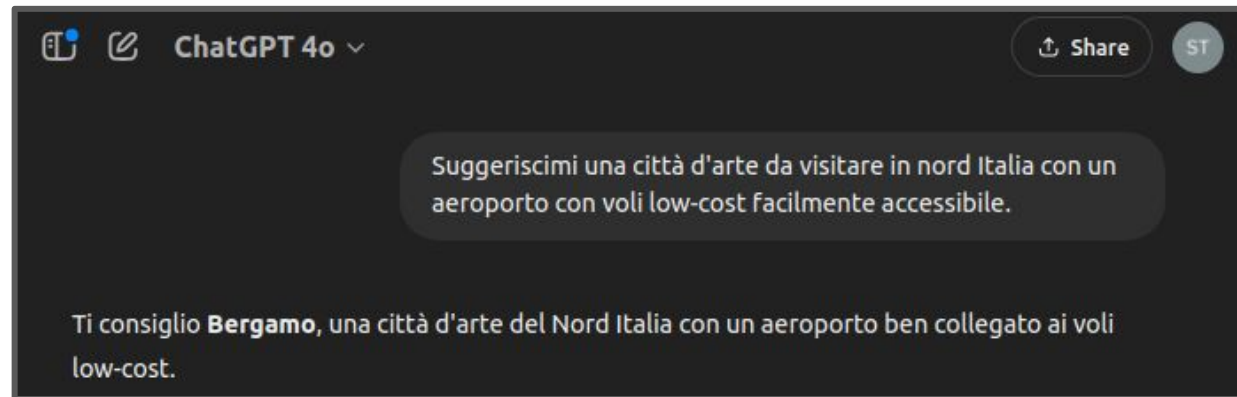


ChatGPT e gli LLM

ChatGPT è un LLM (*large language model*), una rete neurale profonda (cioè funzionalmente strutturata) costituita da miliardi di neuroni artificiali, addestrata all'autocompletamento (quel del vecchio T9 e degli *smartphone* moderni) di letteralmente tutte le frasi trovate in internet:

[Nel mezzo del cammin di nostra] → [vita]

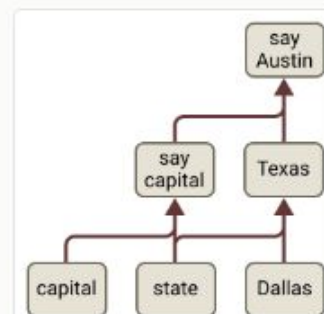
Nell'imparare ad autocompletare, la macchina impara a **parlare...**



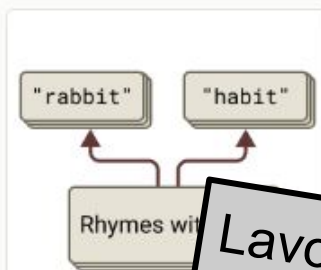
... e, nel farlo, acquisisce **capacità intellettive** (per completare correttamente la frase “L’assassino è...” in un romanzo giallo, bisogna pure aver capito chi è l’assassino!).

On the Biology of a Large Language Model

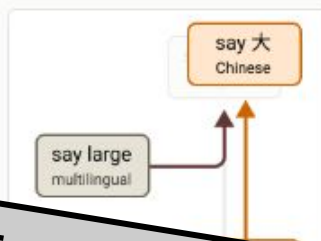
We investigate the internal mechanisms used by Claude 3.5 Haiku — Anthropic's lightweight production model — in a variety of contexts, using our circuit tracing methodology.



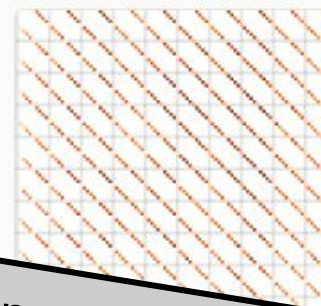
Introductory Example:
Multi-step Reasoning



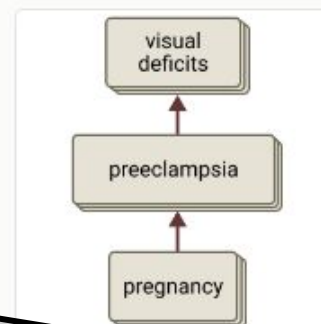
Planning in Reasoning



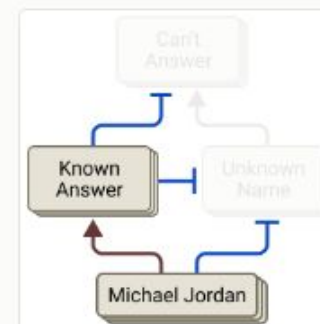
Chain-of-thought
Faithfulness



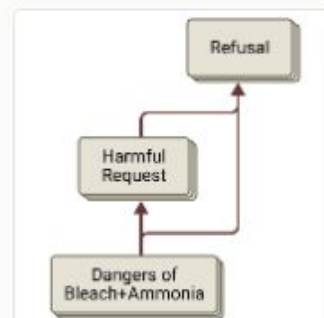
Uncovering Hidden
Goals in a Misaligned
Model



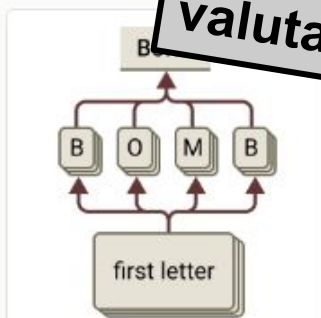
Commonly Observed
Circuit Components
and Structure



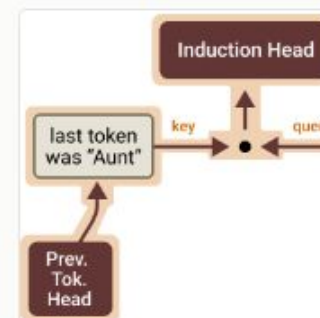
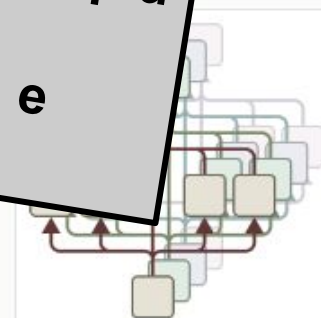
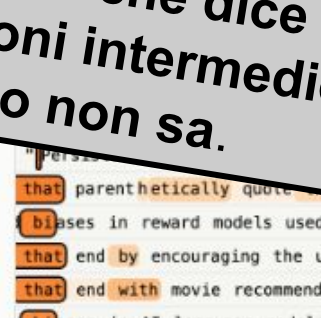
Entity Recognition and
Hallucinations



Refusals



Life of a Jailbreak



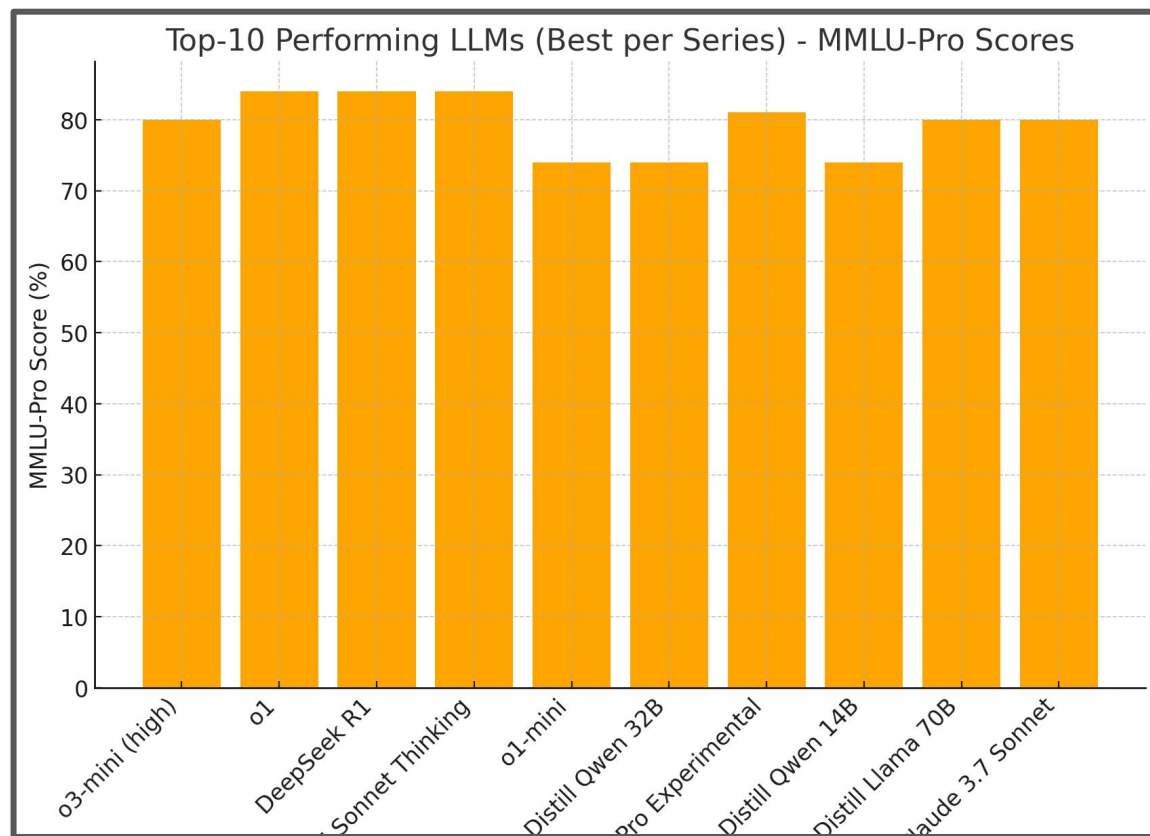
Limitations

Lavori recenti mostrano che, per rispondere, un LLM segue ragionamenti interni su più passi, pianificando quel che dice e arrivando a conclusioni intermedie, e valutando ciò che sa o non sa.

