



Ciclul „INTELIGENȚA ARTIFICIALĂ ÎN MEDICINĂ”

Oct. – Dec. 2025

Coordonator ciclu: Prof. dr. Daniel LIGHEZAN

Coordonator științific: Prof. dr. Gheorghe Ioan MIHALAȘ



Tema 4

Retele Neuronale

Prof. Dr. Gheorghe I MIHALAS

Universitatea de Medicină și Farmacie Victor Babeș Timișoara

Centrul de Modelare a Sistemelor Biologice și Analiza Datelor

Laboratorul de Inteligență Artificială

și

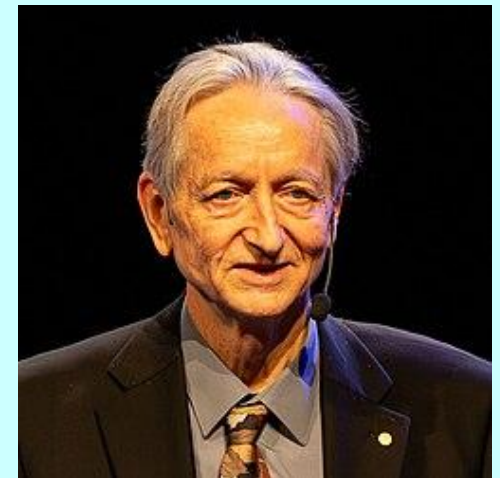
Academia de Științe Medicale

Comisia de Informatică Medicală, Inteligență Artificială și Protecția Datelor

„In artificial intelligence we first learned how to reason; only later did we start to learn how to learn”

„We are trying to build machines that can learn like the brain — not just compute like a calculator”

*Geoffrey Hinton (n. 1947)
Nobel Prize 2024
Univ. of Toronto*



Planul tematic al ciclului de A.I. - 2025



Data	Tema prelegerii
7 oct	Introducere în AI. Definiții, istoric. Abordări cognitive și conexiuniste. Sisteme discriminative și generative. Agenți AI. Aplicații medicale – privire generală, clasificări. Integrarea AI în activitatea medicală.
14 oct	„Machine learning”, învățare supervizată și nesupervizată. Analiza datelor medicale: predicții, clasificări. Date sintetice, simulări. Modelarea proceselor, proiectarea studiilor clinice, evaluare.
28 oct	AI pentru decizia medicală. Reprezentarea cunoștințelor, ontologii medicale, codificări. Sisteme logice, sisteme baysiene. Logica fuzzy. Sisteme expert medicale, decizie asistată. Monitorizarea pacienților.
4 nov	Rețele neuronale: principii, clasificări. Perceptronul, „deep learning”. Aplicații în imagistica medicală, bioinformatică, medicina de precizie. Performanță vs. transparență, criterii de validare.
11 nov	Sisteme generative conversaționale. Prelucrarea limbajului natural, „LLM – large language models”. Construcția și funcționarea unui Chat-bot. Aplicații medicale, chatboți specializați. Prompt engineering. Limite și capcane.
18 nov	AI în învățământul medical și în cercetare. Generare întrebări, cazuri clinice interactive. Utilizarea LLM în redactarea științifică. Construcția și antrenarea de chatboți specializați. Etica utilizării AI în educație și cercetare.
25 nov	Limitele actuale, aspecte etice și juridice. Limite: confidențialitate, bias. Transparență, sisteme explicabile XAI. Tehnologia empatică, riscuri și abuzuri. Competențe AI, responsabilități. Reglementări.

Planul prelegerii

- Drumul către AI – Caracteristicile „inteligenței”
- Modele neuronale
- Perceptronul – rețea cu un singur strat
 - Caracteristicile perceptronului
 - Backpropagation, Noțiunea de „învățare”
- Arhitectura rețelelor neuronale
 - Straturi ascunse
- Tipuri și clasificări
- Deep learning
- Aplicații medicale majore
 - Imagistica medicală
 - „Natural Language Processing”
- Performanțe vs. Transparență
- Criterii de validare

Drumul către AI – Caracteristicile „intelenței”



- **Elementele fundamentale ale intelenței:**
 - Rezolvare de probleme, Recunoaștere de tipare
 - Generalizare, Abstractizare
 - Învățare, Memorare
- **Limitele sistemelor computaționale tradiționale**
 - Sistemele cognitive (bazate pe reguli) – performante la execuția algoritmilor
 - Limite majore:
 - Nu învață singure
 - Nu pot generaliza
- **Provocarea – crearea unor sisteme care pot învăța**
 - Sursa de inspirație - creierul

Modele Neuronale



- **Sursa de Inspirație: Creierul Uman**

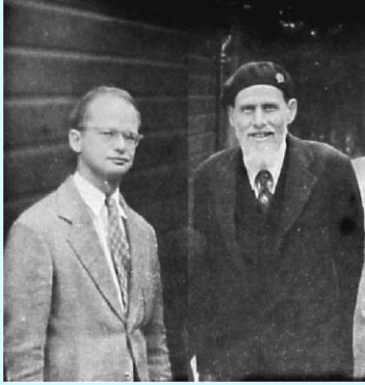
- Cel mai puternic sistem de calcul și învățare cunoscut.
- Este compus din miliarde de **neuroni** (celule nervoase) interconectați.

- **Neuronul Biologic (Simplificat):**

- Primește semnale de la alți neuroni prin **dendrite**.
- Procesează semnalele în **soma** (corpul celulei).
- Dacă stimulul depășește un anumit prag, neuronul "se activează" și trimite un semnal prin **axonul** său către alte neurone.
- Legătura dintre neuroni (sinapsa) este cheia. Ele se pot întări sau slăbi în timp, constituind baza *memoriei* și *învățării*.

