

Intelligenza Artificiale per Sistemi Industriali

Lezione 05 — 10 Marzo 2026

Alice Plebe

www.aliceplebe.com

alice.plebe@unitn.it

Oltre al Perceptron

- *Neocognitron*, Kunihiro Fukushima, 1980
- *Boltzmann Machine*, Geoffrey Hinton & Terry Sejnowski, 1985
- *Self-organizing maps*, Teuvo Kohonen, 1982

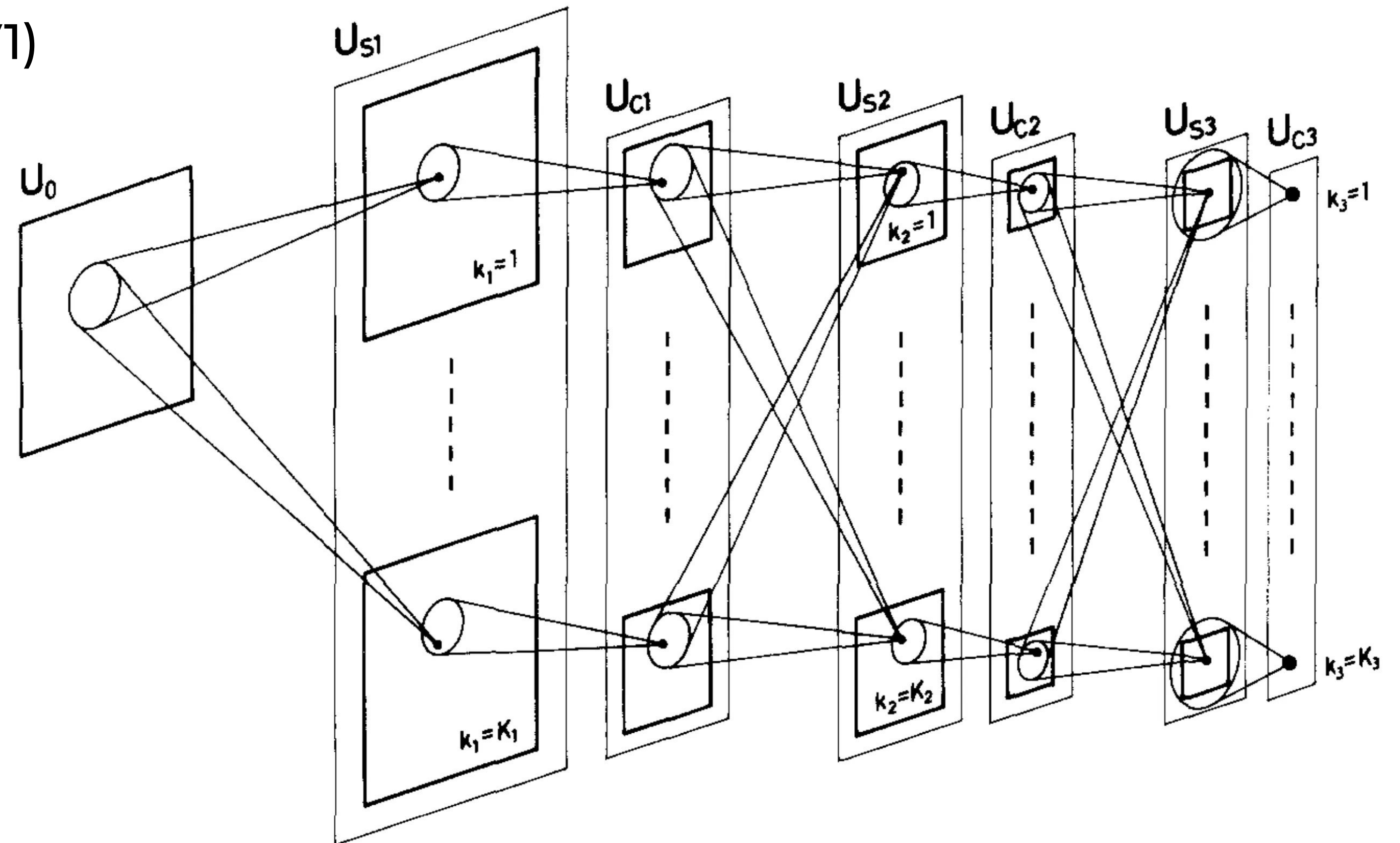
- Per riconoscimento di pattern visivi con una certa invarianza spaziale

