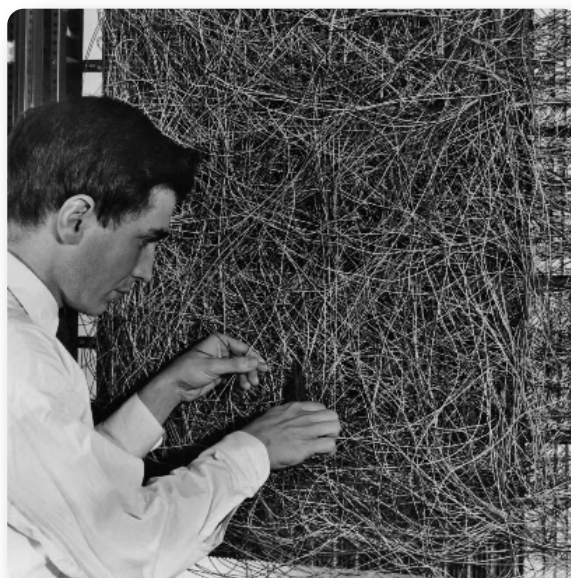


STORIA DELL'INTELLIGENZA ARTIFICIALE



Il primo esempio documentato di un agente autonomo e funzionale si trova nel mito di Talos, dalla mitologia greca, l'indistruttibile dio-automa di bronzo creato da Efesto.

Troviamo quindi sin dall'antichità il **desiderio di replicare la vita attraverso la meccanica**.

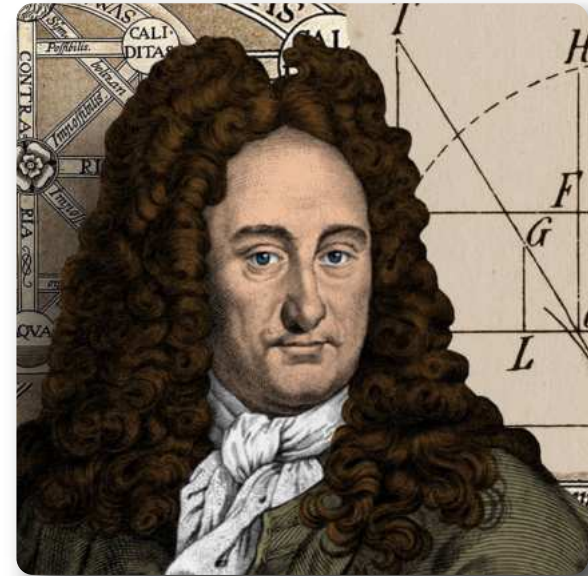


Moneta della grecia antica, con rappresentazione di Talos.

Il filosofo e matematico tedesco **Gottfried Wilhelm Leibniz** sognava una lingua universale del pensiero, una *characteristica universalis*, e un "calcolo del ragionamento", il *calculus ratiocinator*, che potesse ridurre le argomentazioni complesse a semplici calcoli.

Non a caso, fu l'inventore della prima calcolatrice meccanica in grado di eseguire le quattro operazioni e, cosa ancora più fondamentale, del sistema numerico binario nel 1703.

Leibniz non stava cercando di costruire un uomo di bronzo, stava tentando di **codificare la logica del pensiero**.



Gottfried Wilhelm Leibniz

Decimale	Binario	Decimale	Binario
0	0	8	1000
1	1	9	1001
2	10	10	1010
3	11	11	1011
4	100	12	1100
5	101	13	1101
6	110	14	1110
7	111	15	1111

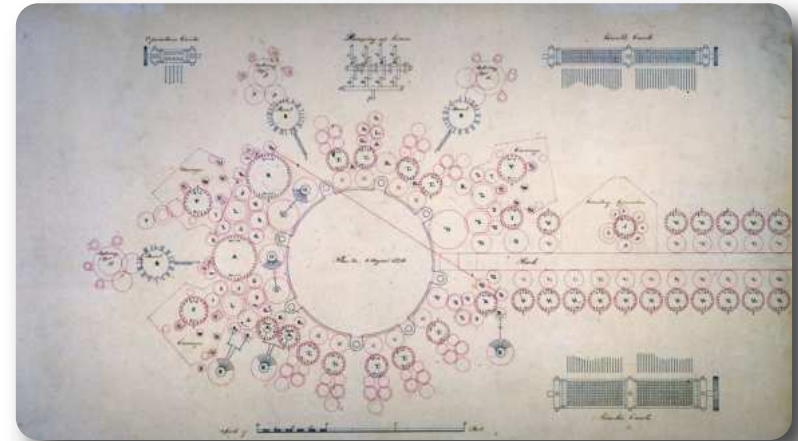
Sistema numerico binario.

Charles Babbage progettò la sua *Macchina Analitica* basandosi anche sul lavoro di Leibniz, un computer meccanico universale che, purtroppo, rimase solo sulla carta.

Fu **Ada Lovelace**, tuttavia, a coglierne il potenziale più profondo, anche grazie ad una lunga corrispondenza epistolare con l'italiano **Luigi Federico Menabrea**.

Capì che una tale macchina poteva manipolare non solo numeri, **ma qualsiasi simbolo che rappresentasse concetti astratti** come la musica o il linguaggio.

Il suo programma per il calcolo dei numeri di Bernoulli è considerato **il primo algoritmo destinato a essere eseguito da una macchina**, rendendola di fatto la prima programmatrice della storia.



Progetto della "Macchina Analitica" per il calcolo differenziale di Charles Babbage.



Ada Lovelace.



700 A.C. - Mito di Talos: Prima volta documentata che viene immaginata una macchina pensante.



1703 - Leibniz: Inventa l' "alfabeto" universale della computazione, il codice binario.