- **Afirmație cheie:** *"Dacă natura a rezolvat problema inteligenței prin interconexiune, de ce nu am încerca și noi aceeași cale?"*

De la Biologie la Model Matematic

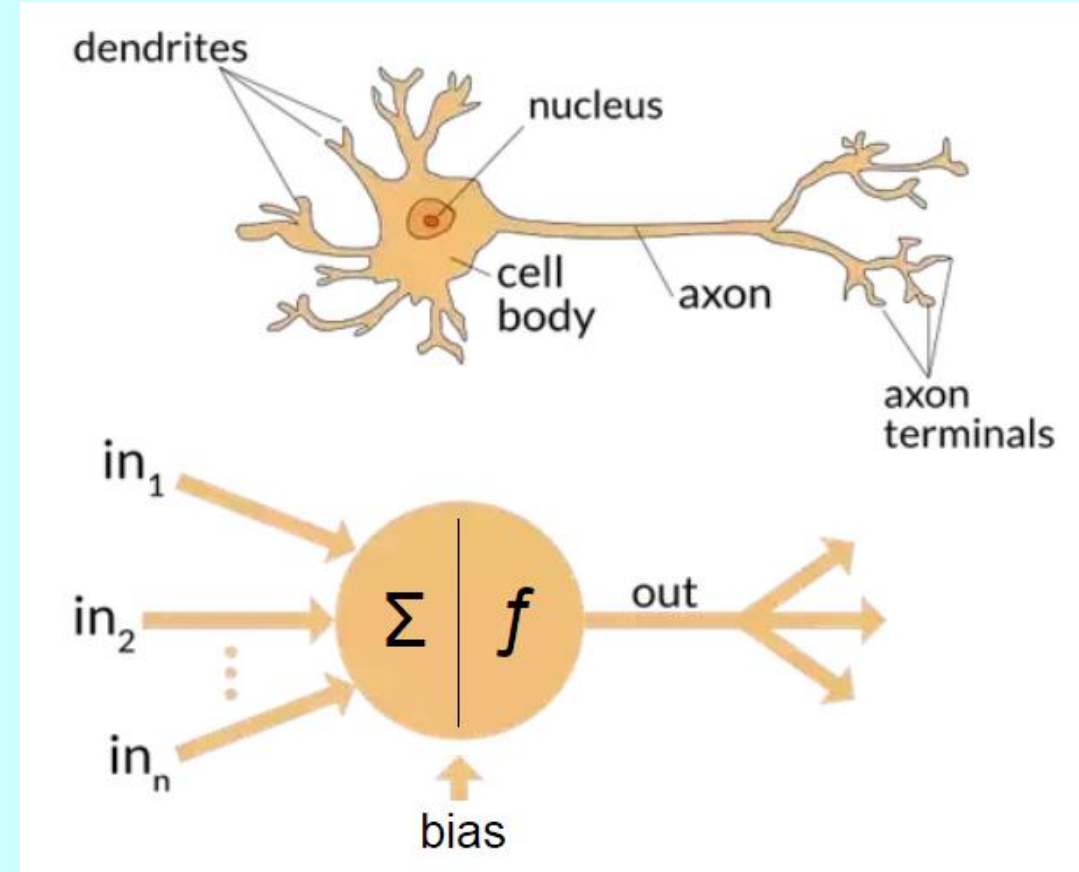


1943: McCulloch și Pitts
primul model de neuron

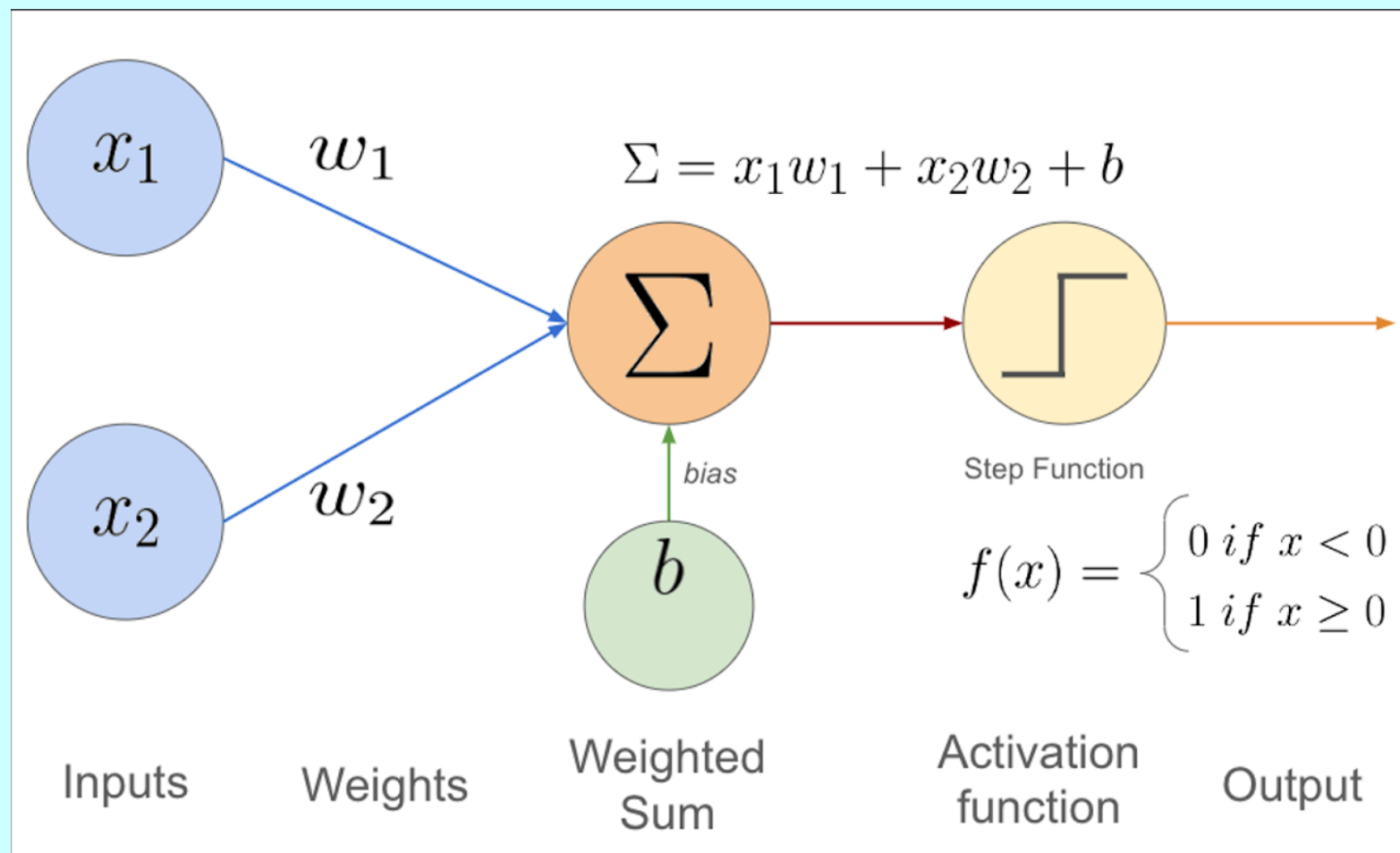
- **Idea Revoluționară:**

- Un neuron = unitate logică:
 - primește intrări (0/1)
 - le însumează și produce o ieșire (0/1) dacă un prag este depășit.

