

# Logica, ragionamento, e *Large Language Models*: uno sguardo critico

---

Guido Sciavicco<sup>1</sup>

`guido.sciavicco@unife.it`

11 Giugno 2025

<sup>1</sup>Applied Computational Logic and Artificial Intelligence (ACLAI) Laboratory,  
Department of Mathematics and Computer Science, University of Ferrara, Italy

L'universo ci presenta ogni giorno dei **problemi**.  
Alcuni di questi sono **problemi computazionali**, e  
sono quelli interessanti per noi,  
altri non sono computazionali, e qui li ignoreremo.

# Introduzione: problemi ed algoritmi

L'universo ci presenta ogni giorno dei **problemi**.  
Alcuni di questi sono **problemi computazionali**, e  
sono quelli interessanti per noi,  
altri non sono computazionali, e qui li ignoreremo.

**Turing** ci ha insegnato che non tutti i problemi  
computazionali sono risolvibili  
e che anche tra quelli risolvibili da una macchina,  
alcuni sono così complessi intrinsecamente  
da avere soluzioni algoritmiche che saranno per  
sempre da noi irraggiungibili

# Introduzione: problemi ed algoritmi

L'universo ci presenta ogni giorno dei **problemi**.  
Alcuni di questi sono **problemi computazionali**, e  
sono quelli interessanti per noi,  
altri non sono computazionali, e qui li ignoreremo.

**Turing** ci ha insegnato che non tutti i problemi  
computazionali sono risolvibili  
e che anche tra quelli risolvibili da una macchina,  
alcuni sono così complessi intrinsecamente  
da avere soluzioni algoritmiche che saranno per  
sempre da noi irraggiungibili

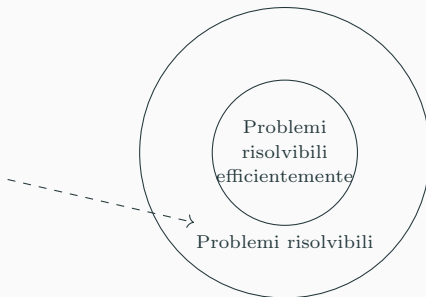
Inoltre, alcuni problemi computazionali  
sono ben definiti ma non hanno  
una caratterizzazione matematica utile.  
Questi problemi non sono neppure giudicabili  
come 'risolvibili' o 'non risolvibili',  
e sfuggono a questa classificazione.

**Esempio:** trovare  
il percorso migliore tra  
due punti



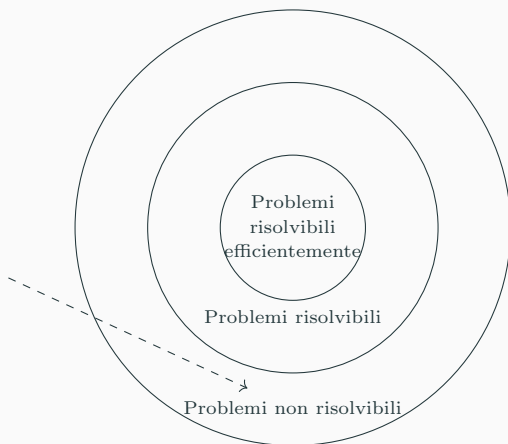
Problemi  
risolvibili  
efficientemente