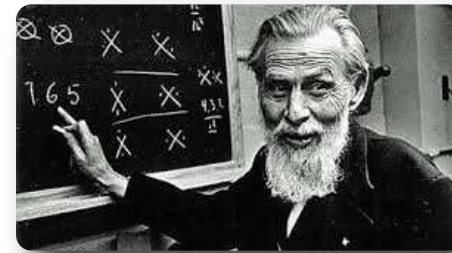


1843 - Lovelace: Intuisce che il codice binario può servire ben oltre i semplici calcoli matematici.

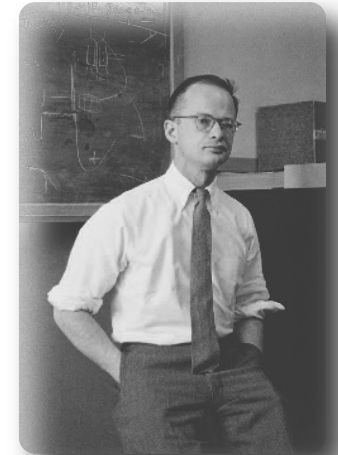
Uno dei pilastri fu il neurone di McCulloch-Pitts. Nel 1943, il neuroscienziato **Warren McCulloch** e il logico **Walter Pitts** proposero il **primo modello matematico di un neurone biologico**.

Il loro modello era elegantemente semplice: un dispositivo binario che riceveva più input e produceva un singolo output. Il neurone "si attivava", emettendo un 1 o uno 0. Dimostrarono che collegando questi semplici neuroni in reti, si potevano costruire circuiti **in grado di eseguire qualsiasi funzione logica fondamentale**.

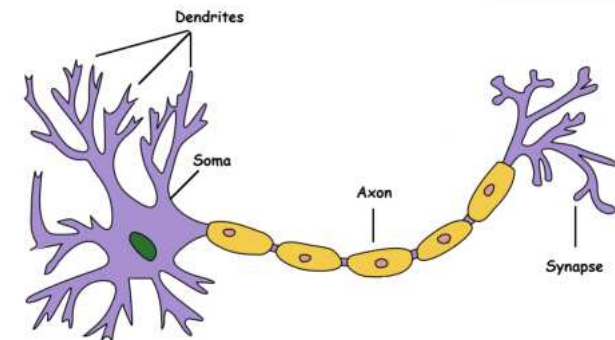
Questo fu un passo da gigante: suggeriva che i processi di pensiero complessi potessero emergere dall'interazione di molti elementi semplici e "stupidi".



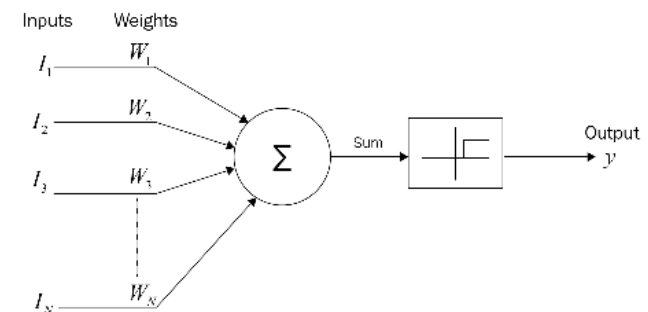
Warren McCulloch.



Walter Pitts.



Neurone biologico.

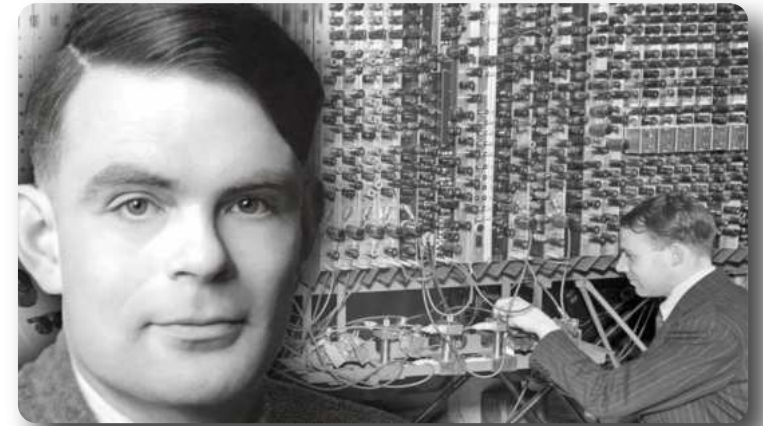


Neurone di Culloch-Pitts.

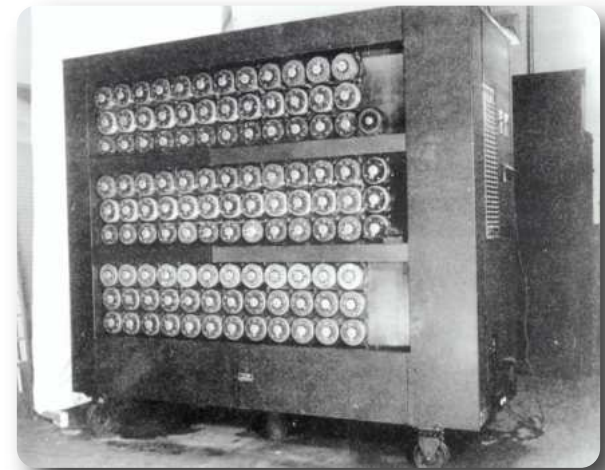
Il secondo pilastro fu fornito da **Alan Turing**. Nel suo fondamentale articolo del 1950, *"Computing Machinery and Intelligence"*.

Invece di chiedersi se una macchina ha una "coscienza", pose una domanda verificabile attraverso l'osservazione: **una macchina può imitare il comportamento umano** in modo così convincente da essere indistinguibile da una persona reale?