- **Semnificație:** o rețea de "neuroni" poate calcula orice funcție logică



Perceptron, 1957, F Rosenblatt



Neural Network: Perceptron

Next topic

Video

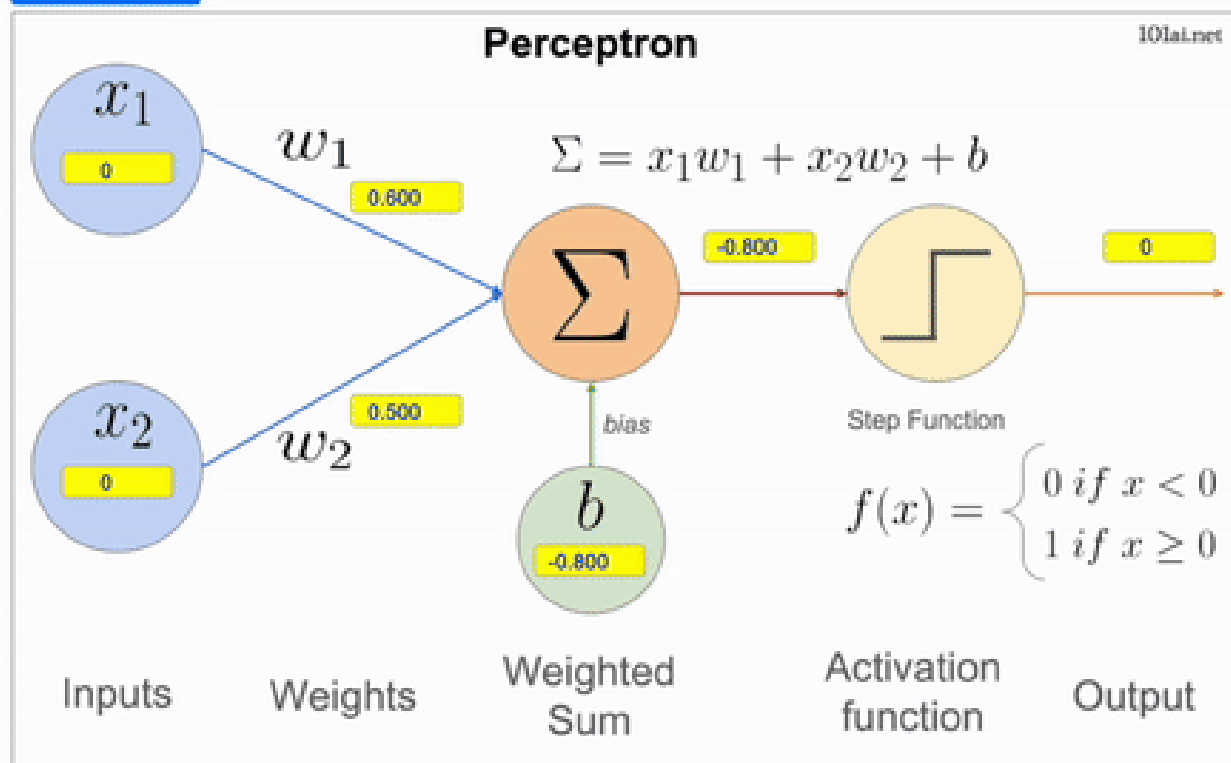
Notes

Code

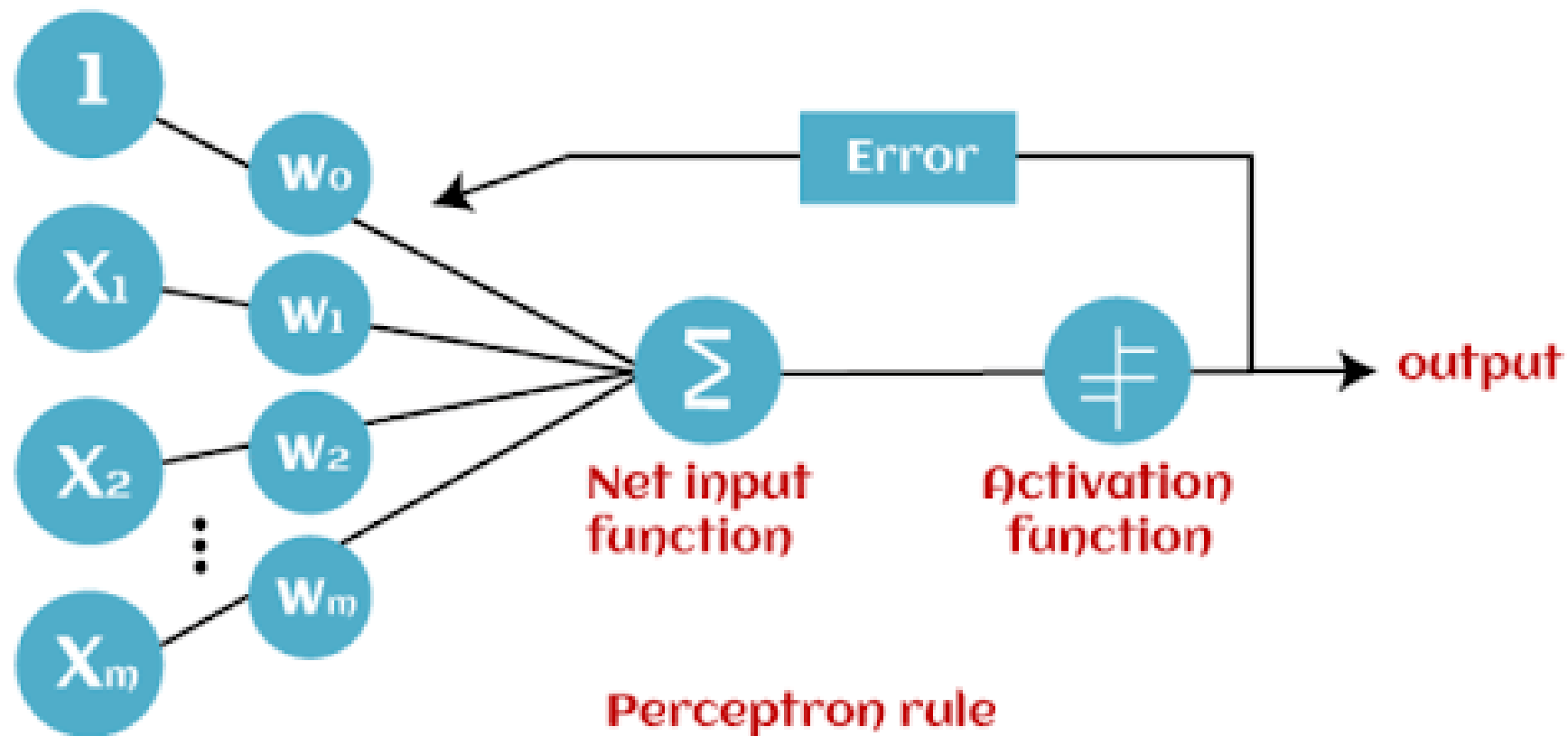
Example

AND gate with Step activation

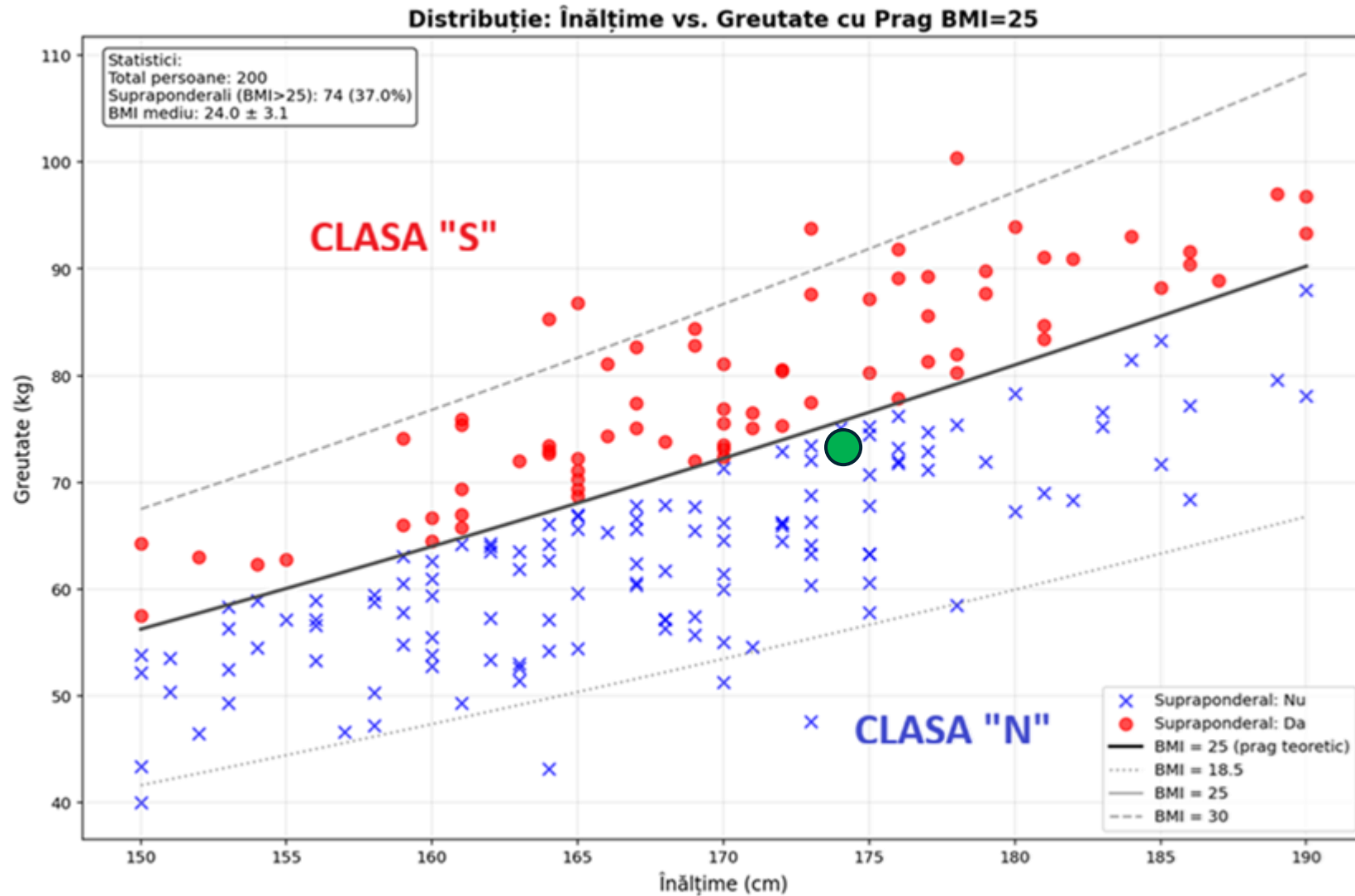
Next data input



Învățare = ajustare ponderi, 1962, F.Rosenblatt



Exemplu



Ce clasă este
un nou caz ?
 $h = 174 \text{ cm}$
 $G = 72 \text{ kg}$

Este în
CLASA N

Verificare:
 $\text{BMI} = 23.79$
 $\text{BMI} < 25$

Problema XOR: Minsky, 1969 → „*AI Winter*” (~ 15 ani)

FUNCȚII DE DECIZIE

Funcția Prag

Pentru operatorii logici

- Prag = 1.5 → AND
- Prag = 0.5 → OR

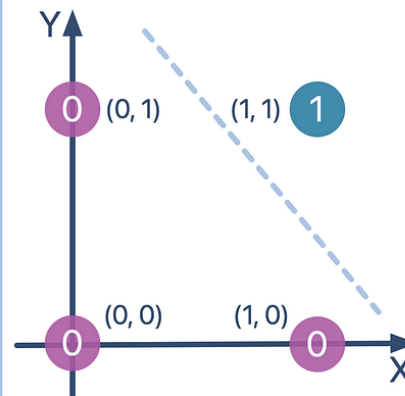
Funcția logistică
→ Probabilități

Pentru XOR –
Perceptronul
NU POATE
ÎNVAȚA

Separabilitate liniară

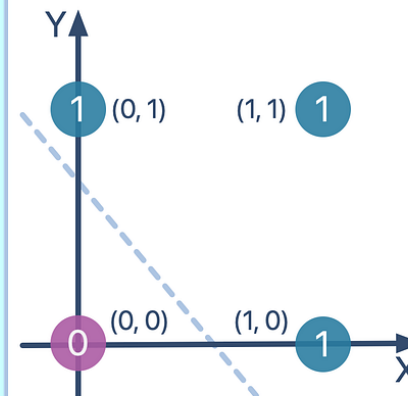
Bitwise AND

X	Y	X and Y
0	0	0
0	1	0
1	0	0
1	1	1



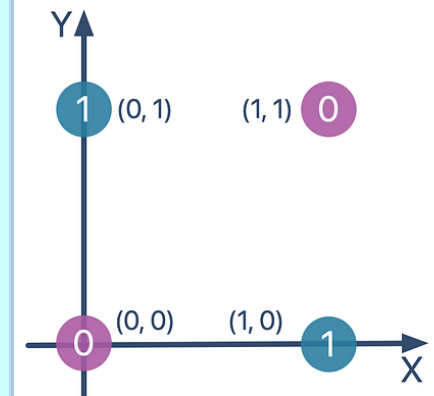
Bitwise OR

X	Y	X or Y
0	0	0
0	1	1
1	0	1
1	1	1



Bitwise XOR