Competenze intellettive

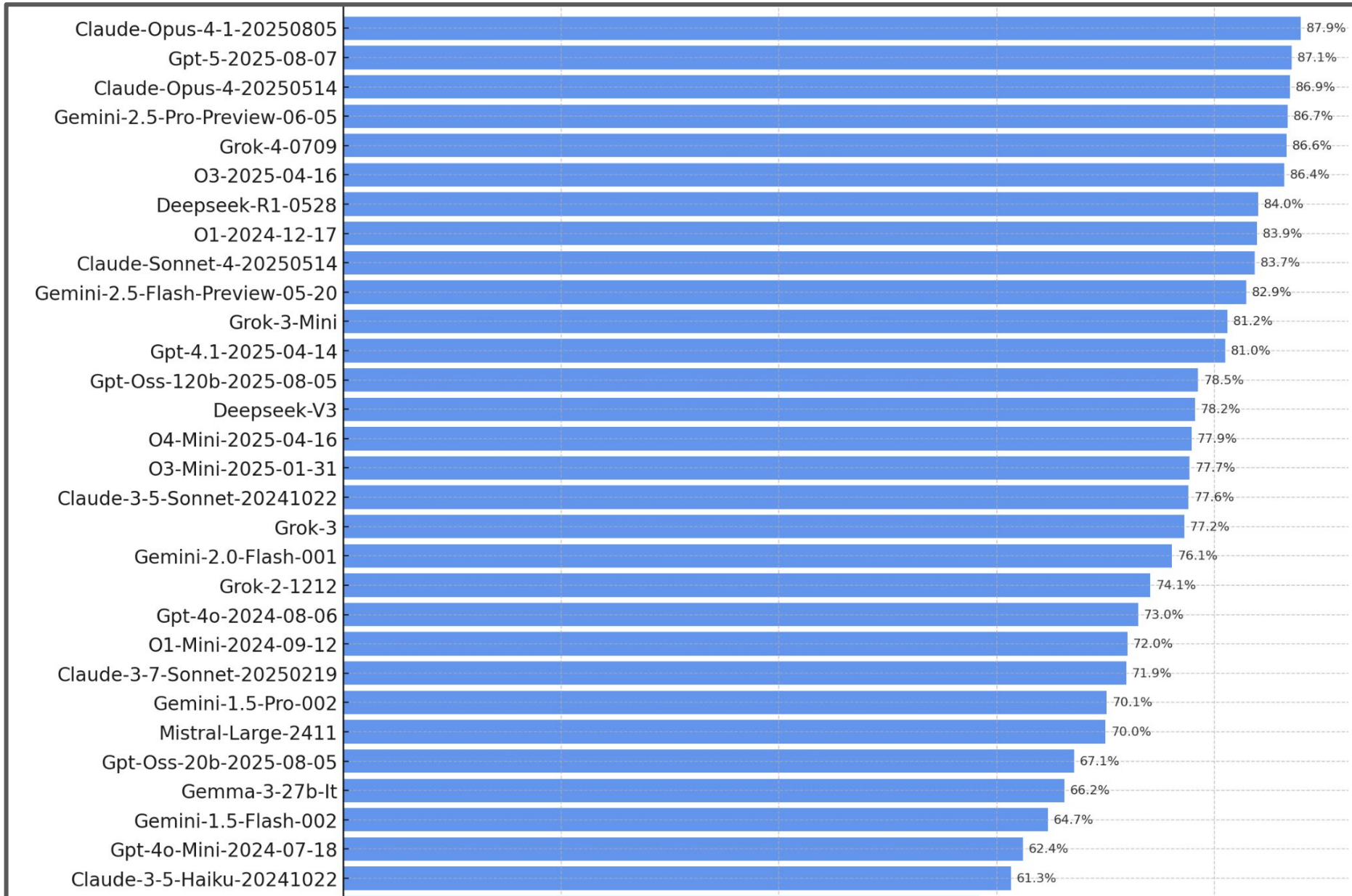
Gli LLM superano brillantemente molti test di “intelligenza” di varia natura.

Un esempio: test **MMLU-Pro** (*Measuring Massive Multitask Language Understanding*): 16.000 domande a scelta multipla che coprono 57 discipline accademiche, tra cui matematica, filosofia, diritto e medicina.



Dati aggiornati

Dati Kaggle, 4 settembre 2025:<https://www.kaggle.com/benchmarks/open-benchmarks/mmlu-pro>

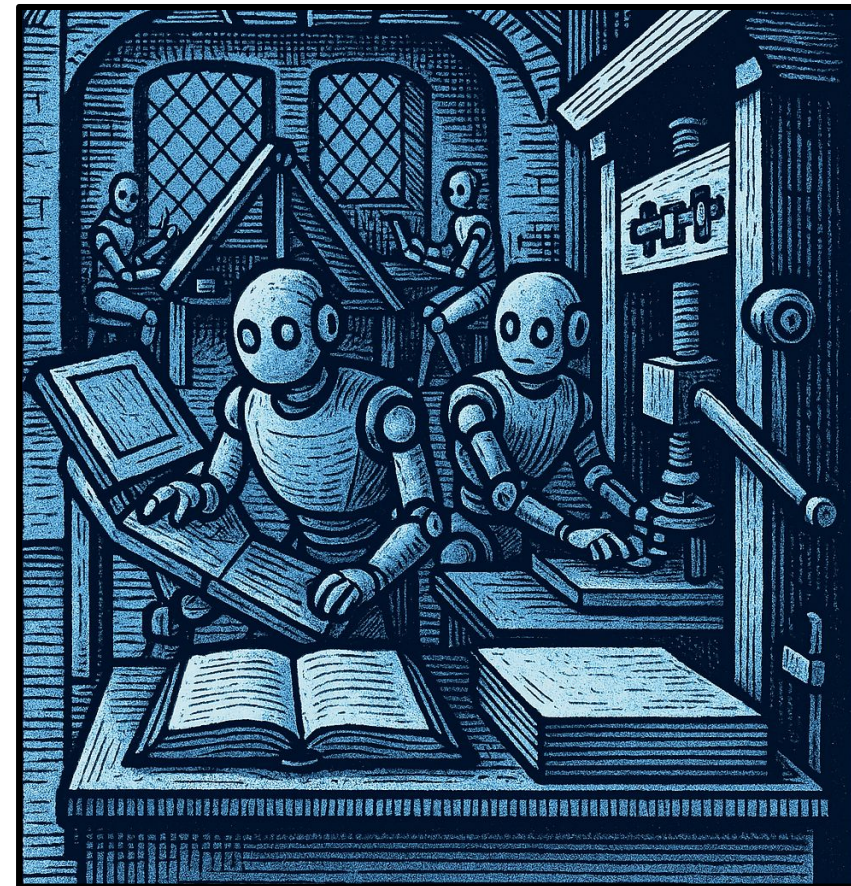


Una rivoluzione epistemologica

ChatGPT e gli LLM sono la più grande rivoluzione nell'interazione dell'uomo col testo scritto dall'invenzione della stampa di Guttenberg (1400).

Prima di ChatGPT, l'interazione col testo (anche quello trovato in internet) era sempre la stessa: lo si leggeva. Con ChatGPT possiamo parlare con il testo.

A ChatGPT possiamo chiedere riassunti e traduzioni, spiegargli quel che abbiamo capito del testo per verificare se l'abbiamo capito bene, chiedere a lui (lei?) di spiegarcelo in meno o più dettagli, magari per metafore...



IA molto potente ma poco sagace Noi serviamo competenti e critici

Stefano Coniglio. «L'Intelligenza Artificiale è la più grande rivoluzione nel sapere dai tempi di Gutenberg. Sistemi versatili, ma talvolta imprecisi: incorrono in errori clamorosi e faticano a far di conto. In alcuni contesti, le loro "allucinazioni" possono avere conseguenze anche gravi. Il ruolo della scuola sarà ancora più fondamentale»



A mio avviso – sostiene l'informatico Stefano Coniglio – gli strumenti messi a nostra disposizione dall'Intelligenza Artificiale moderna (quella di ChatGPT, per intenderci) rappresentano la più grande rivoluzione nelle forme di accesso al sapere dal tempo di Gutenberg, l'inventore della stampa a caratteri mobili. Certo, anche la nascita di Internet ha segnato una svolta. Tuttavia, la nostra fruizione di un testo pubblicato nel Web non è essenzialmente diversa da quella che potremmo avere se lo leggessimo scritto su carta: ora, con ChatGPT e altri sistemi consimili, possiamo "interrogare" il testo, averne una parafrasi o un sunto, chiedere che ne vengano spiegati dei passaggi». Dal 2022, dopo aver condotto a lungo attività di ricerca in Germania e in Inghilterra, Coniglio insegna presso il Dipartimento di Scienze Economiche dell'Università di Bergamo. Gli abbiamo posto alcune domande sulle opportunità e i rischi derivanti dalle applicazioni dell'Intelligenza Artificiale.

Professore, per secoli il desiderio di creare delle «macchine intelligenti» è proceduto di pari passo con i tentativi di «matematizzare» il pensiero umano.

«Sì, in modo da poter elaborare argomentazioni rigorose e inconfutabili. Nel Seicento, Gottfried Wilhelm von Leibniz – che fu anche l'inventore di una delle prime calcolatrici meccaniche – immaginava l'introduzione di un linguaggio simbolico, la *characteristica universalis*, mediante il quale si sarebbero potuti appianare i dissidi su qualsivoglia argomento, pervenendo a conclusioni certe e condivise. Fino agli anni Novanta

Chi è

Apprendimento automatico con reti neurali



UNIVERSITÀ DI BERGAMO

Stefano Coniglio è docente di Informatica e Intelligenza Artificiale presso l'Università degli studi di Bergamo, dove ha preso servizio nel 2022 dopo aver conseguito il dottorato in Tecnologie dell'informazione al Politecnico di Milano (nel 2011) e aver svolto attività di ricerca all'estero, in Germania e nel Regno Unito, per un decennio. I suoi interessi di ricerca vertono sullo sviluppo di modelli e metodi di deep learning per l'apprendimento automatico mediante reti neurali, sulla progettazione e l'analisi di algoritmi e sull'ottimizzazione matematica. Coautore di più di cinquanta articoli scientifici su riviste internazionali e atti di convegno, è intervenuto come relatore in numerose conferenze su temi relativi all'IA.

ASK FT

FINANCIAL TIMES

My Account

Ask FT

YOU ASKED

Summarize Trump's lasted 180 on import tariffs.

Filter FT content from: Last 24 hrs7 days30 days12 months12 Oct 2024 - 10 Apr 2025

ASK FT RESPONSE

Donald Trump has recently made significant changes to his tariff policy. On April 9, 2025, he announced a 90-day pause on many of the tariffs he had introduced just a week earlier on April 2, which he had dubbed "liberation day" ¹⁰. This represents a rapid shift in his trade policy approach.

While the pause applies to many countries, steep tariffs on China remain in place. For most other countries, the tariffs have been reduced to a baseline level of 10% ¹⁰. This indicates a more targeted approach, focusing primarily on trade with China while easing pressure on other trading partners.

The situation with China has escalated further. After the Chinese government announced retaliatory tariffs of 34% on US imports, Trump responded by raising the US tariff rate on most Chinese goods to over 100%. In his April 9 announcement, he stated that this rate would be increased again, pushing it above 125% ¹⁰. This demonstrates a significant intensification of the trade conflict with China.

Despite this partial rollback, the overall impact of Trump's tariff policy remains substantial. According to Yale Budget Lab, even before the latest round of retaliatory tariffs on China, the new duties were expected to raise the US effective tariff rate to at least 22.4% — a level not seen since 1909 ¹⁰.