- Ispirazione dalla corteccia visiva (V1)

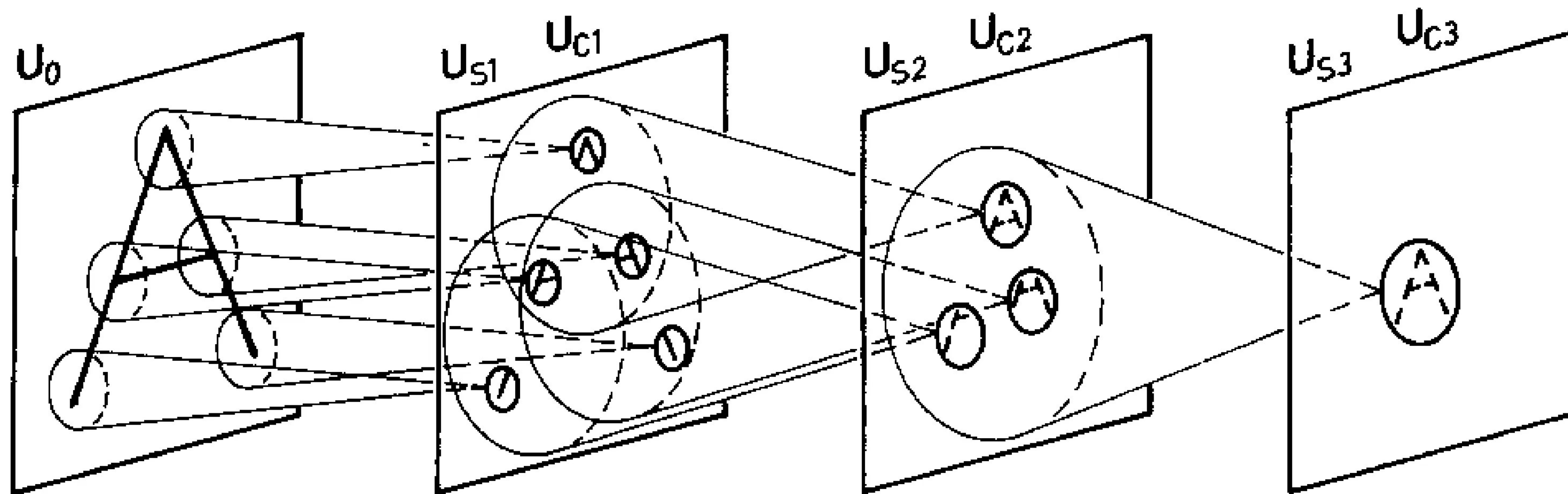
- Antenato delle CNNs

- **S-cells:** caratteristiche locali

- **C-cells:** aggregazione di S-cells



- **Gerarchizzazione** (S-cells) + **pooling** (C-cells)
- Invarianza locale
- Niente backpropagation! Apprendimento ispirato alla regola di Hebb