Nacque così il "Gioco dell'Imitazione", o **Test di Turing**, che spostò l'obiettivo dalla metafisica alla performance. Il suo genio fu di fornire all'intelligenza artificiale non solo uno scopo, ma anche un modo per misurare i propri successi.



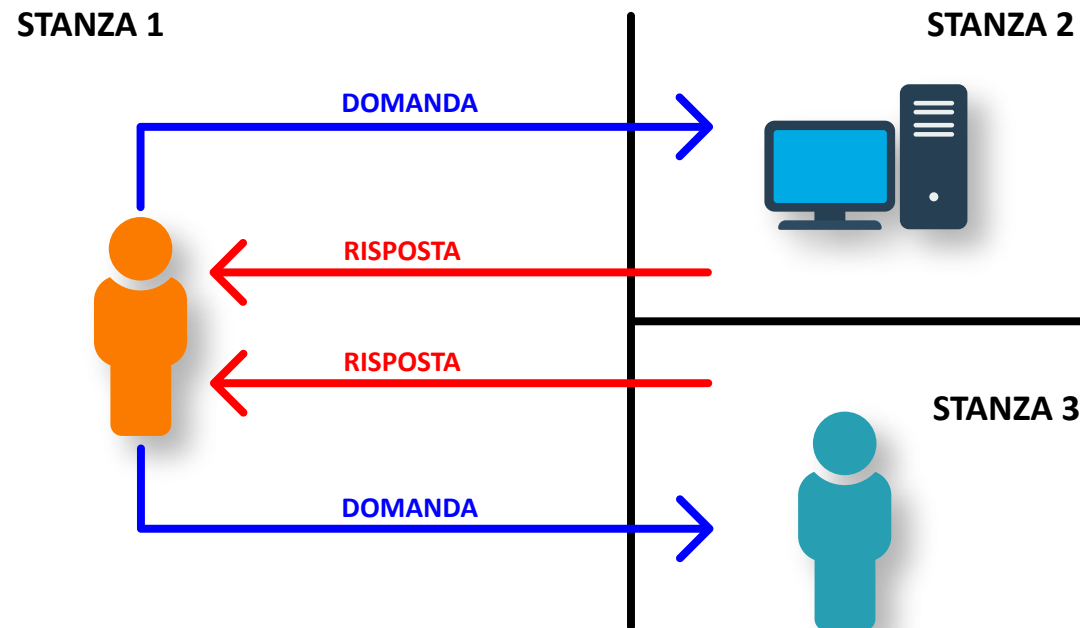
Alan Turing.



"The Bombe" il primo computer a scopo non generale.

Nel **test di Turing**, un interrogatore umano conversa via testo con due interlocutori nascosti: **un umano** e **una macchina**.

Se l'interrogatore non riesce a distinguere in modo affidabile chi dei due sia la macchina, allora quest'ultima ha superato il test, dimostrando un comportamento capace di simulare quello umano.



La Nascita di un Campo e il suo Primo Inverno (Anni '50 - 1980)

Fu nell'estate del 1956, che un'audace visione prese forma al Dartmouth College, questa visione è universalmente riconosciuta come l'**evento fondante del campo dell'intelligenza artificiale**.

Un giovane professore di matematica, **John McCarthy**, riunì un piccolo gruppo di scienziati per un progetto estivo di ricerca, battezzando l'incontro con un nome destinato a fare la storia: "**Intelligenza Artificiale**".

Guidati dall'idea che l'intelligenza potesse essere simulata, i ricercatori **Allen Newell** e **Herbert A. Simon** progettaronο il "**Logic Theorist**", un programma destinato a **dimostrare teoremi matematici**. Prima ancora di poterlo testare su un computer, ne provarono il funzionamento con un'ingegnosa simulazione manuale, che dimostrò la solidità del loro approccio.



Dartmouth College, Hanover, New Hampshire, USA.



*John
McCarthy.*



*Allen
Newell.*



*Herbert A.
Simon.*

Quando il Logic Theorist, considerato **il primo programma di intelligenza artificiale al mondo**, fu eseguito sul computer JOHNNIAC, i risultati furono straordinari.

Non solo dimostrò numerosi teoremi, ma **ne trovò uno con una soluzione persino più elegante dell'originale**, impressionando lo stesso Bertrand Russell (Russell era il co-autore dei *Principia Mathematica*, l'opera monumentale che il Logic Theorist era stato programmato a risolvere.)

Nonostante questo successo, l'accoglienza a Dartmouth fu tiepida. **Quel programma, però, segnò la nascita dell'IA simbolica (basata su regole)** e creò la prima grande spaccatura del settore: da un lato la logica programmata, dall'altro l'idea di macchine che imparano dai dati. Questa dicotomia fondamentale avrebbe definito la ricerca per i decenni a venire.

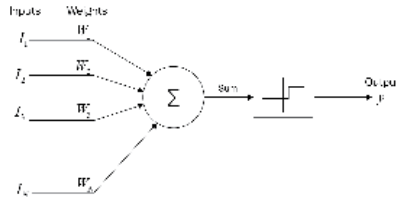


Bertrand Russell, coautore dei *Principia Mathematica*, Premio nobel per la Letteratura.



Principia Mathematica, di A. N. Whitehead e Bertrand Russell.

La Nascita di un Campo e il suo Primo Inverno (Anni '50 - 1980)



1943 - McCulloch-Pitts: Primo modello di un neurone.



1950 - Turing: Pone il traguardo che le macchine intelligenti dovranno raggiungere.