**Esempio:** definire  
l'orario delle lezioni  
rispettando tutti i  
vincoli

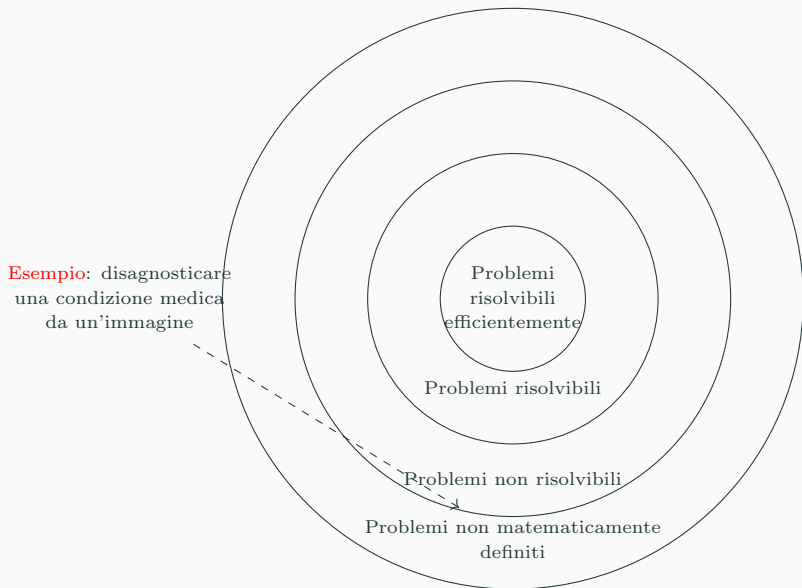


# Introduzione: problemi ed algoritmi

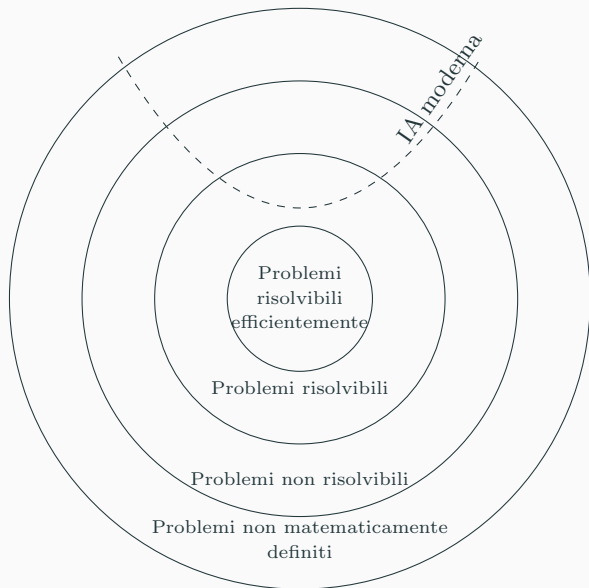
**Esempio:** stabilire  
se un teorema  
matematico è  
corretto



# Introduzione: problemi ed algoritmi



# Introduzione: problemi ed algoritmi



# IA moderna: come funziona?



Nell'informatica classica, fatta principalmente di problemi risolvibili efficientemente queste soluzioni sono **algoritmi** corretti, completi, terminanti, che generano soluzioni in maniera predicibile.

# IA moderna: come funziona?



Nell'informatica classica, fatta principalmente di problemi risolvibili efficientemente queste soluzioni sono **algoritmi** corretti, completi, terminanti, che generano soluzioni in maniera predicibile.

Nella IA moderna ci si basa sui **dati** che possono essere visti come **esempi di soluzione**, che vengono analizzati **statisticamente** e attraverso **meta-algoritmi** diventano essi stessi **algoritmi** creati automaticamente.



# IA moderna: come funziona?

Quindi si parte da dati collezionati  
in un insieme

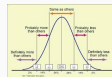


# IA moderna: come funziona?

Quindi si parte da dati collezionati  
in un insieme



Questi vengono analizzati in maniera  
sistematica per comprenderne  
le caratteristiche

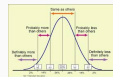


# IA moderna: come funziona?

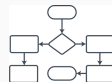
Quindi si parte da dati collezionati  
in un insieme



Questi vengono analizzati in maniera  
sistematica per comprenderne  
le caratteristiche



e viene costruito un algoritmo  
sulla base di queste analisi

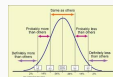


# IA moderna: come funziona?

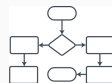
Quindi si parte da dati collezionati  
in un insieme



Questi vengono analizzati in maniera  
sistematica per comprenderne  
le caratteristiche



e viene costruito un algoritmo  
sulla base di queste analisi



Nel caso dei moderni *Large Language Models* (come chatGPT!)  
il dato con cui i meta-algoritmi si allenano  
è composto principalmente da **testo**.

Nel caso dei moderni *Large Language Models* (come chatGPT!)  
il dato con cui i meta-algoritmi si allenano  
è composto principalmente da **testo**.