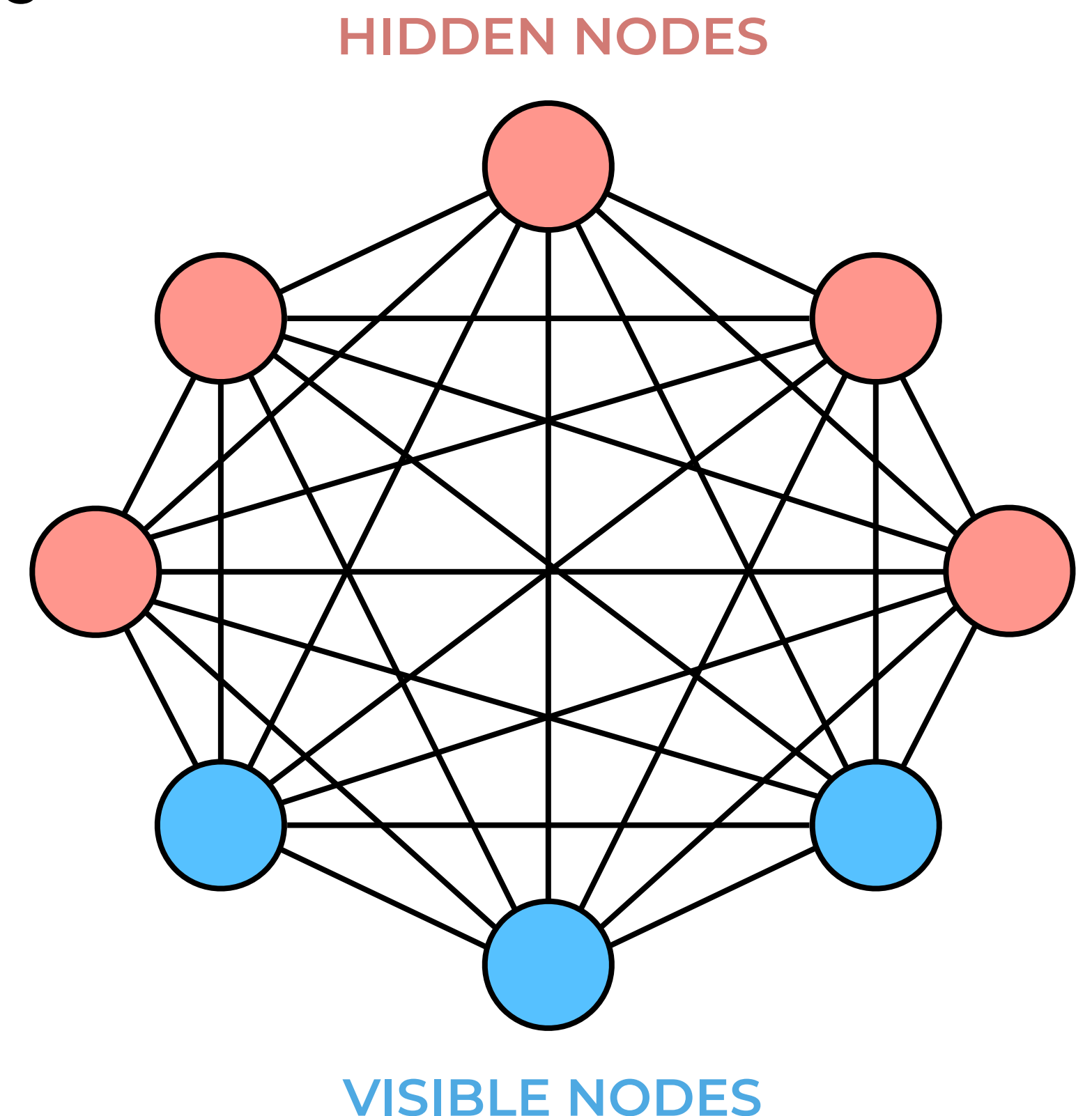
- Per modellare distribuzione di dati, no classificazione
- Rete completamente connessa, non direzionata, con pesi simmetrici
- Funzione di energia, dove $\mathbf{s} = (s_1, \dots, s_n)$ è lo stato della rete

$$E(\mathbf{s}) = - \sum_{i < j} w_{ij} s_i s_j - \sum_i b_i s_i$$

- La probabilità sulle configurazioni è determinata dalla distribuzione di Boltzmann

$$P(\mathbf{s}) = \frac{1}{Z} e^{-E(\mathbf{s})}$$

dove $Z = \sum_{\mathbf{s}} e^{-E(\mathbf{s})}$



FUNZIONAMENTO

- Pesi già fissati
- Generare configurazioni tramite Gibbs sampling che tendono all'equilibrio termico

$$P(s_i = 1 \mid s_{-i}) = \frac{1}{1 + \exp\left(-\sum_j w_{ij} s_j - b_i\right)}$$