This rapid shift in policy, from introducing sweeping tariffs to partially pausing them within a week, reflects the volatile nature of Trump's trade strategy and its potential impact on global trade relations.

Was this answer helpful?

Yes

No

SOURCES

ASK FT

FINANCIAL TIMES

SOURCES

Trump tariffs so far

8 DAYS AGO

1 These are the trade levies US President Donald Trump has announced in the 10 weeks since taking office: 20 per cent tariffs on imports from China25 per cent tariffs — with exemptions — on imports from Canada and Mexico25 per cent tariffs on steel and aluminium imports, with an end to previous exemp...

Show full extract

Trump's tariffs in brief: universal levies and targeted retaliation

8 DAYS AGO

2 Donald Trump announced sweeping new tariffs on US imports on Wednesday, taking the global economy into a new era of trade competition. The measures, announced from the White House's Rose Garden during what Trump dubbed "liberation day", include a baseline universal levy as well as targeted reciproc...

Show full extract

Trump tariffs

Donald Trump's tariffs in brief: universal levies and targeted retaliation

7 DAYS AGO

Aime Williams

3 To reduce the US trade deficit, which has been a long-standing goal of Trump's. Officials blamed "massive" and "chronic" US trade deficits for eroding manufacturing capabilities, suppressing wages, and "transferring . . . assets into foreign hands".To force companies to move manufacturing to the US...

Show full extract

5 Donald Trump announced sweeping new tariffs on US imports on Wednesday, taking the global economy into a new era of trade competition. The measures, announced from the White House's Rose Garden during what Trump dubbed "liberation day", include a baseline universal levy as well as targeted reciproc...

Show full extract

6 A 20 per cent tariff has been imposed on the EU, reflecting Trump's long-standing animosity to what he says are the bloc's "unfair trade practices".Trump said goods from China would be hit with a 34 per cent reciprocal tariff. Crucially, this will be on top of an existing 20 per cent tariff, taking ...

Show full extract

7 Mexico and Canada, which have been frequent targets of Trump's attacks in recent weeks, will dodge reciprocal tariffs. Duties of 25 per cent for goods that do not comply with the terms of their 2020 trade deal with the US will remain in place. Cars — sort of. The White House said cars and car parts...

Show full extract



ASK

THE POST

AI

What happened to Charlie Kirk?

Charlie Kirk, a conservative activist and founder of Turning Point USA, was fatally shot at a Utah university on Wednesday. The suspect, 22-year-old Tyler Robinson, has been taken into custody and is being held without bail on suspicion of aggravated murder, obstruction of justice, and discharge of a weapon causing serious bodily harm. Robinson had sent a text message before the shooting saying he had "an opportunity to take out Charlie Kirk" and that he was "going to take it." The motive behind Kirk's killing is still unclear, but authorities believe Robinson has a "leftist ideology." Vice President JD Vance has been honoring Kirk's legacy and has taken on a role in leading the MAGA movement's mourning for the activist. A memorial for Kirk is planned for Sunday in Arizona, and many prominent leaders, including President Donald Trump, are expected to attend.

If you have a follow-up question, feel free to ask.

[Sources](#)



New chat

History

...
More



Sources

5 articles were used to generate the answer.

1. JD Vance to host 'Charlie Kirk Show,' as he mourns a friend and ally

September 15, 2025



[See passage >](#)

2. What to know as Kirk shooting suspect awaits first court appearance

September 15, 2025



[See passage >](#)

3. The political divides after political violence

September 12, 2025



[See passage >](#)

4. 'My kid has seen this. Now what?': Parents reel as Charlie Kirk video goes viral

September 12, 2025



[See passage >](#)

5. He was asking Charlie Kirk a question. Then a shot rang out.

September 12, 2025



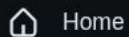
[See passage >](#)

Charlie Kirk, a conservative activist and founder of Turning Point USA, was fatally shot at a Utah university on Wednesday. The suspect, 22-year-old Tyler Robinson, has been taken into custody and is being held without bail on suspicion of aggravated murder, obstruction of justice, and discharge of a weapon causing serious bodily harm. Robinson had sent a text message before the shooting saying he had "an opportunity to take out Charlie Kirk" and that he was "going to take it." The motive behind Kirk's killing is still unclear, but authorities believe Robinson has a "leftist ideology." Vice President JD Vance has been honoring Kirk's legacy and has taken on a role in leading the MAGA movement's mourning for the activist. A memorial for Kirk is planned for Sunday in Arizona, and many prominent leaders, including President Donald Trump, are expected to attend.

If you have a follow-up question, feel free to ask.

[Sources](#)





Home



Popular

Answers **BETA**

Explore



All

GAMES ON REDDIT

MODERATION



Mod Queue



Mod Mail



r/Mod



r/the_blackest_... ☆

CUSTOM FEEDS



Create Custom Feed

RECENT

COMMUNITIES

 New question

What did "scientists" think of the book "The MANIAC" by Labatut?



Sources: r/RSbookclub, r/IReadABookAndAdoredIt, r/ThomasPynchon +4 more >

"The MANIAC" by Benjamin Labatut is a fascinating blend of fiction and non-fiction that explores the life and mind of John von Neumann and other scientific luminaries. Here's a summary of what scientists and readers in general think about the book:

Positive Reception

- **Engaging and Gripping:** Many readers found the book to be incredibly engaging and thought-provoking. One Redditor mentioned, "I read this like a junky over the course of two days." >
- **Insightful and Inspiring:** The book is praised for its deep insights into the lives and motivations of scientists. "This is an inspiring book. It talks a lot about the painstaking efforts that go into research." >
- **Unique Narrative Style:** Labatut's unique narrative style, which mixes fictionalized accounts with real historical events, is highly appreciated. "It was unlike anything I've read before and is still simmering in my brain." >

Critical Views

- **Fictionalization Concerns:** Some readers were critical of the book's fictional elements, especially when it came to the portrayal of scientists. "The book is more literature than science history to me." >
- **Repetitive Themes:** A few readers felt that the book was repetitive and lacked originality. "I think he does a sort of 'Great Men of History' thing with mathematicians that's flawed." >
- **Intellectual Fawning:** One reader mentioned that the book felt like "intellectual fawning/fan fiction," which might be a turn-off for those looking for a purely factual account. "While I really enjoyed it, she found it a bit too intellectually fawning/fan fiction like." >

Ask a followup



There are many pitfalls in modern AI, which is “happy to make up citations,” said Steven Piantadosi, a professor in psychology and neuroscience at the University of California at Berkeley.

“The problem with current AI is that it’s not trustworthy, so it’s just based on statistical associations and dependencies,” he said. “It has no notion of ground truth, no notion of ... a rigorous logical or statistical argument. It has no notions of evidence and how strongly to weigh one kind of evidence versus another.”

The Post previously reported that the document stretched the boundaries of science with some of its conclusions. Several sections offer misleading representations of findings in scientific papers.

What readers are saying

The comments overwhelmingly criticize the use of artificial intelligence (AI) in generating the White House's MAHA Report, highlighting issues such as inaccurate citations and misinterpretation of data. Many commenters express skepticism about AI's reliability in producing... [Show more](#)

This summary is AI-generated. AI can make mistakes and this summary is not a replacement for reading the comments.

Comments 2,744

Highlights ✨

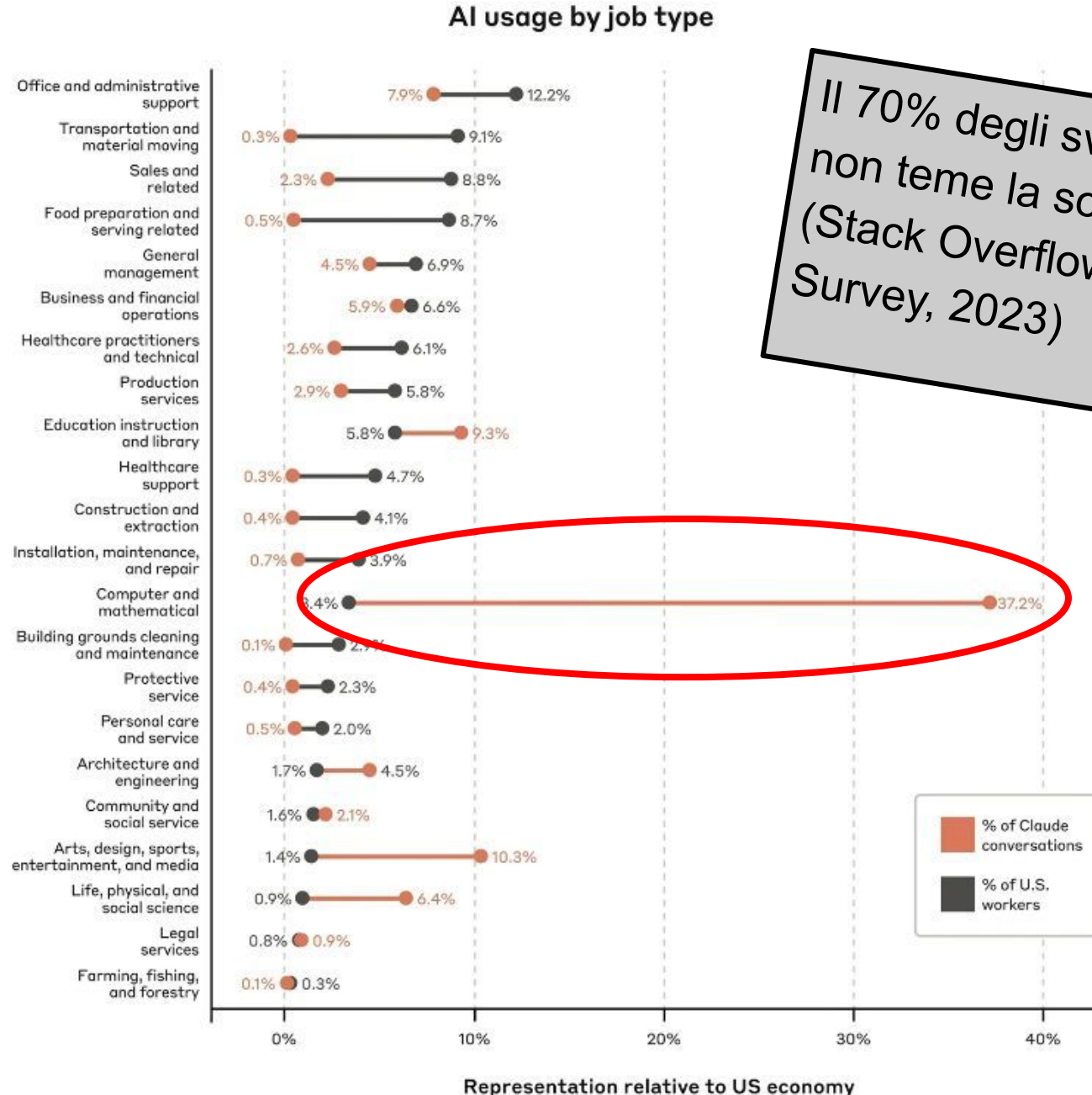
Summary

The comments overwhelmingly criticize the use of artificial intelligence (AI) in generating the White House's MAHA Report, highlighting issues such as inaccurate citations and misinterpretation of data. Many commenters express skepticism about AI's reliability in producing scientific reports, emphasizing the need for human oversight and expertise. The report is seen as emblematic of broader incompetence within the Trump administration, with some suggesting that the reliance on AI reflects a lack of genuine expertise and accountability. [Show less](#)

ⓘ This summary is AI-generated.