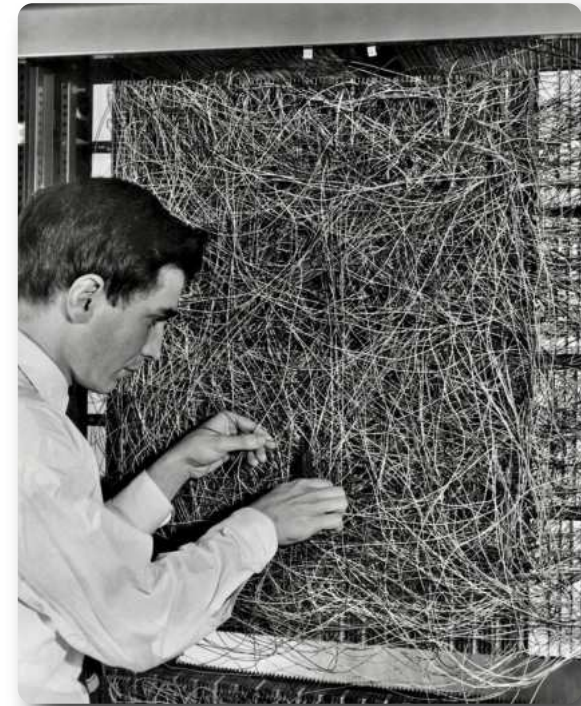


1956 - McCarthy, Newell, Simon: Viene coniato il termine “Intelligenza artificiale”, viene creato il primo programma IA, il “Logic Theorist”.

Negli anni '50, l'IA era dominata da un approccio basato su logica e regole. Ma nel 1958, lo psicologo **Frank Rosenblatt** presentò una rivoluzione: il **Perceptron**. Sostenuto dalla Marina USA, l'evento mediatico fu senza precedenti. Un computer IBM 704, davanti a un pubblico sbalordito, imparò in soli 50 tentativi a distinguere delle schede perforate.

L'entusiasmo fu esplosivo. Il New York Times lo definì "l'embrione di un computer elettronico che si aspetta sarà in grado di camminare, parlare, vedere, scrivere, riprodursi ed essere consapevole della sua esistenza". Rosenblatt stesso dichiarò che era la "prima macchina capace di avere un'idea originale". Non si trattava più di eseguire istruzioni, ma di apprendere dall'esperienza.

L'alba del machine learning era ufficialmente iniziata.



Frank Rosenblatt.

La Nascita di un Campo e il suo Primo Inverno (Anni '50 - 1980)

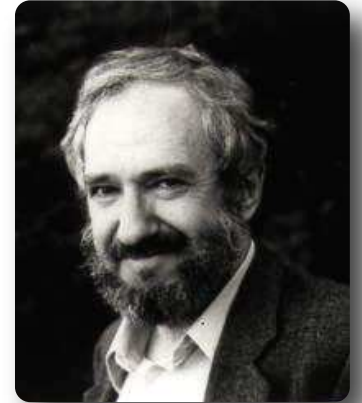
L'entusiasmo per il percettrone fu esplosivo, e vennero fatte grandi promesse al pubblico, le macchine avrebbero presto camminato parlato, visto e sarebbero state autoconsapevoli.

Tuttavia presto venne sollevata una **pesante critica sui limiti di questa tecnologia**, l'IA era capace di risolvere problemi “giocattolo” ma non quelli del mondo reale per diversi motivi.

Cominciò così il **primo inverno dell'IA**, i finanziamenti smisero di affluire, la ricerca si interruppe quasi completamente. Alla fine l'IA si era scontrata contro gli enormi limiti della tecnologia del tempo, congelando così il progresso nel campo per una decade.



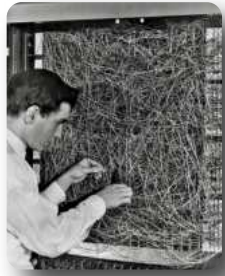
Marvin Minsky, MIT, criticò la tecnologia IA.



Seymour Papert, MIT, criticò la tecnologia IA.



Sir James Lighthill, redisse un rapporto pessimistico per il regno unito, influenzò gli investimenti globali nel campo IA.



1958 - Rosenblatt: Successo del percettrone e approccio connessionista nel campo IA.



1969-1973 - Minsky-Papert-Lighthill: Critica verso percettrone e approccio connessionista. Primo inverno dell'IA.

L'approccio connessionista era caduto in disgrazia, ma gli anni '80 videro esplodere l'approccio simbolico.

Il boom scaturì dai nuovi “**sistemi esperti**”. Se l'IA non poteva imparare da sola, allora bastava dargli la conoscenza umana ed insegnargli tutte le regole per ottenere i risultati desiderati.

Un esempio fu MYCIN del 1970. Una macchina a cui era stato insegnato a riconoscere le infezioni batteriche.

Questa ottenne dei risultati straordinari con un **accuratezza nelle diagnosi del 70%** superando i risultati degli studenti di medicina.

Tuttavia MYCIN non fu mai adottata dal mondo della medicina per le sue **limitazioni all'adattamento situazionale** e ai suoi frequenti errori grossolani.



Sistema esperto.

Mentre i sistemi esperti dominavano, diversi ingegneri lavoravano per risolvere il problema delle reti neurali basate sui perceptron.

Per superare i limiti delle prime reti neurali, venne sviluppata la “**backpropagation**” (retropropagazione dell'errore) inventata da **Geoffrey Hinton** e altri ricercatori, che risolse il problema di come addestrare reti complesse. Questo metodo calcola l'errore finale e lo propaga all'indietro per aggiustare i pesi di tutta la rete. La vera svolta avvenne quando questa tecnica fu abbinata a una nuova architettura creata da **Yann LeCun**: la Rete Neurale Convolutionale (CNN) nel 1988.