X	Y	X xor Y
0	0	0
0	1	1
1	0	1
1	1	0

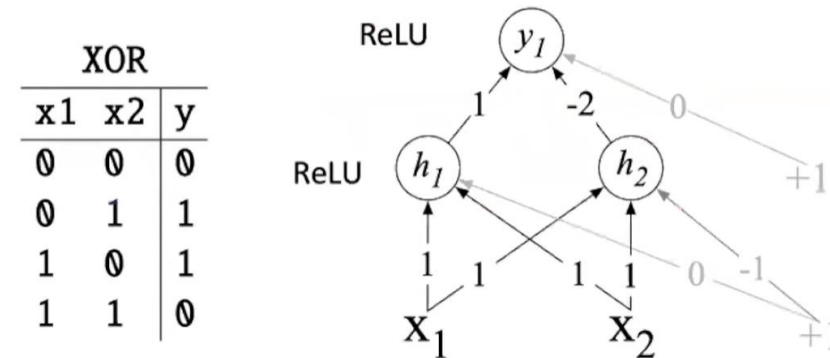


Multilayer perceptron → Rețele neuronale

- **Rețele Neuronale = Neural Networks (NN)**
- *Hinton, Rumelhart & Williams (1986):*
 - Rețea multistrat („hidden layers”)
 - Backpropagation
 - Deep learning
- **Proprietățile NN:**
 - Procesare paralelă și distribuită
 - Învățare prin Adaptare
 - Robustețe

Solution to the XOR problem

XOR **can't** be calculated by a single perceptron
XOR **can** be calculated by a layered network of units.