Questo permette di ottenere dei modelli statistici del linguaggio  
basati su un'idea chiamata **semantica distribuzionale**  
che può essere riassunta in una frase:  
**il significato di un termine dipende dal suo contesto**

Attraverso quest'idea, possiamo imitare un processo intelligente nel quale la macchina dà l'impressione di comprendere un discorso e rispondere in maniera coerente.

Attraverso quest'idea, possiamo imitare un processo intelligente nel quale la macchina dà l'impressione di comprendere un discorso e rispondere in maniera coerente.

"I gatti sono così intuitivi che si può dire  
abbiano un sesto senso."

Attraverso quest'idea, possiamo imitare un processo intelligente nel quale la macchina dà l'impressione di comprendere un discorso e rispondere in maniera coerente.

"I gatti sono così intuitivi che si può dire  
abbiano un sesto **senso**"



Attraverso quest'idea, possiamo imitare un processo intelligente nel quale la macchina dà l'impressione di comprendere un discorso e rispondere in maniera coerente.

”Su questa strada si può procedere in entrambe le direzioni:  
si dice che è a doppio senso.”

Attraverso quest'idea, possiamo imitare un processo intelligente nel quale la macchina dà l'impressione di comprendere un discorso e rispondere in maniera coerente.

”Su questa strada si può procedere in entrambe le direzioni:  
si dice che è a doppio **senso**”



E' possibile però affermare che questo comportamento  
denoti **intelligenza** nel senso in cui siamo abituati?  
Qual è la definizione? Ne esiste una sola? Qual è il legame  
con il processo di ragionamento?

# IA moderna e intelligenza generale ?

E' possibile però affermare che questo comportamento  
denoti **intelligenza** nel senso in cui siamo abituati?  
Qual è la definizione? Ne esiste una sola? Qual è il legame  
con il processo di ragionamento?

**Ragionamento:**[enc. Treccani]"Nella logica antica,  
ogni processo discorsivo della mente o ragione,  
che, muovendo da alcune premesse, perviene ad una conclusione."



Recentemente, molti ricercatori  
si sono posti questa domanda,  
e cercato modi di rispondere

## GSM8K

Recentemente, molti ricercatori  
si sono posti questa domanda,  
e cercato modi di rispondere

*K. Cobbe, V. Kosaraju, M. Bavarian, and others.  
Training verifiers to solve math word problems. 2021*

GSM8K

Math comp.

Recentemente, molti ricercatori  
si sono posti questa domanda,  
e cercato modi di rispondere

*D. Hendrycks, C. Burns, and others.*

*Measuring mathematical problem solving with the MATH dataset. 2021*

GSM8K

Math comp.

Recentemente, molti ricercatori  
si sono posti questa domanda,      Rule Taker  
e cercato modi di rispondere

*P. Clark, O. Tafjord, and K. Richardson.  
Transformers as soft reasoners over language. 2020*

GSM8K

Math comp.

Recentemente, molti ricercatori  
si sono posti questa domanda,     Rule Taker  
e cercato modi di rispondere

Proof Writer

*O. Tafjord, B. Dalvi, and P. Clark.  
ProofWriter: Generating implications, proofs, and abductive  
statements over natural language. 2021*

GSM8K

Math comp.

Recentemente, molti ricercatori  
si sono posti questa domanda,      Rule Taker  
e cercato modi di rispondere

Proof Writer

FOLIO

*S. Han, H. Schoelkopf, Y. Zhao, and others.  
FOLIO: natural language reasoning with first-order logic. 2024*

GSM8K

Math comp.

Recentemente, molti ricercatori  
si sono posti questa domanda,      Rule Taker  
e cercato modi di rispondere

LTLBench

Proof Writer

FOLIO

*W. Tang and V. Belle.*

*Ltlbench: Towards benchmarks for evaluating temporal logic  
reasoning in large language models. 2024*

GSM8K

Math comp.

Recentemente, molti ricercatori  
ProntoQA si sono posti questa domanda, Rule Taker  
e cercato modi di rispondere

LTLBench

Proof Writer

FOLIO

*A. Saparov and H. He.*

*Language models are greedy reasoners:*

*A systematic formal analysis of chain-of-thought. 2023*

GSM8K

FLD

Math comp.