Questo ebbe un'enorme applicazione pratica, venne usata per leggere milioni di assegni bancari, dimostrando in modo definitivo la superiorità dell'AI rispetto alla programmazione classica.



Geoffrey Hinton



Yann LeCun.

La Rete Neurale Convoluzionale (CNN) di LeCun era eccellente per il riconoscimento di immagini, tuttavia era drammaticamente carente in termini di memoria, una pessima notizia nel caso ci fosse stato bisogno di creare un AI in grado di gestire dati sequenziali come il testo.

Nel 1997, i ricercatori **Sepp Hochreiter** e **Jürgen Schmidhuber** hanno introdotto le **reti LSTM** (Long Short-Term Memory) per superare un limite fondamentale delle precedenti intelligenze artificiali: l'incapacità di ricordare informazioni su lunghi periodi.

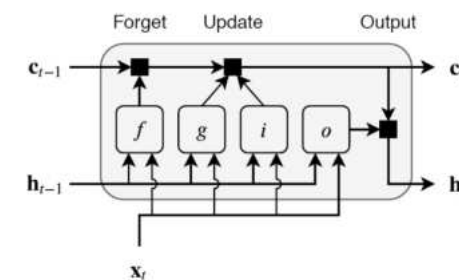
Le LSTM imparano dinamicamente a memorizzare i dati rilevanti, a scartare quelli inutili e a decidere quando utilizzare le informazioni conservate, permettendo così di analizzare dati complessi come testi lunghi o serie storiche finanziarie.



Sepp Hochreiter.



Jürgen Schmidhuber.



Schema di un LSTM.

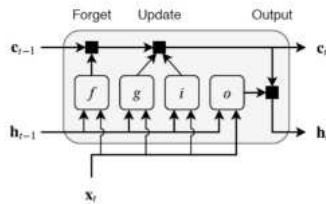
X_t : Input
 h_t : Hidden state
 C_t : Cell state
 f : Forget gate
 g : Memory cell
 i : Input gate
 o : Output gate



1970 - Stanford University: Esplosione dei “Sistemi esperti”.



1988 - LeCun: Invenzione della “Backpropagation” e delle reti neurali Convoluzionali.



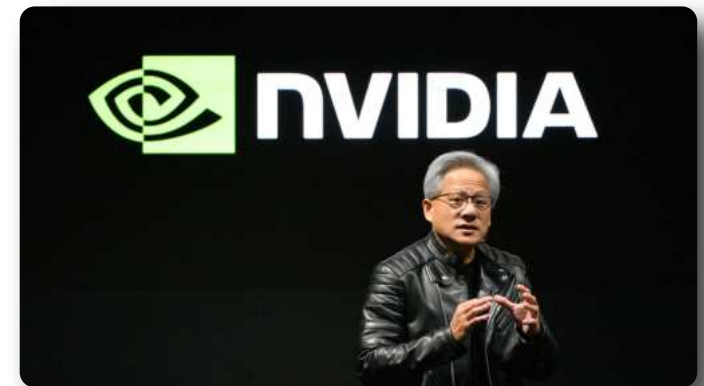
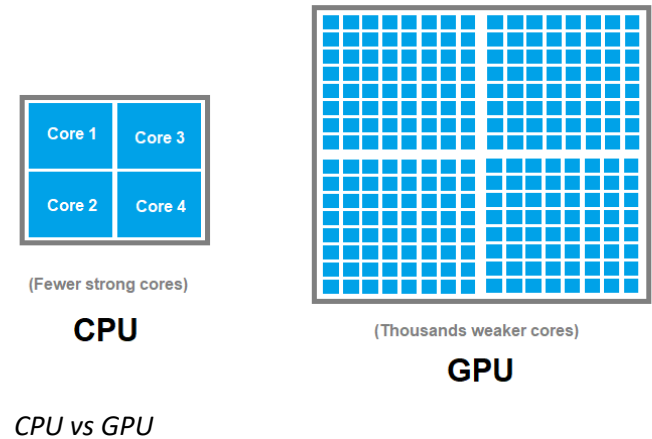
1997 - Hochreiter, Schmidhuber: Invenzione delle reti “LSTM” per dati sequenziali.

L'Esplosione del Deep Learning (2000 - 2025)

Per molto tempo, l'intelligenza artificiale è stata come un'auto da corsa bloccata nel traffico. Il "motore" di allora, il processore del computer (CPU), era bravo a fare una cosa complessa alla volta, ma addestrare un'IA richiede di fare migliaia di calcoli semplici contemporaneamente.

La svolta arrivò al mondo dei videogiochi. **Le schede grafiche (GPU)**, nate per creare mondi virtuali complessi, erano progettate per eseguire un'enorme quantità di operazioni in parallelo.

Nel 2006, l'azienda **NVIDIA** con a capo **Jen-Hsun "Jensen" Huang**, rilasciò una piattaforma software chiamata **CUDA**, dove gli sviluppatori potevano usare le GPU per qualsiasi tipo di calcolo. Nel 2014, NVIDIA rese l'accesso a questa super-potenza di calcolo facile per tutti i ricercatori. Processi di addestramento che prima richiedevano mesi potevano essere completati in poche ore.



Jen-Hsun "Jensen" Huang. CEO di NVIDIA.

L'Esplosione del Deep Learning (2000 - 2025)

Avere un motore potente non serve a nulla senza carburante. L'AI, specialmente quella che impara da sola (il deep learning), è affamata di dati. Più ne "vede", più diventa intelligente. La seconda svolta avvenne grazie a una ricercatrice: **Fei-Fei Li**.