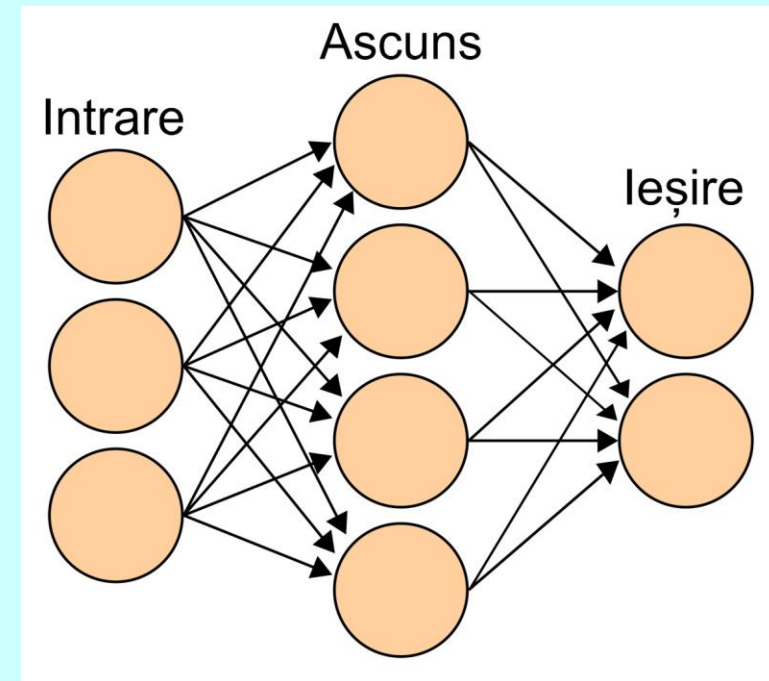
Stanford

Arhitectura Rețelelor Neuronale

Ideea centrală: O rețea neuronală este un graf orientat, format din straturi de neuroni interconectați

Componente:

- **Stratul de intrare** (Input layer)
 - Nr. neuroni depinde de date
- **Straturi ascunse** (Hidden layer)
 - „Inima rețelei”
 - Conexiuni multiple
- **Stratul de ieșire** (Output layer)
 - Nr. neuroni depinde de sarcină
 - Predicție sau Clasificare binară – 1 neuron
 - Clasificare N clase – N neuroni



- **Fluxul informațional:** *forward propagation*

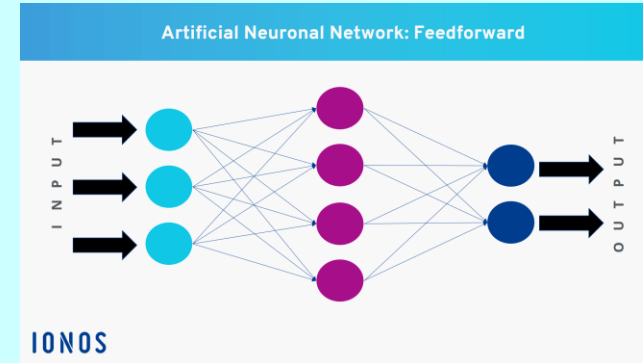
- **Ajustarea ponderilor:** *backpropagation*

Clasificări ale NN

După direcția fluxului informațional

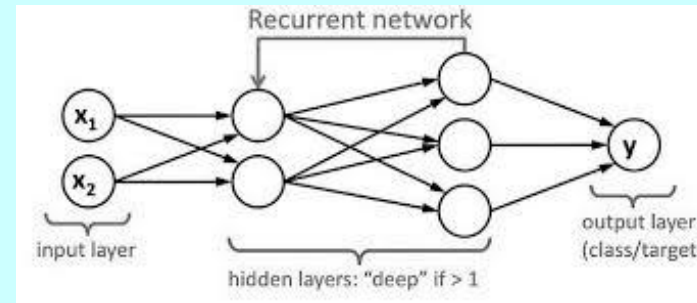
- **Rețele Feed-Forward (FFN)**

- Flux unidirecțional
- Clasificări date, Predicții, Clasificări imagini



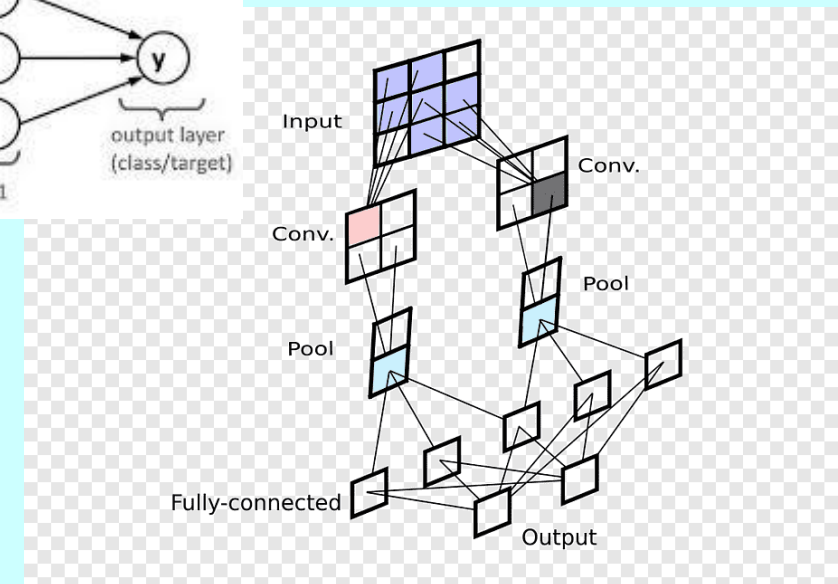
- **Rețele Recurente (RNN)**

- Au bucle spre straturi anterioare
- Analize secvențiale (ECG)