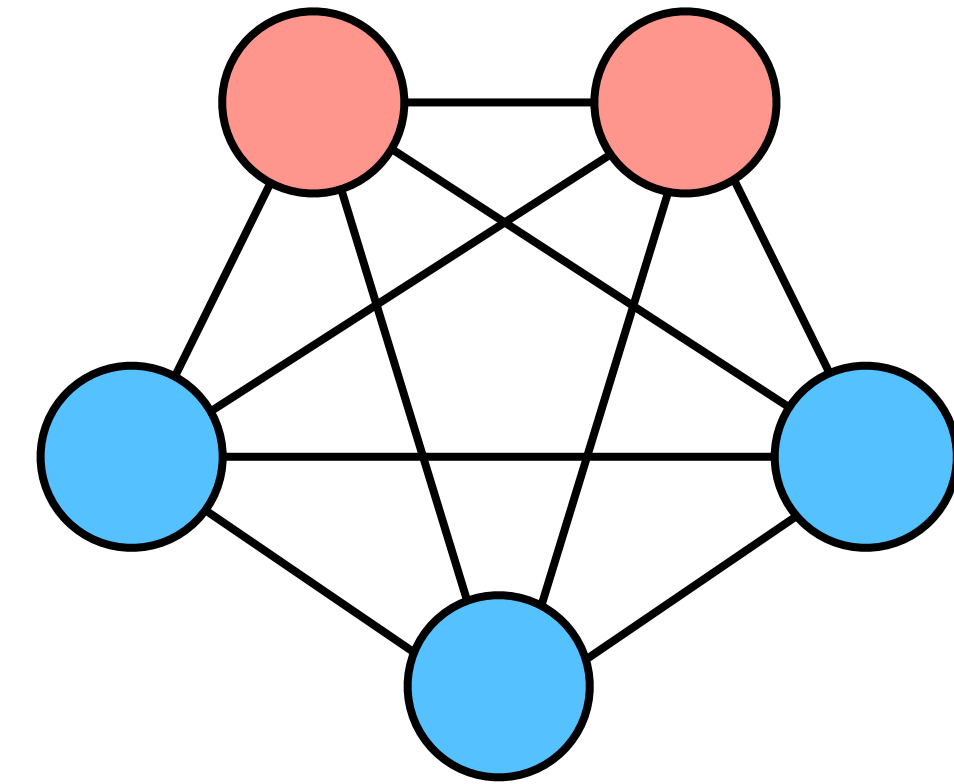
ADDESTRAMENTO

- Modificare i pesi per far coincidere distribuzione di equilibrio con distribuzione dei dati
- Differenza tra correlazione tra i e j con rete "vincolata" e rete "libera"