Recentemente, molti ricercatori  
ProntoQA si sono posti questa domanda, Rule Taker  
e cercato modi di rispondere

LTLBench

Proof Writer

FOLIO

*T. Morishita, G. Morio, A. Yamaguchi, and others.  
Learning deductive reasoning from synthetic corpus  
based on formal logic. 2023*

# IA moderna e intelligenza generale ?

GSM8K

FLD

Math comp.

ProntoQA

E cosa hanno in comune  
tutti i risultati a queste  
ricerche?

Rule Taker

LTLBench

Proof Writer

FOLIO

*P. Shojaei, I. Mirzadeh, and others.*

*The Illusion of Thinking: Understanding the Strengths and Limitations of Reasoning Models via the Lens of Problem Complexity. 2025*

**Primo.** La capacità di arrivare a conclusioni corrette si assesta su una certa percentuale che poi tende a stabilizzarsi. Questa percentuale varia tra il 60% e l'80%.  
(prenderemmo un aereo che atterra correttamente l'80% delle volte?)

**Primo.** La capacità di arrivare a conclusioni corrette si assesta su una certa percentuale che poi tende a stabilizzarsi. Questa percentuale varia tra il 60% e l'80%.  
(prenderemmo un aereo che atterra correttamente l'80% delle volte?)

**Secondo.** Analizzando i ragionamenti ci si accorge che dopo una certa complessità la semantica distribuzionale fallisce generando allucinazioni, che sono errori **imprevedibili, inevitabili, e incorreggibili.**

**Primo.** La capacità di arrivare a conclusioni corrette si assesta su una certa percentuale che poi tende a stabilizzarsi. Questa percentuale varia tra il 60% e l'80%.  
(prenderemmo un aereo che atterra correttamente l'80% delle volte?)

**Secondo.** Analizzando i ragionamenti ci si accorge che dopo una certa complessità la semantica distribuzionale fallisce generando allucinazioni, che sono errori **imprevedibili, inevitabili, e incorreggibili.**

**Terzo.** Questo effetto non è destinato a cambiare e pure moltiplicando di molte volte le già enormi risorse destinate all'allenamento, si inizia a sospettare che i limiti siano stati già raggiunti.

Un sistema che **non ragiona** ma **imita il ragionamento**

è uno strumento che **non ha alcuna responsabilità**.

Non si può accusare nessuno di aver sbagliato, di aver inventato un dato,  
di aver risolto incorrettamente un esercizio, di aver costruito  
una politica inesistente.

Un sistema che **non ragiona** ma **imita il ragionamento**  
è uno strumento che **non ha alcuna responsabilità**.

Non si può accusare nessuno di aver sbagliato, di aver inventato un dato,  
di aver risolto incorrettamente un esercizio, di aver costruito  
una politica inesistente.

**Make America ChatGPT again: Experts  
say AI was used to create RFK Jr health  
report that cited false studies**

Un sistema che **non ragiona** ma **imita il ragionamento**

è uno strumento che **non ha alcuna responsabilità**.

Non si può accusare nessuno di aver sbagliato, di aver inventato un dato,  
di aver risolto incorrettamente un esercizio, di aver costruito  
una politica inesistente.

**L'IA di Musk è impazzita: Grok parla ovunque di  
"genocidio bianco"**

Un sistema che **non ragiona** ma **imita il ragionamento**

è uno strumento che **non ha alcuna responsabilità**.

Non si può accusare nessuno di aver sbagliato, di aver inventato un dato,  
di aver risolto incorrettamente un esercizio, di aver costruito  
una politica inesistente.

AI Gone Wild: Airline Has to Honor a  
Refund Policy Its Chatbot Fabricated

Cosa ci sta portando dove siamo? Questa IA (non tutta la IA è questa!) è caratterizzata, oltre che dalla sua natura tecnica, e per colpa della sua natura tecnica, anche dal tipo di training di hanno bisogno perchè la semantica ditribuzionale abbia senso. I costi, anche limitandoci ai soli costi vivi di hardware e di energia, sono sempre dell'ordine delle **decine** e delle **centinaia** di milioni di dollari statunitensi.