Nel **2006**, capì che la comunità scientifica era troppo concentrata sugli algoritmi e stava trascurando i dati. Iniziò così un progetto colossale chiamato **ImageNet**: creare un gigantesco archivio di milioni di immagini, tutte etichettate a mano per descrivere cosa contenessero. Mobilitò quasi 50.000 persone nel mondo tramite **Amazon Mechanical Turk**.

Nel 2010 lanciò una competizione (**ILSVRC**) per vedere quale algoritmo riuscisse a riconoscere le immagini di ImageNet con maggiore precisione. Questa sfida divenne il "campionato mondiale" della visione artificiale, questo spinse i ricercatori a sviluppare sistemi sempre più potenti e precisi.



ImageNet.



Fei-Fei Li.

Nel 2019, il pioniere dell'IA **Rich Sutton** ha formalizzato un principio chiave dell'era moderna nel suo saggio "**La Lezione Amara**".

La lezione sostiene che i tentativi di programmare manualmente la conoscenza umana in regole fisse sono destinati a essere superati.

L'approccio vincente, reso possibile da GPU potenti e dati massicci, consiste nel dare alla macchina algoritmi di apprendimento generali e lasciare che impari da sola.

Ad esempio, invece di descrivere le caratteristiche di un gatto, è più efficace mostrargliene milioni di immagini.

Questo metodo, che migliora con l'aumentare della potenza di calcolo e dei dati, è il fondamento del successo del deep learning.



Richard S. Sutton, e il suo saggio "The bitter lesson".

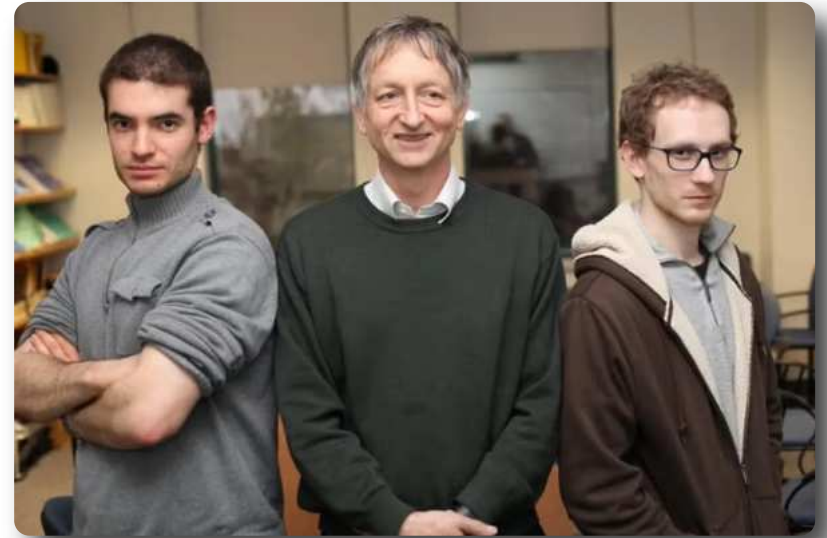
L'Esplosione del Deep Learning (2000 - 2025)

Nel 2012, un modello chiamato **AlexNet** creato da **Alex Krizhevsky, Ilya Sutskever** e **Geoffrey Hinton**, demolì la concorrenza nell'ILSVRC.

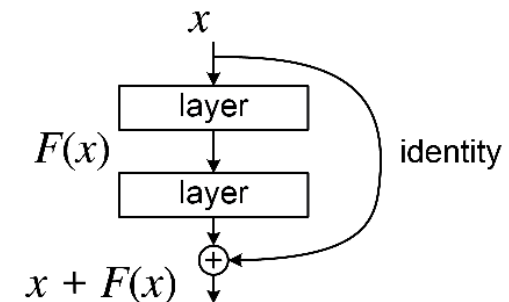
La sua architettura introduceva dei concetti che divennero pilastri nel deep learning: Una maggiore profondità della rete, una nuova funzione (ReLU), e una tecnica di spegnimento casuale dei neuroni (Dropout).

Questo approccio innescò un'esplosione di innovazioni nel campo AI, i ricercatori dovettero inventare numerosi “trucchi” strutturali per migliorare le prestazioni, una versione sfumata della “The Bitter Lesson”.

ResNet-152 nel 2015 ottenne un tasso di errore del 3.57% **superando per la prima volta le prestazioni umane**. ResNet è il predecessore del Vision Transformer (ViT), il sistema usato oggi.



Da sinistra a destra: Ilya Sutskever, Geoffrey Hinton e Alex Krizhevsky.



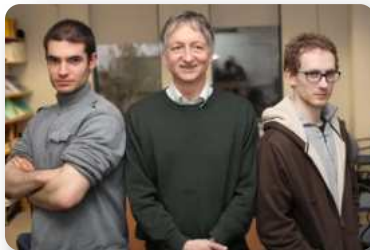
Schema rete neurale residuale (ResNet)



2006 - NVIDIA: Lancio di CUDA ed esplosione di progetti AI indipendenti grazie alle GPU.



2006 - Fei-Fei Li: Lancio di ImageNet, pone un obiettivo da raggiungere per la Machine Vision e rende accessibile il “labeling” per apprendimento supervisionato.



2012 - Krizhevsky, Sutskever e Hinton: Con AlexNet aprono la strada alla corsa alla profondità delle reti neurali.