Software 1.5?

Forte impatto: industria del software



Il 70% degli sviluppatori software non teme la sostituzione (Stack Overflow Developer Survey, 2023)

The "vibe coding" mind virus explained...

**“in 12 months,
nearly all code may be
generated by AI”**

▶ ▶ 🔊 0:48 / 4:46

⏸ ⌨ ⚙ ⛶

Vibe coding



Andrej Karpathy

@karpathy



There's a new kind of coding I call "vibe coding", where you fully give in to the vibes, embrace exponentials, and forget that the code even exists. It's possible because the LLMs (e.g. Cursor Composer w Sonnet) are getting too good. Also I just talk to Composer with SuperWhisper so I barely even touch the keyboard. I ask for the dumbest things like "decrease the padding on the sidebar by half" because I'm too lazy to find it. I "Accept All" always, I don't read the diffs anymore. When I get error messages I just copy paste them in with no comment, usually that fixes it. The code grows beyond my usual comprehension, I'd have to really read through it for a while. Sometimes the LLMs can't fix a bug so I just work around it or ask for random changes until it goes away. It's not too bad for throwaway weekend projects, but still quite amusing. I'm building a project or webapp, but it's not really coding - I just see stuff, say stuff, run stuff, and copy paste stuff, and it mostly works.

12:17 AM · Feb 3, 2025 · 4.9M Views



leo  @leojr94_ · Mar 15

my saas was built with Cursor, zero hand written code

AI is no longer just an assistant, it's also the builder

Now, you can co

P.S. Yes, people

 81



leo  @leojr94_ · Mar 17

guys. i'm under attack

even though I built my SaaS using Cursor

random things are happening on api keys, people
bypassing the subscription, creating issues hit on db

as you know, I'm not technical so this is taking me longer than usual to figure out

for now, I will stop sharing what I do publicly on X

there are just some weird ppl out there

 642

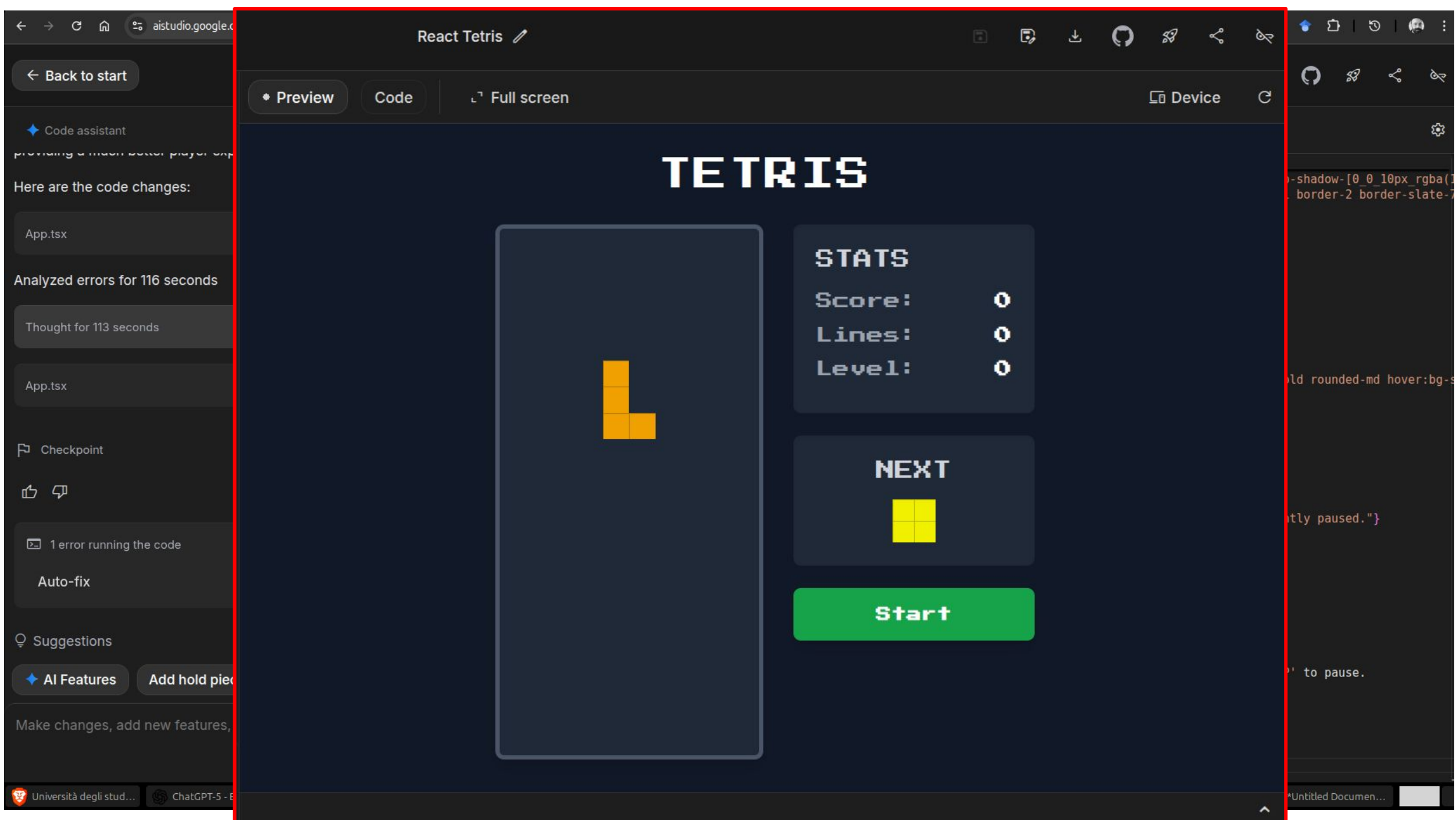
 1K

 6.2K

 2.2M

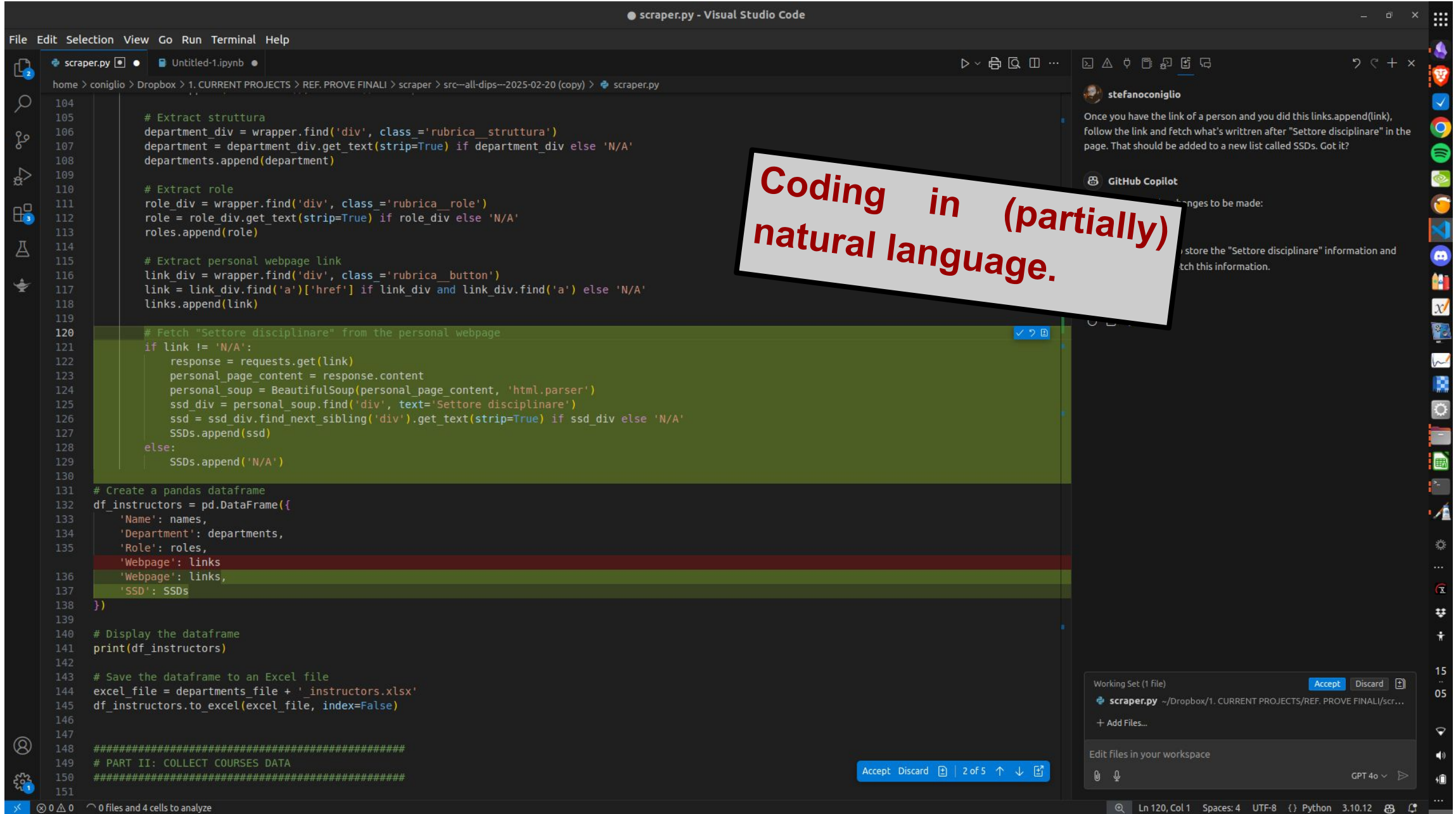


Demo to product gap



Chat-oriented programming (COP)