- **Rețele Convoluționale (CNN)**

- Au straturi specializate
- Reduc parametrii prin filtre
- Analiza imaginilor, recunoașteri



După arhitectura conexiunilor

- **Rețele Complet conectate**
 - Clasificări
- **Rețele convoluționale**
 - Conexiuni locale (filtre) – Analize de imagini
- **Autoencodere**
 - Comprimă datele – Analize genetice
- **Generative Adversarial Networks (GAN)**
 - Două rețele – una generează, cealaltă discriminează
 - Generare de imagini
- **Transformere**
 - Bazate pe „Atenție”, gestionează secvențe lungi
 - NLP – Prelucrarea Limbajului Natural (toate LLM)
 - Interpretare text, traduceri

După funcție:

- Clasificare, Regresie, Segmentare, Generare, Traducere etc

După mod de învățare

- Supervizată, Nesupervizată, Reinforcement

După „profundimea” rețelei

- Shallow (1-2 straturi ascunse)
- Deep Networks
 - > 3 straturi ascunse
 - Aplicații: imagistică, NLP, sisteme generative, predicții complexe

Deep Learning

Conținut:

- *Deep learning* = utilizarea rețelelor neuronale cu **multe straturi ascunse**.
- Fiecare strat învață **reprezentări progresiv mai abstracte** ale datelor.
- Permite recunoașterea automată a tiparelor complexe fără programare explicită.
- **Imagine:** *date brute* → *trăsături* → *concepte* → *decizie*.

Cum învață o rețea profundă:

Fiecare strat ascuns extrage caracteristici mai abstracte.

În imagini:

Stratul 1 → detectează margini

Stratul 2 → recunoaște forme

Stratul 3 → identifică organe sau leziuni

În text:

Stratul 1 → recunoaște cuvinte

Stratul 2 → sensul propoziției

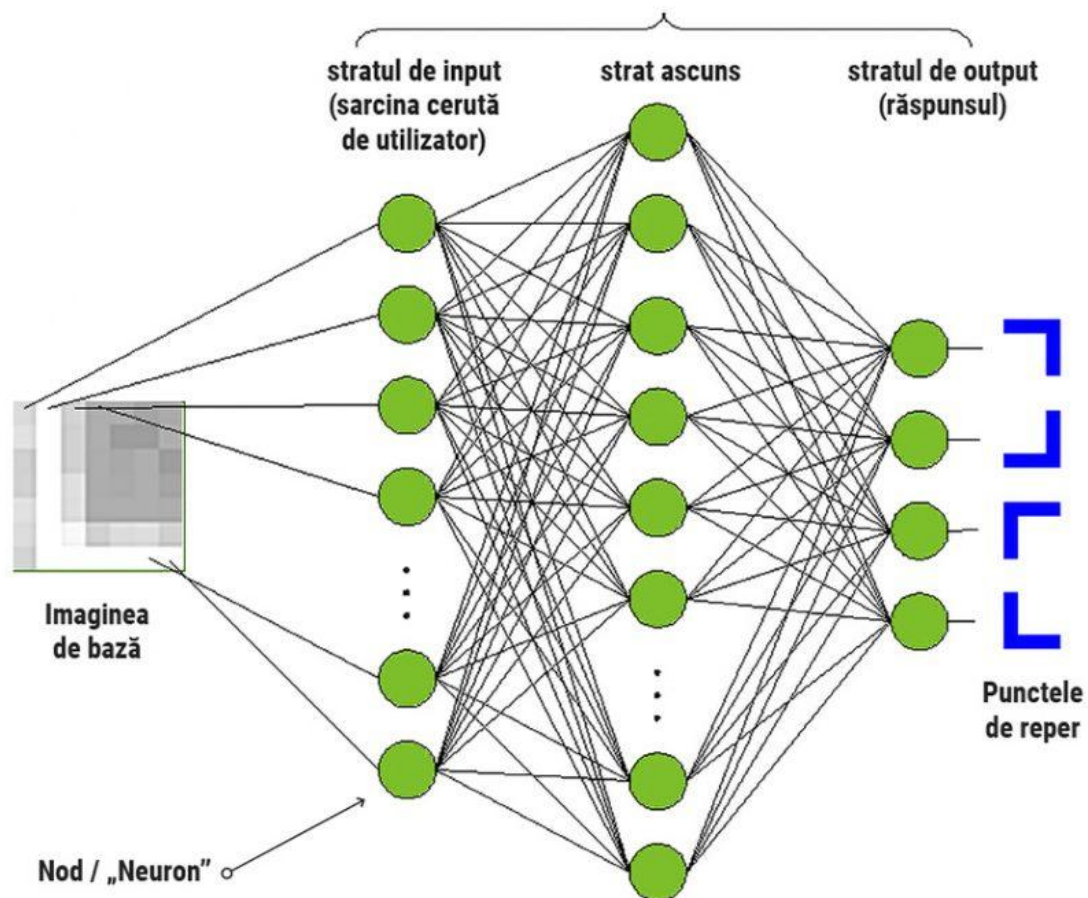
Stratul 3 → intenția / semnificația clinică

Aplicații medicale reprezentative

- **Exemple structurate pe domenii:**
 - ***Imagistică medicală*** – detectarea leziunilor, segmentare, clasificare tumorală.
 - ***Bioinformatică*** – predicția structurii proteinelor (AlphaFold).
 - ***Semnale fiziologice*** – analiza ECG, EEG, monitorizare intensivă.
 - ***Patologie digitală*** – clasificarea imaginilor histologice.
 - ***NLP medical*** – extragerea automată de date din rapoarte clinice.