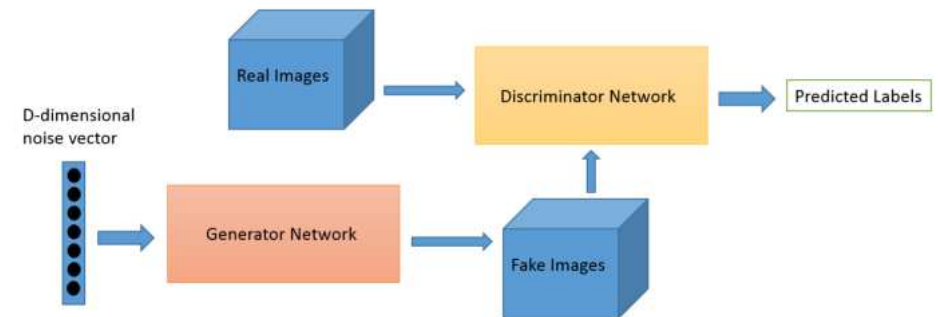
Le **Reti Generative Avversarie (GAN)**, ideate da **Ian Goodfellow** nel 2014, sono una forma di intelligenza artificiale che impara a creare dati nuovi e originali.

Questa tecnologia si basa sulla competizione tra due reti neurali. Una rete, il Generatore, produce dati artificiali, come l'immagine di un volto. La seconda rete, il Discriminatore, viene addestrata per distinguere i dati reali da quelli falsi creati dal Generatore.

Attraverso questa sfida continua, il Generatore migliora costantemente nel tentativo di ingannare il Discriminatore, fino a quando le sue creazioni diventano così realistiche da essere quasi indistinguibili da quelle autentiche.



Ian Goodfellow.



Schema di una GAN.

L'Esplosione del Deep Learning (2000 - 2025)

Nel 2017, un team di ricercatori di Google ha introdotto un'architettura rivoluzionaria chiamata **Transformer**. Prima di allora, i modelli di intelligenza artificiale leggevano il testo una parola alla volta, un processo lento e inefficiente. Il Transformer ha cambiato tutto introducendo il meccanismo di "self-attention" (auto-attenzione).

Invece di procedere in sequenza, questo sistema **analizza tutte le parole di una frase simultaneamente, pesando l'importanza di ciascuna parola rispetto alle altre per capirne il contesto**. Ad esempio, capisce a chi si riferisce un pronome esaminando l'intera frase in un colpo solo.

Questa capacità di elaborazione parallela ha reso l'addestramento molto più veloce ed efficiente, aprendo la strada alla creazione dei moderni e potentissimi modelli linguistici come GPT.



Il team dietro l'ideazione dei modelli "Transformer".



ChatGPT, primo modello Transformer commercializzato.



2014 - Goodfellow: Invenzione delle GAN, le AI ora possono competere tra di loro per perseguire risultati sempre più accurati.



2017 - Google: Invenzione dei modelli Transformer potenziandone di molto la capacità di comprensione ed ottimizzando i costi.

L'Esplosione del Deep Learning (2000 - 2025)

Nel marzo 2016, l'intelligenza artificiale **AlphaGo di DeepMind** ha sconfitto il campione del mondo Lee Sedol nel complesso gioco del Go, un evento che molti credevano a decenni di distanza. Il Go era considerato l'ultima frontiera dell'intuizione umana, inattaccabile dai computer per il suo numero quasi infinito di mosse possibili.

Sebbene AlphaGo abbia vinto la serie per 4-1, il match è ricordato per due momenti iconici. Il primo è la "Mossa 37" di AlphaGo, un colpo così inaspettato e creativo da essere inizialmente scambiato per un errore, ma che si rivelò geniale. Il secondo è il "Tocco Divino" di Lee Sedol, una mossa umana così brillante da mandare in crisi l'algoritmo e assicurargli l'unica vittoria. L'evento ha dimostrato che l'IA poteva raggiungere una creatività sovrumana, ma ha anche celebrato il picco dell'ingegno umano.



Go (Gioco).

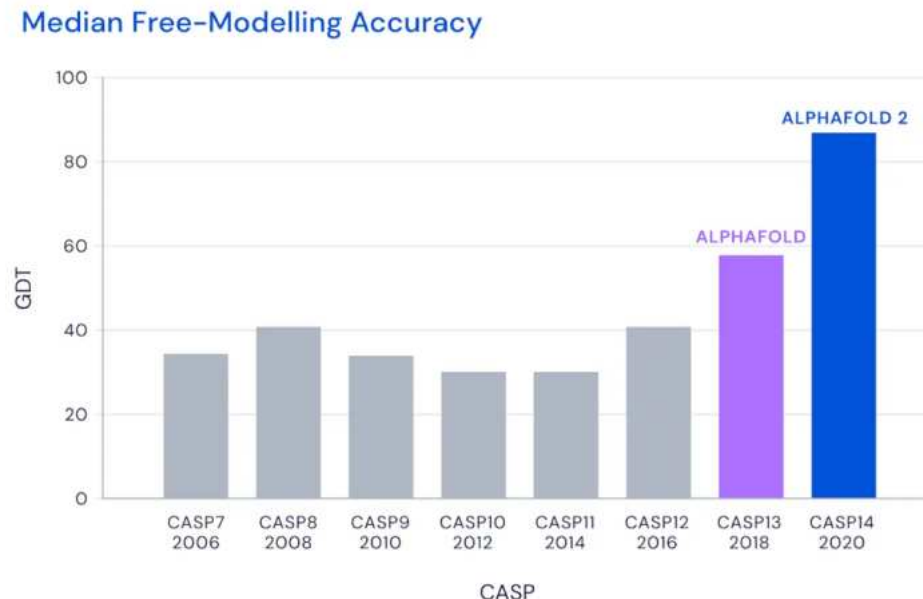


AlphaGo ed esecutore umano VS Lee Sedol, campione del mondo.