Copilot in
VSCode



The image shows a Visual Studio Code editor window with a Python script named `scraper.py`. The script is a web scraper that extracts data from a webpage and saves it to an Excel file. The code is written in a mix of English and Italian comments. A large, tilted text box is overlaid on the code, stating: "Coding in (partially) natural language." The right sidebar shows the GitHub Copilot chat interface, which is currently displaying a conversation with a user named `stefanoconiglio`. The chat history shows a previous message about appending links and a new message asking for help with extracting "Settore disciplinare" information. The bottom status bar indicates the current file is `scraper.py` at line 120, column 1, using UTF-8 encoding, and the Python interpreter is set to 3.10.12.

```
104
105
106 # Extract struttura
107 department_div = wrapper.find('div', class_='rubrica_struttura')
108 department = department_div.get_text(strip=True) if department_div else 'N/A'
109 departments.append(department)
110
111 # Extract role
112 role_div = wrapper.find('div', class_='rubrica_role')
113 role = role_div.get_text(strip=True) if role_div else 'N/A'
114 roles.append(role)
115
116 # Extract personal webpage link
117 link_div = wrapper.find('div', class_='rubrica_button')
118 link = link_div.find('a')['href'] if link_div and link_div.find('a') else 'N/A'
119 links.append(link)
120
121 # Fetch "Settore disciplinare" from the personal webpage
122 if link != 'N/A':
123     response = requests.get(link)
124     personal_page_content = response.content
125     personal_soup = BeautifulSoup(personal_page_content, 'html.parser')
126     ssd_div = personal_soup.find('div', text='Settore disciplinare')
127     ssd = ssd_div.find_next_sibling('div').get_text(strip=True) if ssd_div else 'N/A'
128     SSDs.append(ssd)
129 else:
130     SSDs.append('N/A')
131
132 # Create a pandas dataframe
133 df_instructors = pd.DataFrame({
134     'Name': names,
135     'Department': departments,
136     'Role': roles,
137     'Webpage': links,
138     'SSD': SSDs
139 })
140
141 # Display the dataframe
142 print(df_instructors)
143
144 # Save the dataframe to an Excel file
145 excel_file = departments_file + '_instructors.xlsx'
146 df_instructors.to_excel(excel_file, index=False)
147
148 #####
149 # PART II: COLLECT COURSES DATA
150 #####
151
```

Coding in (partially) natural language.

stefanoconiglio

Once you have the link of a person and you did this `links.append(link)`, follow the link and fetch what's written after "Settore disciplinare" in the page. That should be added to a new list called SSDs. Got it?

GitHub Copilot

Changes to be made:

store the "Settore disciplinare" information and fetch this information.

Working Set (1 File) Accept Discard

`scraper.py` ~/Dropbox/1. CURRENT PROJECTS/REF. PROVE FINALI/scr...

+ Add Files...

Edit files in your workspace

Accept Discard 2 of 5

GPT 4o

Someone reported this bug. Can you figure out if real and fix + add test if so ?

name_override not respected when use_docstring_info=False due to missing parentheses in func_name assignment in function_schema . Issue #860

Completed >

Summary

- Ensured that an override name is applied even if docstrings are ignored by enclosing the conditional expression with parentheses `()`
- Added a regression test verifying that name overrides work when docstring parsing is disabled `()`

Testing

- ☒ make format
- ☒ make lint
- ☒ make mypy
- ☒ make tests

Files (2)

src/agents/function_schema.py +2 -1

tests/test_function_schema.py +12 -0

tests/test_function_schema.py

+12 -0

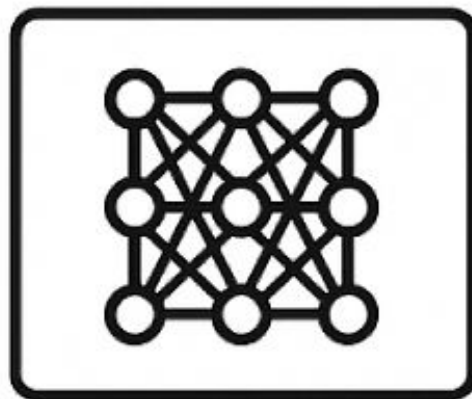
```
@@ -417,25 +417,37 @@ def test_var_positional_tuple_annotation():
417 417
418 418     def test_var_keyword_dict_annotation():
419 419         # Case 3:
420 420         # When a function has a var-keyword parameter annotated with a dict type,
421 421         # function_schema() should convert it into a field with type Dict[<key>, <value>].
422 422         def func(**kwargs: dict[str, int]):
423 423             return kwargs
424 424
425 425         fs = function_schema(func, use_docstring_info=False, strict_json_schema=False)
426 426
427 427         properties = fs.params_json_schema.get("properties", {})
428 428         # The name of the field is "kwargs", and it's a JSON object i.e. a dict.
429 429         assert properties.get("kwargs").get("type") == "object"
430 430         # The values in the dict are integers.
431 431         assert properties.get("kwargs").get("additionalProperties").get("type") == "integer"
432 432
433 433     def test_schema_with_mapping_raises_strict_mode_error():
434 434         """A mapping type is not allowed in strict mode. Same for dicts. Ensure we raise a UserError."""
435 435
436 436         def func_with_mapping(test_one: Mapping[str, int]) -> str:
437 437             return "foo"
438 438
439 439         with pytest.raises(UserError):
440 440             function_schema(func_with_mapping)
441 441
442 442 +
443 443 +
444 444 + def test_name_override_without_docstring() -> None:
445 445 +     """name_override should be used even when not parsing docstrings."""
446 446 +
447 447 +     def foo(x: int) -> int:
448 448 +         return x
449 449 +
450 450 +     fs = function_schema(foo, use_docstring_info=False, name_override="custom")
451 451 +
452 452 +     assert fs.name == "custom"
453 453 +     assert fs.params_json_schema.get("title") == "custom_args"
```

SOFTWARE 1.0

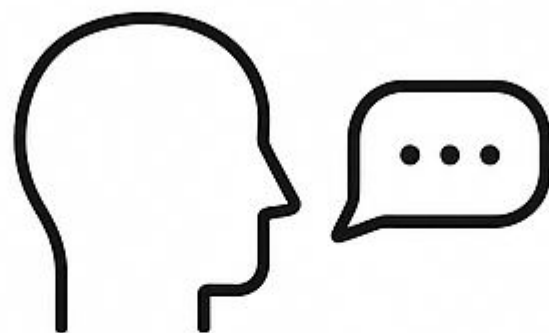


```
if (x > 0) {  
  y = 1;  
} else {  
  y = -1;  
}
```

SOFTWARE 2.0



SOFTWARE 1.5



```
if (x > 0) {  
  y = 1;  
} else {  
  y = -1;  
}
```


La decade degli agenti

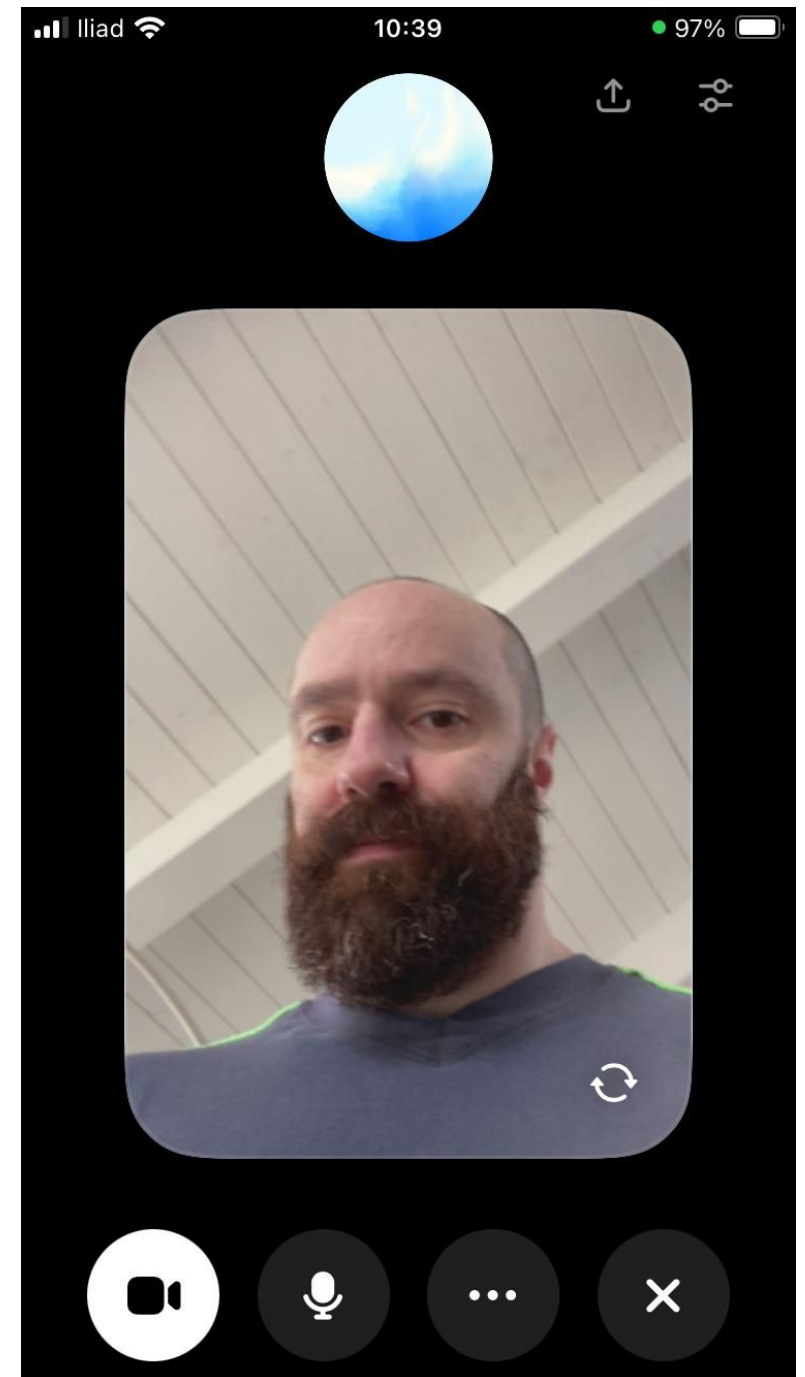
IA ad agenti

Gli LLM stanno diventando il motore inferenziale di agenti più potenti che, oltre a leggere e scrivere:

- parlano modulando dinamicamente il tono della voce
- ascoltano il parlato umano e ne capiscono il tono
- interpretano immagini
- interpretano video
- disegnano con la stessa facilità (e tecnologia) con cui scrivono

ma anche

- navigano in rete
- usano il browser
- scrivono ed eseguono codice al computer (macchine che programmano altre macchine!)