CUM FUNCȚIONEAZĂ O REȚEA NEURONALĂ ARTIFICIALĂ

Structura unei rețele neuronale artificiale folosită
pentru a identifica puncte de reper într-o fotografie



SURSĂ: Cyberbotics Ltd.
/ Wikimedia Commons

Avantaje, limite și provocări

Avantaje	Limite
Recunoaște tipare complexe	„Cutie neagră” – decizie greu explicabilă
Învăță direct din date	Necesită volume mari de date etichetate
Performanță ridicată (imagini, text)	Risc de overfitting / bias
Extensibilitate (multimodală)	Costuri mari de antrenare

De vorbă cu AI despre ...

Cât de multe detalii privind Perceptronul ar fi potrivit să prezint auditoriului meu - rezidenți și studenți ani mari.

Excelentă întrebare — și foarte bine pusă în context.

Pentru un auditoriu de **rezidenți și studenți avansați la medicină**, scopul principal al secțiunii despre **Perceptron** nu este formarea în calcul numeric, ci **înțelegerea conceptuală**: *cum gândește* o rețea neuronală la nivelul cel mai simplu și *ce limite are acest model*.

Așadar, detalierea optimă este una **semi-formală**, fără matematică avansată, dar suficientă ca să le ofere o „intuiție inginerească” despre cum funcționează învățarea.

Iată un ghid clar, împărțit pe niveluri și recomandarea pentru nivelul ideal în cursul dvs.:

Nivel 1 – Esențial (obligatoriu pentru toți)

 5–6 minute – 2 slide-uri

Scop: să înțeleagă ideea.

Ce trebuie inclus:

- Reprezentarea grafică a perceptronului (2 intrări, un neuron, o ieșire binară).
- Ecuația conceptuală:

$$y = f(w_1x_1 + w_2x_2 + b)$$

fără a intra în calcule.

- Sensul fiecărui element:
 - w_i : importanța semnalului,
 - b : prag de activare,
 - $f()$: funcție decizională (treaptă).
- Explicarea vizuală (planul cu linia de separare între DA/NU).
- Un exemplu medical simplu: fumat + IMC $\rightarrow$ risc de hipertensiune (DA/NU).

⚙ Nivel 2 – Explicativ (recomandat pentru rezidenți și studenți avansați)

🕒 8–10 minute – 2 slide-uri suplimentare

Scop: să înțeleagă mecanismul de învățare.

Ce se poate adăuga:

- Formula simplă a **regulii de învățare**:

$$w_i^{nou} = w_i^{vechi} + \eta(y_{dorit} - y_{obținut})x_i$$

explicată intuitiv, nu numeric:

- Dacă perceptronul a greșit, el „corectează” puțin ponderile în direcția greșelii.
- η = „viteza de învățare”.
- Tabel simplificat cu 3–4 exemple de actualizări.
- Grafic cu linia de decizie care se deplasează după fiecare eroare.

💬 **Comentariu oral:**

„În mod esențial, perceptronul învață din greșeli: de fiecare dată când clasifică greșit un exemplu, își ajustează 'sinapsele'. Acesta e începutul ideii de *învățare automată*.”

Nivel 3 – Teoretic (opțional, pentru curioși sau întrebări din sală)

 doar menționat, fără detalieri

Scop: să arătați că știți unde e limita, dar fără a încărca.

Ce se poate menționa scurt:

- Perceptronul poate învăța doar **relații linear separabile**.
- Problema XOR: imposibil de rezolvat de un perceptron simplu.
- Necesitatea introducerii **straturilor ascunse** → nașterea *rețelelor neuronale multistrat*.

Comentariu scurt:

„Aceasta este limita perceptronului — exact punctul în care a fost criticat de Minsky și Papert. Ca să depășim această limitare, au apărut rețelele cu mai multe straturi — pasul natural spre deep learning.”

Cum aş putea verifica dacă, sau în ce măsură, audiența a apreciat pozitiv faptul că am prezentat și funcționarea perceptronului, adică dacă au pătruns "esența funcționării rețelelor neuronale" sau nu ?
Un răspuns scurt și concis ar fi bine apreciat!

Cel mai eficient mod: o **întrebare de verificare conceptuală**, imediat după secțiune.

Exemple:

„Care este ideea principală a perceptronului?”

- a) aplică reguli logice fixe
- b) învață prin ajustarea ponderilor
- c) memorează toate cazurile

→ Dacă majoritatea aleg (b) sau o formulează spontan, **au înțeles esența**.

Alternativ, la final:

„Explicați în 10 secunde cum 'învață' o rețea neuronală.”

Dacă răspunsurile conțin „corectează erorile / modifică ponderile”, mesajul a fost asimilat.



ș.a.m.d...

„A system that is not understood cannot be trusted”

*Math. Timothy Gowers (n. 1963)
Univ. of Cambridge & Paris*