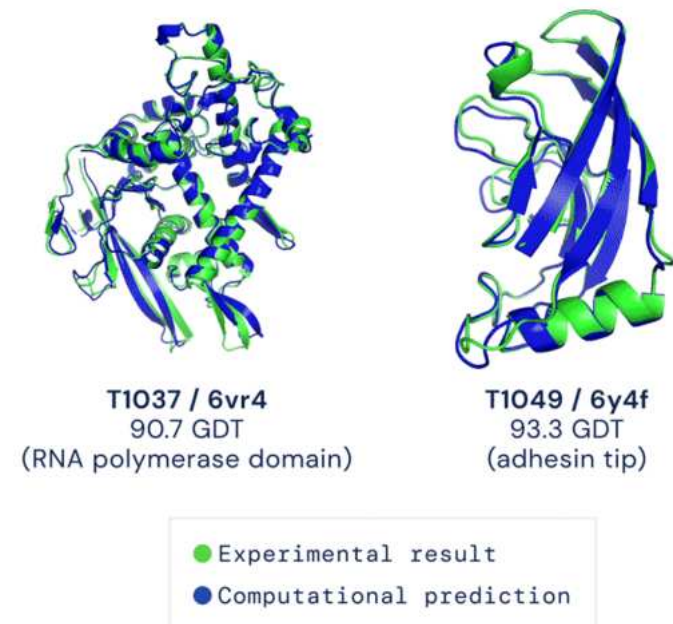
L'Esplosione del Deep Learning (2000 - 2025)

Per 50 anni, una delle sfide più grandi della biologia è stata il "problema del ripiegamento proteico": predire la forma 3D di una proteina partendo dalla sua sequenza, un'informazione cruciale per creare farmaci e capire le malattie. Questo compito era considerato quasi impossibile a causa della sua enorme complessità.

Nel 2020, **AlphaFold 2**, un'intelligenza artificiale creata da DeepMind, ha risolto questa sfida. Il sistema è in grado di predire la struttura delle proteine con una precisione paragonabile a quella dei metodi sperimentali, che però richiedono mesi o anni. **DeepMind ha reso pubbliche le strutture di oltre 200 milioni di proteine, accelerando in modo drastico la ricerca su malattie** come l'Alzheimer e il cancro e aprendo una nuova era nella scoperta di farmaci



Risultati ad alta precisione sul ripiegamento proteico, AlphaFold 2.



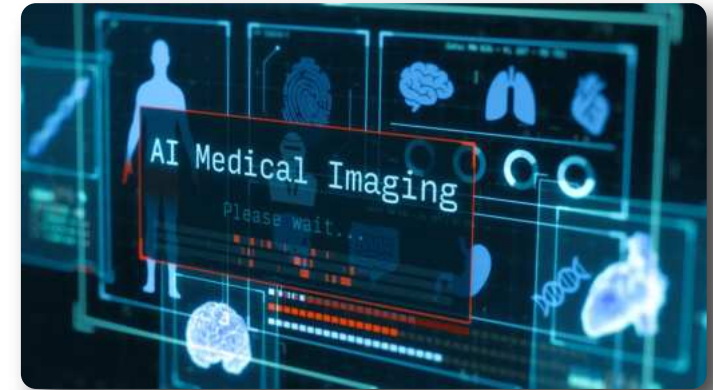
L'Esplosione del Deep Learning (2000 - 2025)

L'intelligenza artificiale non è più solo una tecnologia per imprese eccezionali, ma è già integrata in molti aspetti della nostra vita quotidiana.

In **medicina**, aiuta i dottori a individuare tumori da immagini come radiografie e TAC con una precisione sempre maggiore.

In **finanza**, sistemi automatici la usano per prevedere i movimenti di mercato e per scovare transazioni fraudolente in tempo reale.

Nell'**intrattenimento**, l'IA è il motore dietro i sistemi di raccomandazione, come quello di Netflix, che personalizza i suggerimenti di film e serie TV per ogni utente. Questi esempi dimostrano che l'IA sta diventando uno strumento fondamentale, che ci permette di comprendere sistemi (biologici, economici, sociali) troppo complessi per la mente umana.



L'Esplosione del Deep Learning (2000 - 2025)

700 A.C. - Mito di Talos: Prima volta documentata che viene immaginata una macchina pensante.

1703 - Leibniz: Inventa l' "alfabeto" universale della computazione, il codice binario.

1843 - Lovelace: Intuisce che il codice binario può servire ben oltre i semplici calcoli matematici.

1943 - McCulloch-Pitts: Primo modello di un neurone.

1950 - Turing: Pone il traguardo che le macchine intelligenti dovranno raggiungere.

1956 - McCarthy, Newell, Simon: Viene coniato il termine "Intelligenza artificiale", viene creato il primo programma IA, il "Logic Theorist".

1958 - Rosenblatt: Successo del perceptrone e approccio connessionista nel campo IA.

1969-1973 - Minsky-Papert-Lighthill: Critica verso perceptrone e approccio connessionista. Primo inverno dell'IA.

1970 - Stanford University: Esplosione dei "Sistemi esperti".

1988 - LeCun: Invenzione della "Backpropagation" e delle reti neurali Convoluzionali.

1997 - Hochreiter, Schmidhuber: Invenzione delle reti "LSTM" per dati sequenziali.

2006 - NVIDIA: Lancio di CUDA ed esplosione di progetti AI indipendenti grazie alle GPU.

2006 - Fei-Fei Li: Lancio di ImageNet, pone un obiettivo da raggiungere per la Machine Vision e rende accessibile il "labeling" per apprendimento supervisionato.

2012 - Krizhevsky, Sutskever e Hinton: Con AlexNet aprono la strada alla corsa alla profondità delle reti neurali.

2014 - Goodfellow: Invenzione delle GAN, le AI ora possono competere tra di loro per perseguire risultati sempre più accurati.

2017 - Google: Invenzione dei modelli Transformer potenziandone di molto la capacità di comprensione ed ottimizzando i costi.