Introducing ChatGPT agent: bridging research and action

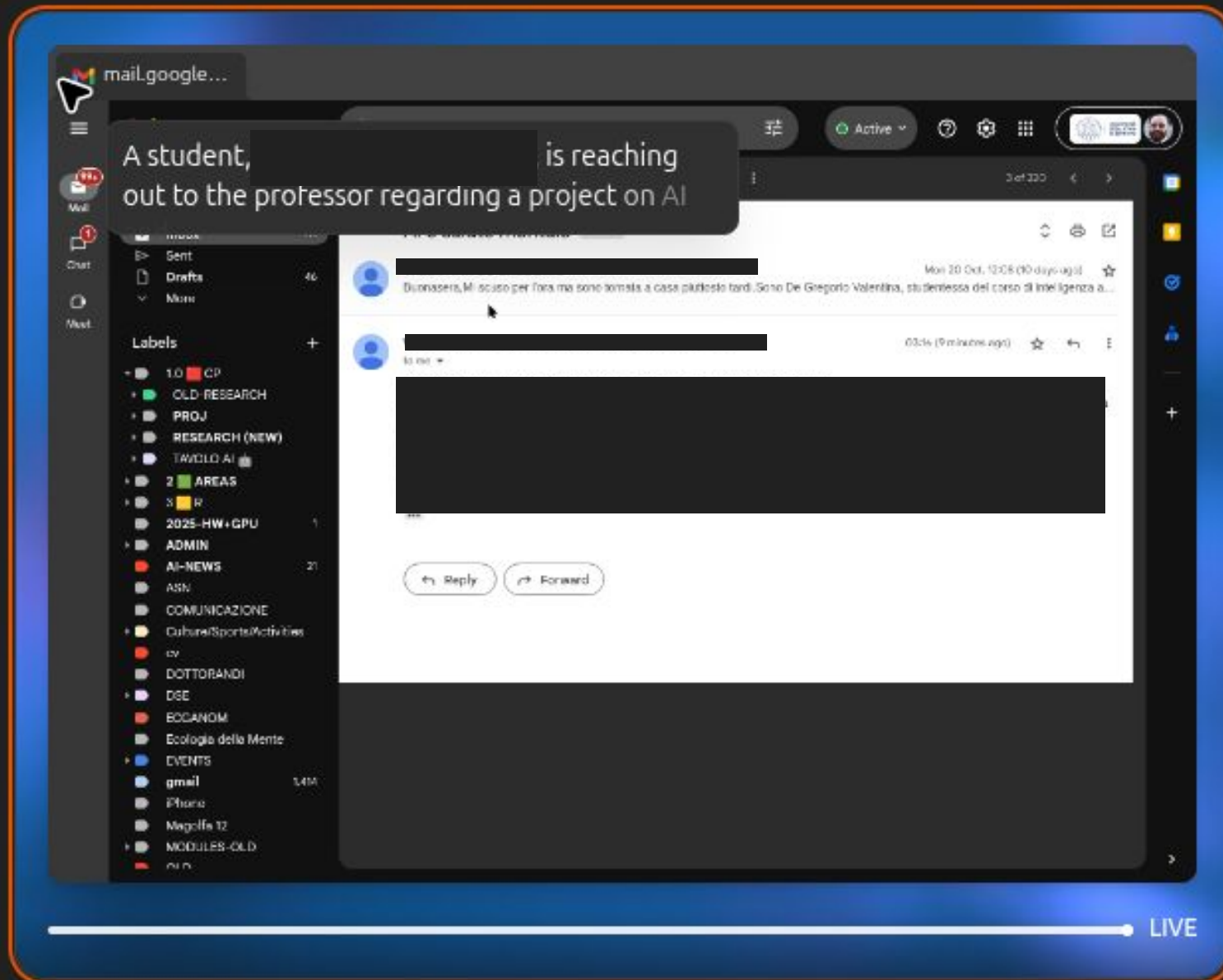
ChatGPT now thinks and acts, proactively choosing from a toolbox of agentic skills to complete tasks for you using its own computer.

Try on ChatGPT ↗

Visiting Expedia to plan your trip

The screenshot shows the Expedia website interface. At the top, there's a navigation bar with the Expedia logo and links for 'Get the app', 'USD', 'Set your privacy', 'Support', 'Link', and 'Sign in'. Below the navigation bar, the search section is active, showing 'Flying from San Francisco, CA (SFO-San...' and 'Flying to New York, NY (NYC>All Airp...'. The departure date is set to 'Aug 4' and the return date is 'Aug 11'. A 'Search' button is visible. Below the search bar, there's a 'Price Tracking' section with a toggle switch. The 'Filter by' section is expanded, showing 'Popular filters' and 'More filters'. The 'Popular filters' section includes 'Nonstop' (selected), 'Carry-on bag included', 'Morning departure', and 'Afternoon arrival'. The 'More filters' section includes 'Stop' (Nonstop (3M), 1 Stop (20)), 'Airlines' (United (24), Delta (14)), and 'From' (100+ airports, 100+ airports, 100+ airports). The 'Recommended' section shows flight options from San Francisco (SFO) to New York City (NYC) with prices starting from \$343. A chat bubble overlay on the right side of the screen contains the text: 'I'm going to look for nonstop routes to NYC as well as flig'.

Student contacting professor for meeting



ⓘ Sensitive data · Something not right? [Stop agent](#)

Computer use (Anthropic)

Announcements

Introducing computer use, a new Claude 3.5 Sonnet, and Claude 3.5 Haiku

Oct 22, 2024 • 5 min read

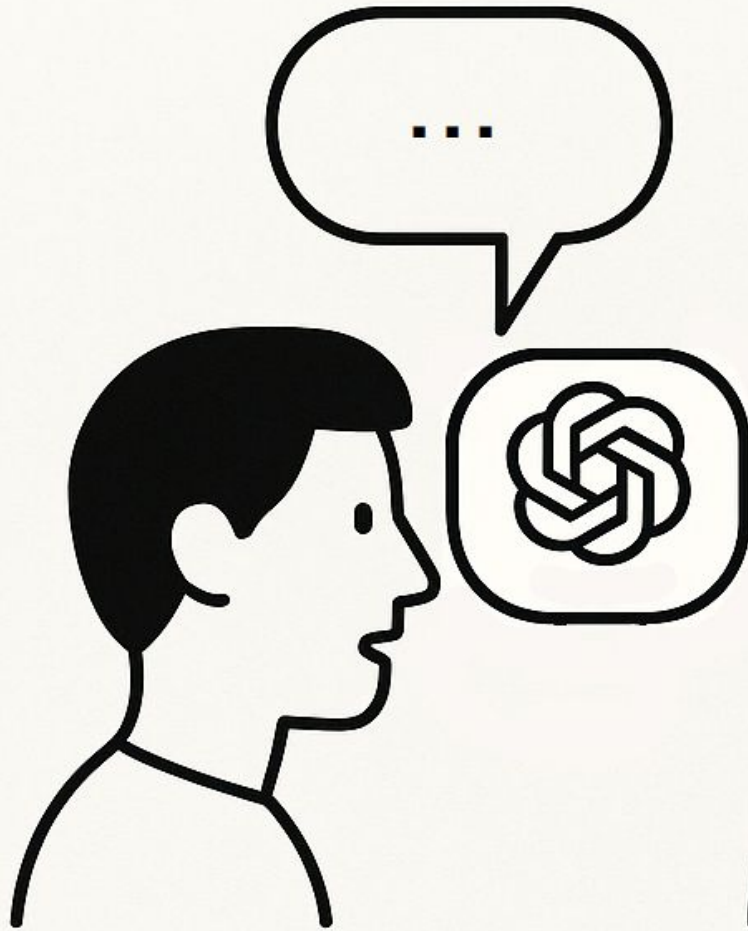


We're also introducing a groundbreaking new capability in public beta: **computer use**. Available today on the API, developers can direct Claude to use computers the way people do—by looking at a screen, moving a cursor, clicking buttons, and typing text. Claude 3.5 Sonnet is the first frontier AI model to offer computer use in public beta. At this stage, it is still experimental—at times cumbersome and error-prone. We're releasing computer use early for feedback from developers, and expect the capability to improve rapidly over time.

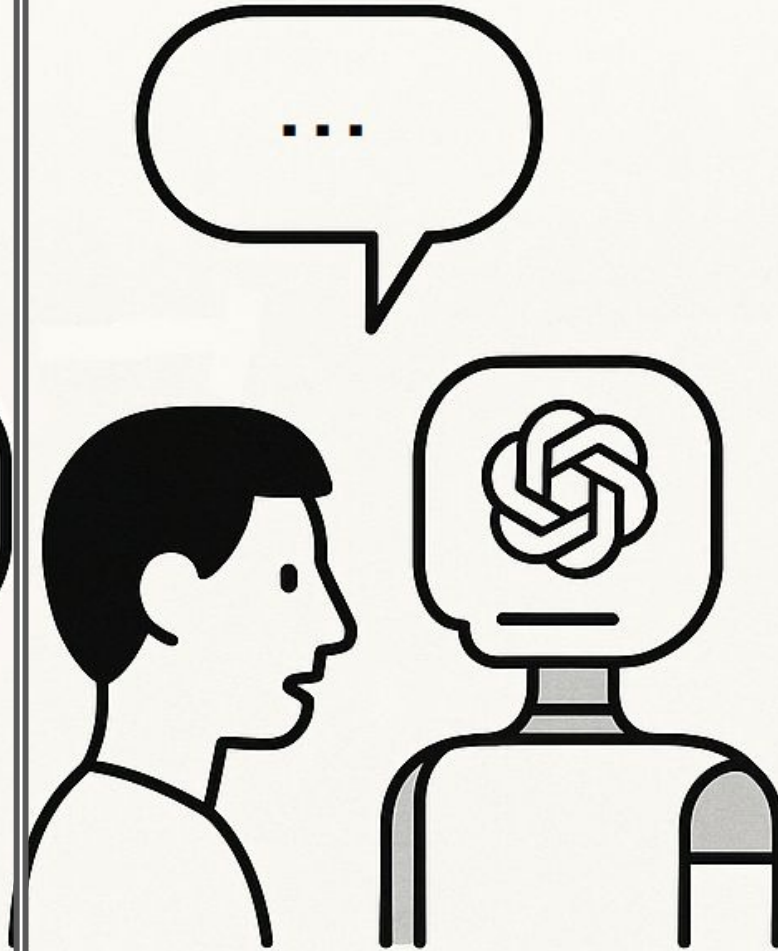
**Works with your screen,
keyboard and mouse**