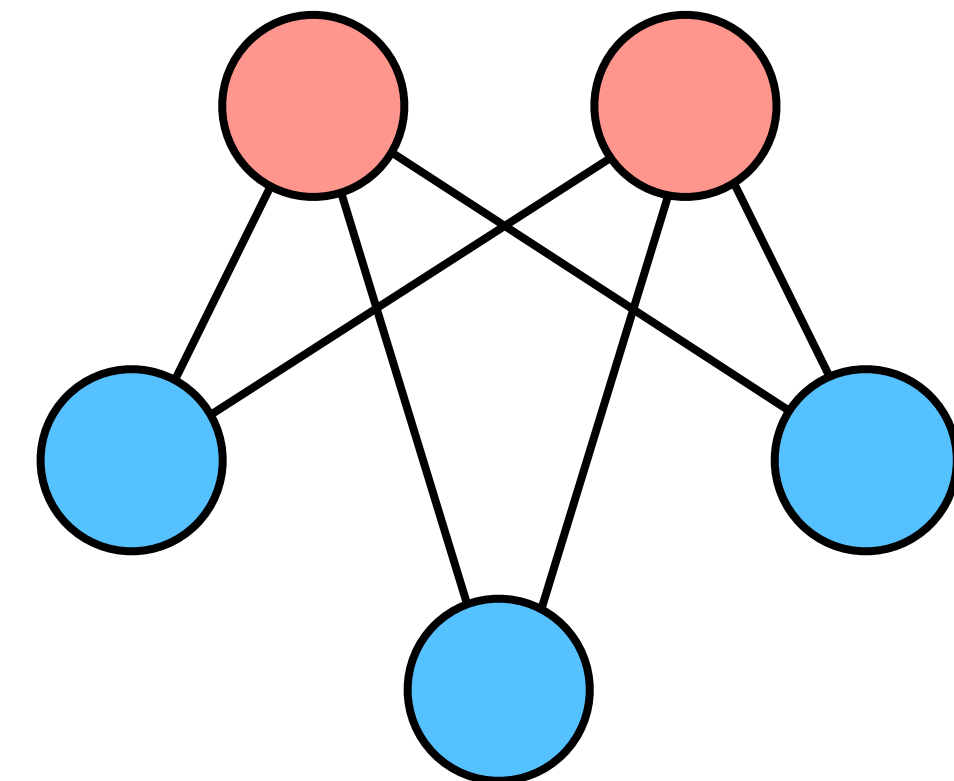
$$\Delta w_{ij} = \eta \left(\langle s_i s_j \rangle_V - \langle s_i s_j \rangle_L \right)$$

- Addestramento di BM "complete" instabile e computazionalmente troppo costoso
- *Restricted* Boltzmann machine come alternativa trattabile
- Separazione tra layer visibile e layer nascosto