Cosa ci sta portando dove siamo? Questa IA (non tutta la IA è questa!) è caratterizzata, oltre che dalla sua natura tecnica, e per colpa della sua natura tecnica, anche dal tipo di training di hanno bisogno perchè la semantica ditribuzionale abbia senso. I costi, anche limitandoci ai soli costi vivi di hardware e di energia, sono sempre dell'ordine delle **decine** e delle **centinaia** di milioni di dollari statunitensi.

INNOVATION > BIG DATA

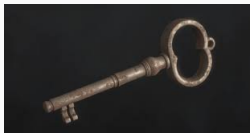
# The Extreme Cost Of Training AI Models

By [Katharina Buchholz](#), Contributor. ⓘ Katharina Buchholz is a Statista journ... ▼

[Follow Author](#)

Questo porta a che solo pochi, grandissimi players possono costruire, mettere a disposizione, e gestire questi modelli, che però adesso sono essenzialmente ovunque. Anche gli sviluppatori piccoli possono inserirli nei loro software per renderli migliori (?) sebbene la **chiave** rimanga sempre in mano ai creatori.

Questo porta a che solo pochi, grandissimi players possono costruire, mettere a disposizione, e gestire questi modelli, che però adesso sono essenzialmente ovunque. Anche gli sviluppatori piccoli possono inserirli nei loro software per renderli migliori (?) sebbene la **chiave** rimanga sempre in mano ai creatori.



E cosa accade se questi player staccano la spina?  
Che (come in parte già succede) il bisogno, una volta creato  
diventa a pagamento, allo scopo di rientrare nell'incredibile  
investimento. Questo crea anche delle situazioni limite!

E cosa accade se questi player staccano la spina?  
Che (come in parte già succede) il bisogno, una volta creato  
diventa a pagamento, allo scopo di rientrare nell'incredibile  
investimento. Questo crea anche delle situazioni limite!

A screenshot of a news headline on a dark background. The text is white and bold, reading: '\$1.5 Billion AI Unicorn Collapse, All Indian Programmers Impersonating!'.

**\$1.5 Billion AI Unicorn Collapse, All Indian  
Programmers Impersonating!**

## Conclusione: è tutto sbagliato?

Questa IA basata sulla semantica distribuzionale  
non è certamente inutile, se correttamente usata.  
L'idea è che dovrebbe essere impiegata unicamente per  
task **ripetitivi**, **analogici**, con un alto  
tasso di **errore ammesso**.

## Conclusione: è tutto sbagliato?

Questa IA basata sulla semantica distribuzionale non è certamente inutile, se correttamente usata. L'idea è che dovrebbe essere impiegata unicamente per task **ripetitivi, analogici**, con un alto tasso di **errore ammesso**.

Inoltre, dovrebbe essere impiegata unicamente per task le cui soluzioni **siamo in grado di controllare** perchè in assenza di un algoritmo responsabile e fatto di **passaggi logici** gli errori non solo sono parte del processo, ma non possono essere evitati.

Esempi di **task** che sono safe includono  
la **lettura** automatica e riassunto di documenti,  
la **scrittura** di documenti anche ufficiali che riportano dati  
e nozioni, e la realizzazione di attività e creazioni anche  
artistiche: se sono in grado di descrivere una situazione, e voglio  
rappresentarla graficamente, sono certamente autorizzato  
a lasciare che un sistema automatico lo faccia per me.

Esempi di **task** che sono safe includono  
la **lettura** automatica e riassunto di documenti,  
la **scrittura** di documenti anche ufficiali che riportano dati  
e nozioni, e la realizzazione di attività e creazioni anche  
artistiche: se sono in grado di descrivere una situazione, e voglio  
rappresentarla graficamente, sono certamente autorizzato  
a lasciare che un sistema automatico lo faccia per me.

Esempi di **task** che **non** sono safe includono  
la realizzazione di prodotti e soluzione di problemi  
che **altrimenti** non sarei in grado di risolvere,  
perchè non può essere attribuito alcun grado di responsabilità,  
diretta o indiretta, ad un sistema statistico che  
non può, per sua natura, esplicitare le decisioni.