L'internet del futuro

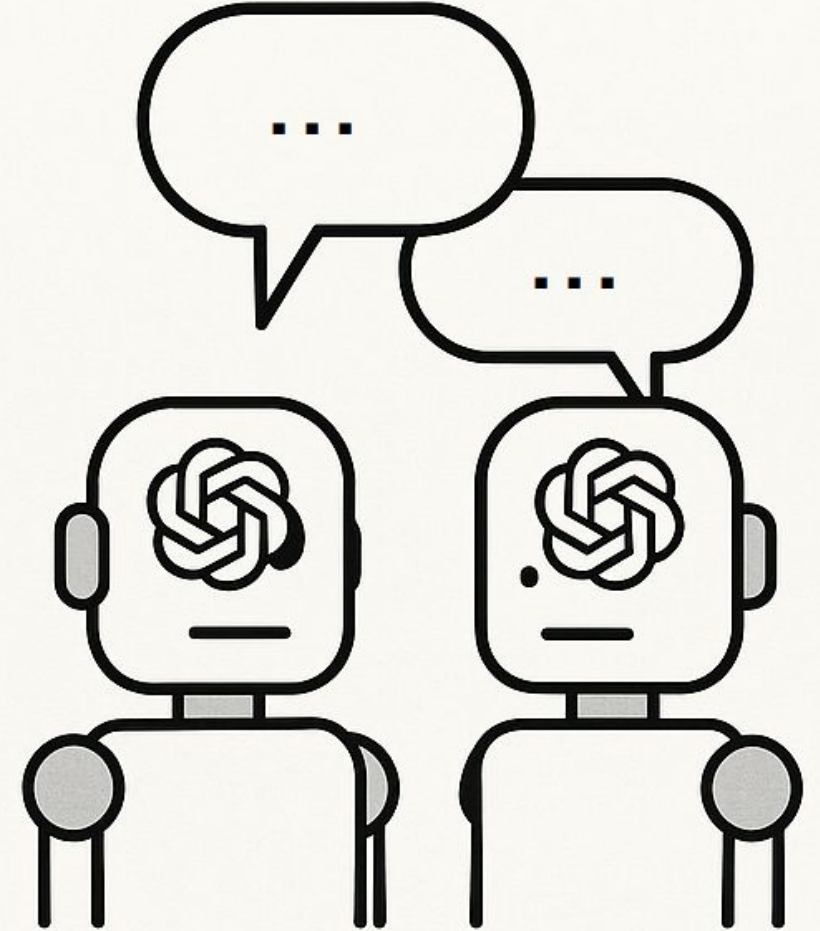
PHASE 1



PHASE 2



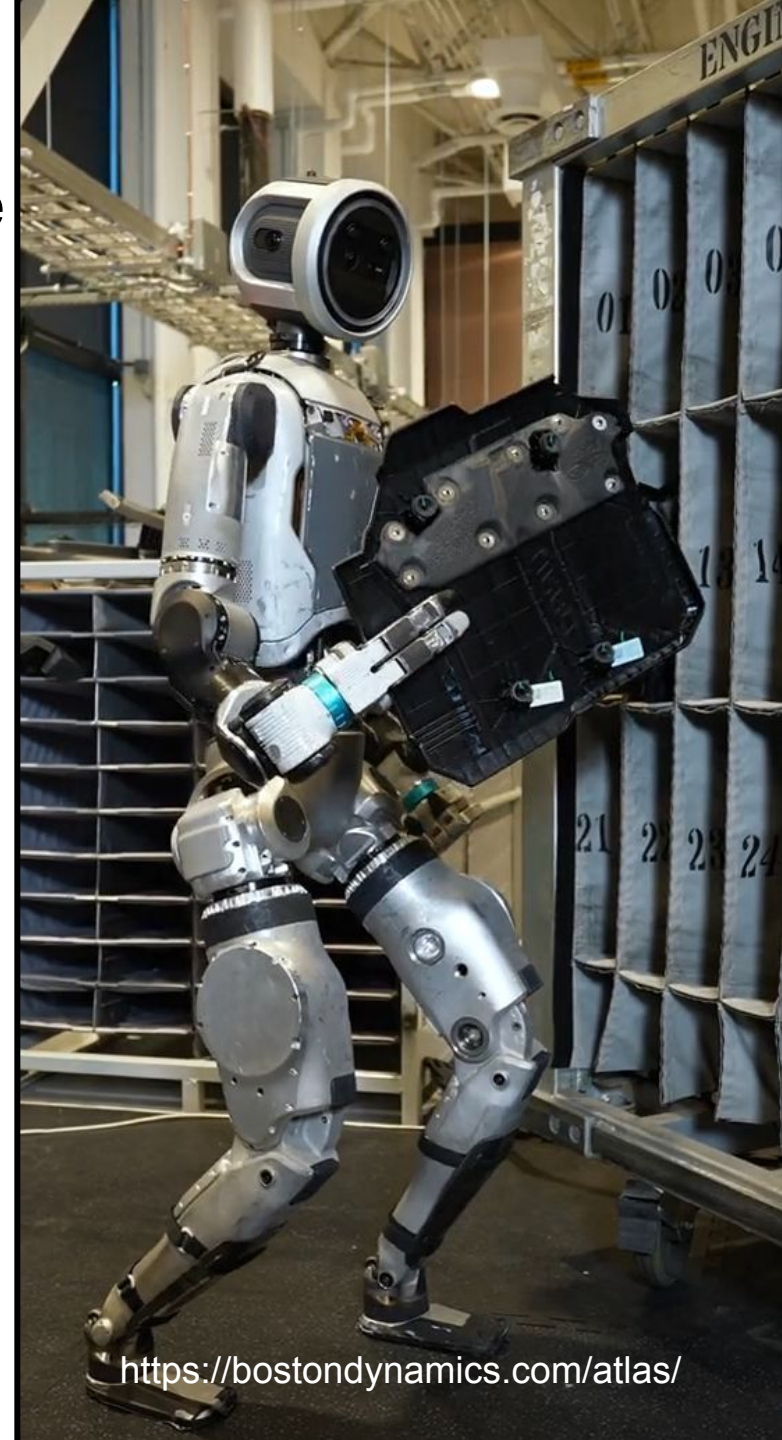
PHASE 3



IA e robotica

L'idea di apprendimento automatico la troviamo oggi anche nei **robot** , che imparano a “fare” senza che il progettista gli dica “come fare”.

Se pensiamo a un bambino: impara a camminare da sé e non perché i genitori gli hanno detto: “metti avanti questo piede, poi metti avanti l'altro”, ma osservando chi già ne è capace e facendo molti tentativi.



Conclusioni

Effetti a breve termine



Oren Etzioni

- Prof. Emerito a U. Washington
- primo studente a laurearsi in CS a Harvard
- CEO Ai2 (centro di ric. no prof in IA fondato da Paul Allen)

Per ora, l'IA degli LLM sta affiancando il lavoratore automatizzando compiti ripetitivi e aumentandone la produttività, permettendo di concentrarsi su attività più “alte”.

0:17

0:20

0:23

0:27

<<Il motivo per cui promuovo l'uso dell'IA sul lavoro non è perché possa sostituire il lavoratore, ma per far capire che il lavoratore che non la sa usare rischia di essere sostituito da uno che ne è capace.>>

ou.

Effetti a medio termine



Artificial intelligence (AI)

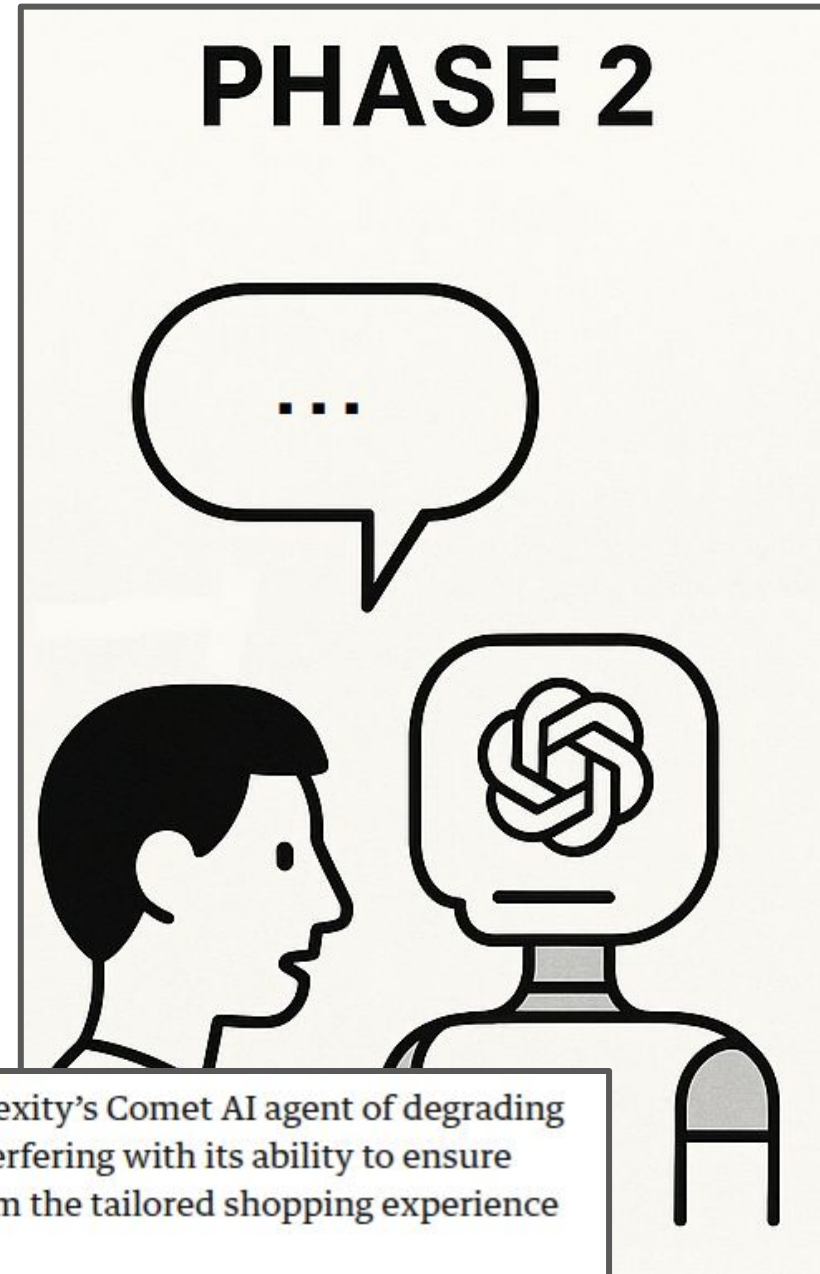
Amazon sues AI startup over browser's automated shopping and buying feature

Amazon accuses Perplexity of covertly accessing customer accounts and disguising AI activity as human browsing