GENERAL
BOLTZMANN MACHINE



RESTRICTED
BOLTZMANN MACHINE



SELF-ORGANIZING MAPS

9

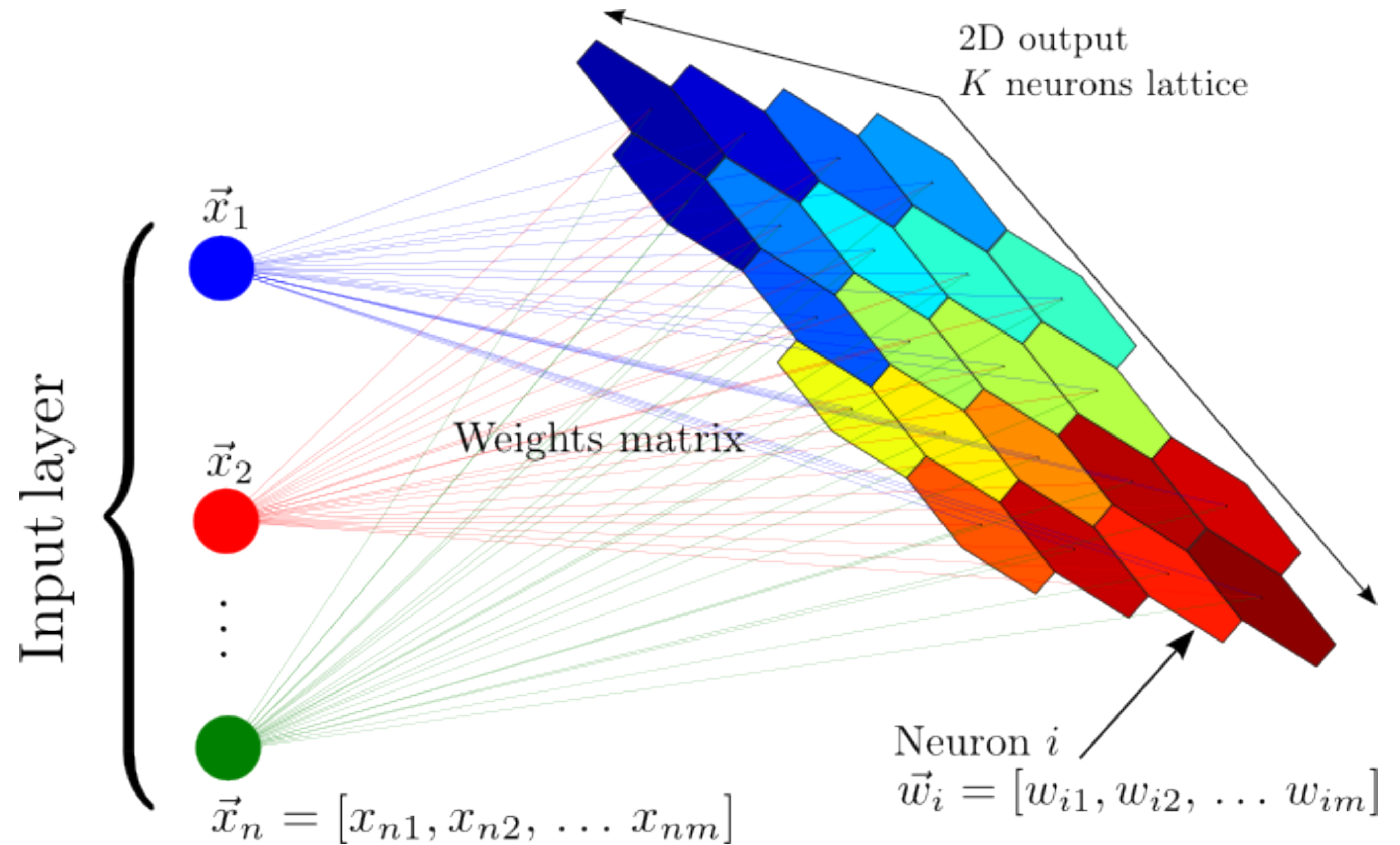
- Per riduzione dimensionale
- No layer, singola griglia di neuroni