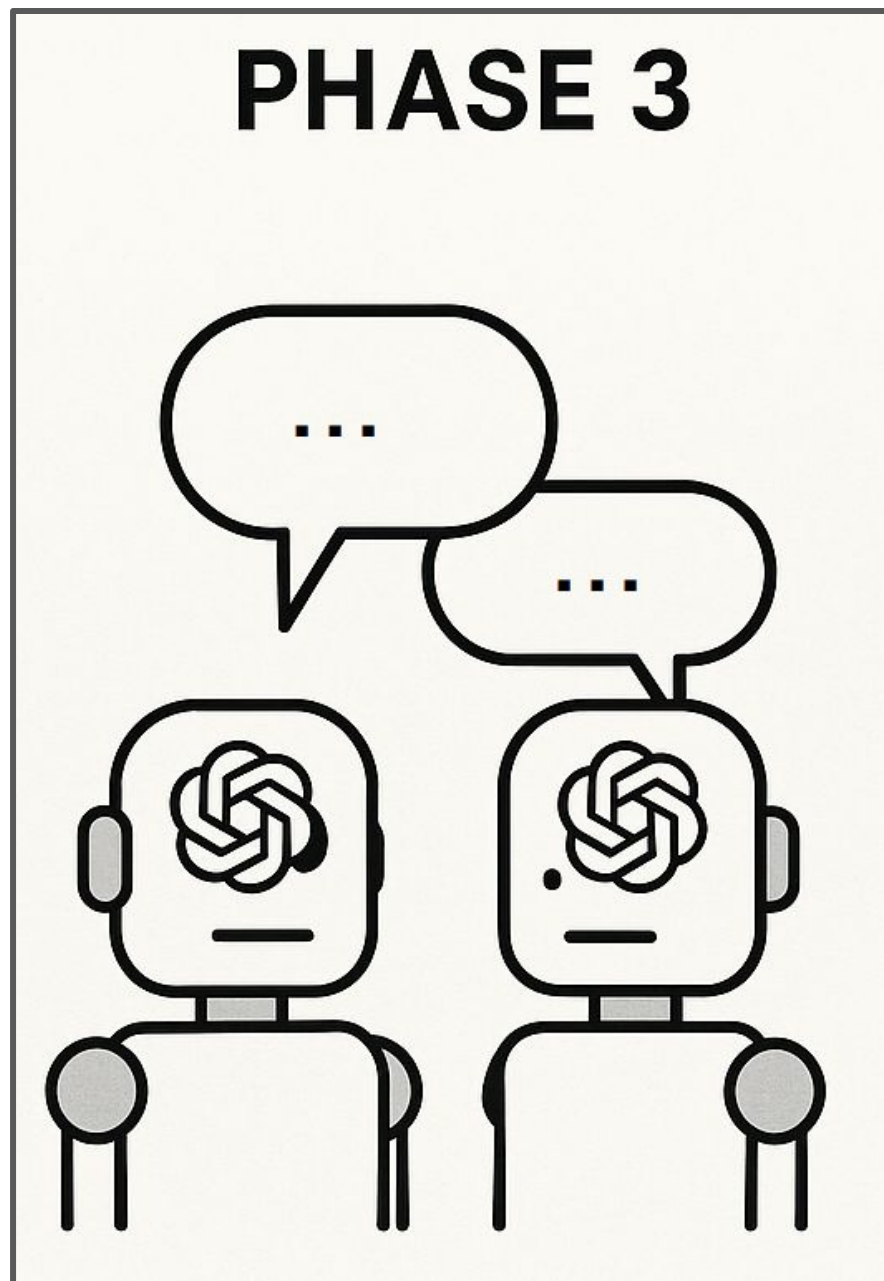
Amazon sued a prominent artificial intelligence startup on Tuesday over a shopping feature in the company's browser, which can automate placing orders for users. [Amazon](#) accused Perplexity AI of covertly accessing customer accounts and disguising AI activity as human browsing.

The image shows the Perplexity logo, which is a stylized white star-like shape on a dark background, and the word "perplexity" in a white, lowercase, sans-serif font.

In the complaint, Amazon accused Perplexity's Comet AI agent of degrading customers' shopping experience and interfering with its ability to ensure customers who use the agent benefit from the tailored shopping experience Amazon curated over decades.



Effetti a lungo termine



Effetti a lungo termine

I lavori più impattati i *white collar jobs*: programmatori, analisti, giuristi, ..

Opinion **AI and jobs**

Wealthy cities may be surprise losers from AI automation

The geographic pattern of digital dislocation's impact may be very different from that of previous waves of automation

JOHN THORNHILL

+ Add to myFT

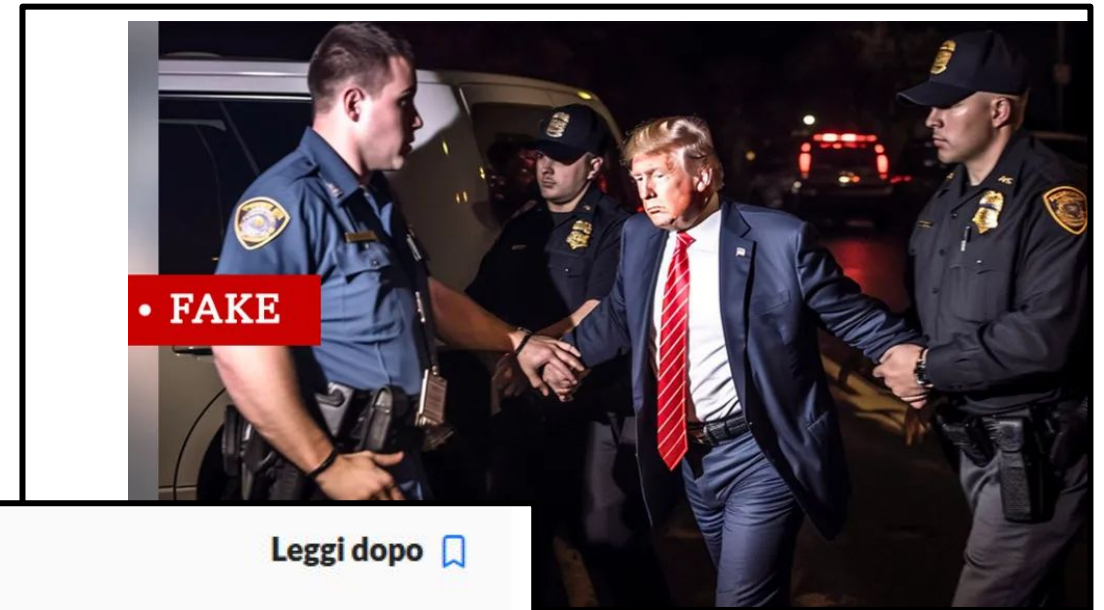


But a recent [study](#) from the [Brookings Institution](#) suggests that the communities most exposed to AI-driven job dislocation will be white-collar information workers. The researchers studied the usage of OpenAI's generative AI tools across more than 1,000 occupations and mapped this against where those jobs were most commonly located.

Their analysis suggests that many [coders, lawyers, financial analysts and bureaucrats](#) in cities such as San Jose, San Francisco, Durham, New York and Washington DC [might want to rethink their futures](#). But non-office-bound workers in places such as Las Vegas, Toledo, Ohio and Fort Wayne, Indiana may be less exposed to AI disruption.

Rischi dell'IA moderna

- L'abbandono del simbolico per il connessionista la rende potente ma a rischio di errori (allucinazioni) – quel che produce va controllato
- può essere usata per *fake news*, furto di identità, *phishing*, ...
- l'uso esteso e pervasivo porta al rischio della decapacitazione.



⚡ POLITICA

Leggi dopo

Truffa del finto Crosetto, anche i Del Vecchio nella trappola. Come funzionava il sistema usato dalla banda

di Franco Luigi Sani

🕒 2 min

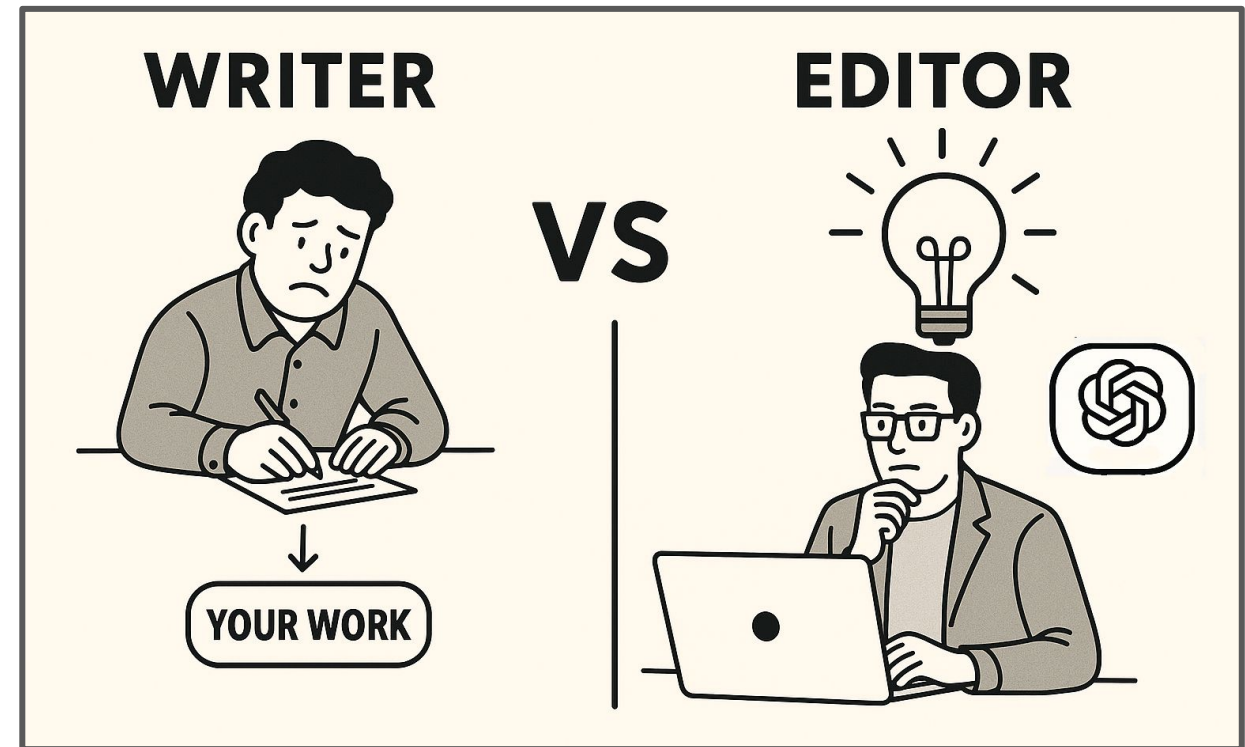
10 febbraio 2025, 11:47

Scrittori vs editori

Con gli LLM, da autori unici del nostro lavoro (**scrittori**) diventiamo manager (**editori**) che valutano e aggiustano i contenuti generati dall'IA (o chiedono all'IA di perfezionarli).

Questo richiede capacità cognitive elevate: **analisi** critica, **valutazione** del lavoro altrui e competenze solide nel fornire **istruzioni** per migliorarlo o correggerlo.

Dobbiamo imparare a scrivere i **prompt** (**e a contestualizzarli**), ma occorrono anche forti competenze professionali specializzate nel settore di appartenenza.



Uno strumento potentissimo per l'autoapprendimento

Citando (non troppo precisamente) il Prof. Michael Crow (rettore della Arizona State University negli Stati Uniti):

“AI and LLMs are a fantastic tool to create a **universal learner** capable of learning anything.”



Grazie per l'attenzione!