- Apprendimento *winner-takes-all*

$$c = \arg \min_i \| \mathbf{x} - \mathbf{w}_i \|$$

- Ma si aggiornano i pesi di tutti i neuroni vicino al vincitore

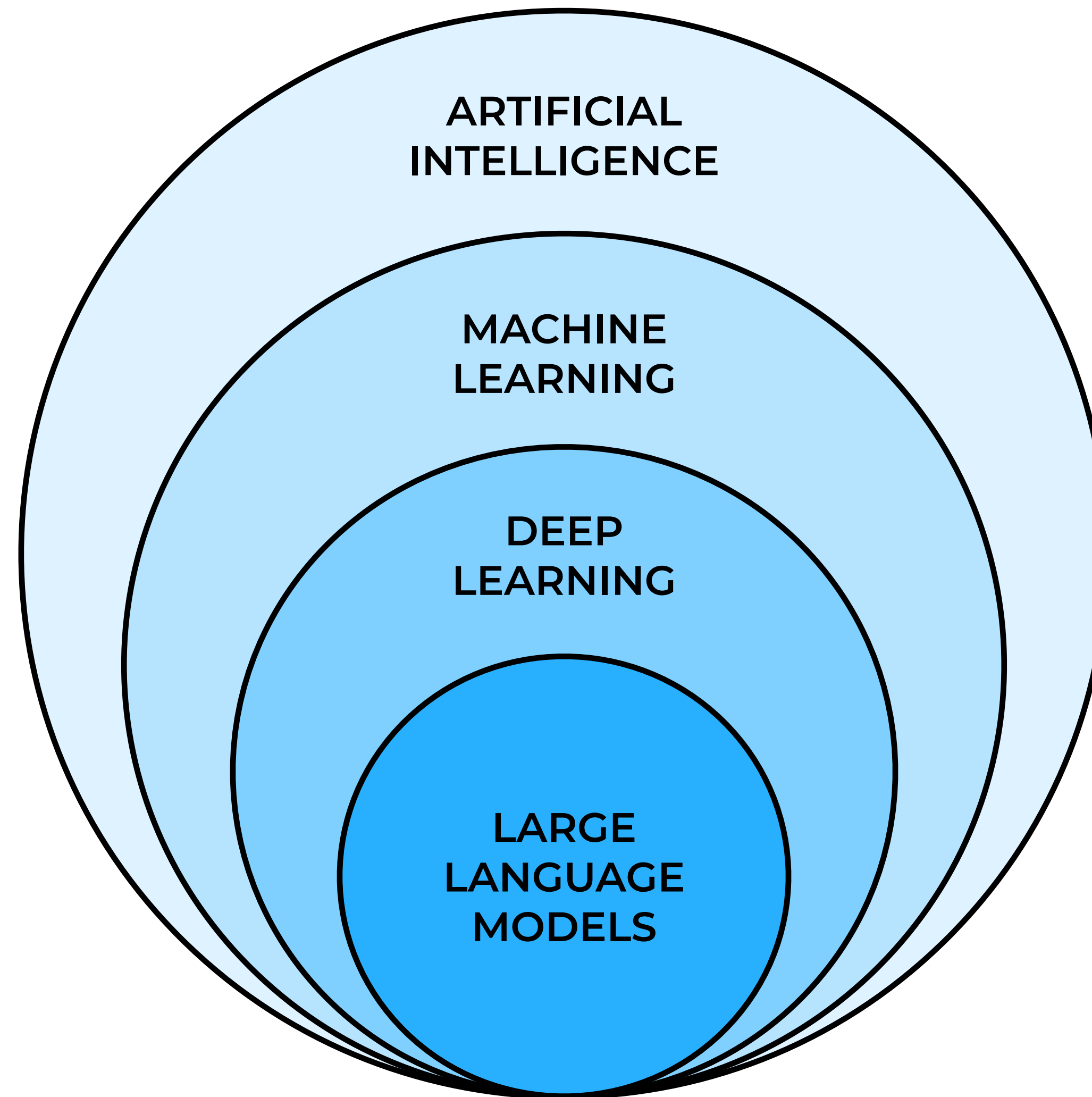
$$\Delta \mathbf{w}_i = \eta \exp \left(-\frac{\| \mathbf{r}_c - \mathbf{r}_i \|^2}{2\sigma^2} \right) (\mathbf{x} - \mathbf{w}_i)$$



Reti artificiali moderne

CLASSIFICAZIONE ODIERNA

11



RETE COME FUNZIONE PARAMETRICA

$$\hat{y} = f_{\theta}(\mathbf{x})$$

FUNZIONE DI LOSS

$$L(\theta) = \frac{1}{N} \sum_{i=1}^N \ell(f_{\theta}(\mathbf{x}_i), \mathbf{y}_i)$$

COMPOSIZIONE DI HIDDEN LAYERS

$$\mathbf{h}_1 = \varphi(\mathbf{W}_1 \mathbf{x} + \mathbf{b}_1)$$

$$\mathbf{h}_2 = \varphi(\mathbf{W}_2 \mathbf{h}_1 + \mathbf{b}_2)$$

...

$$\mathbf{h}_M = \varphi(\mathbf{W}_M \mathbf{h}_{M-1} + \mathbf{b}_M)$$

OBIETTIVO DEL TRAINING

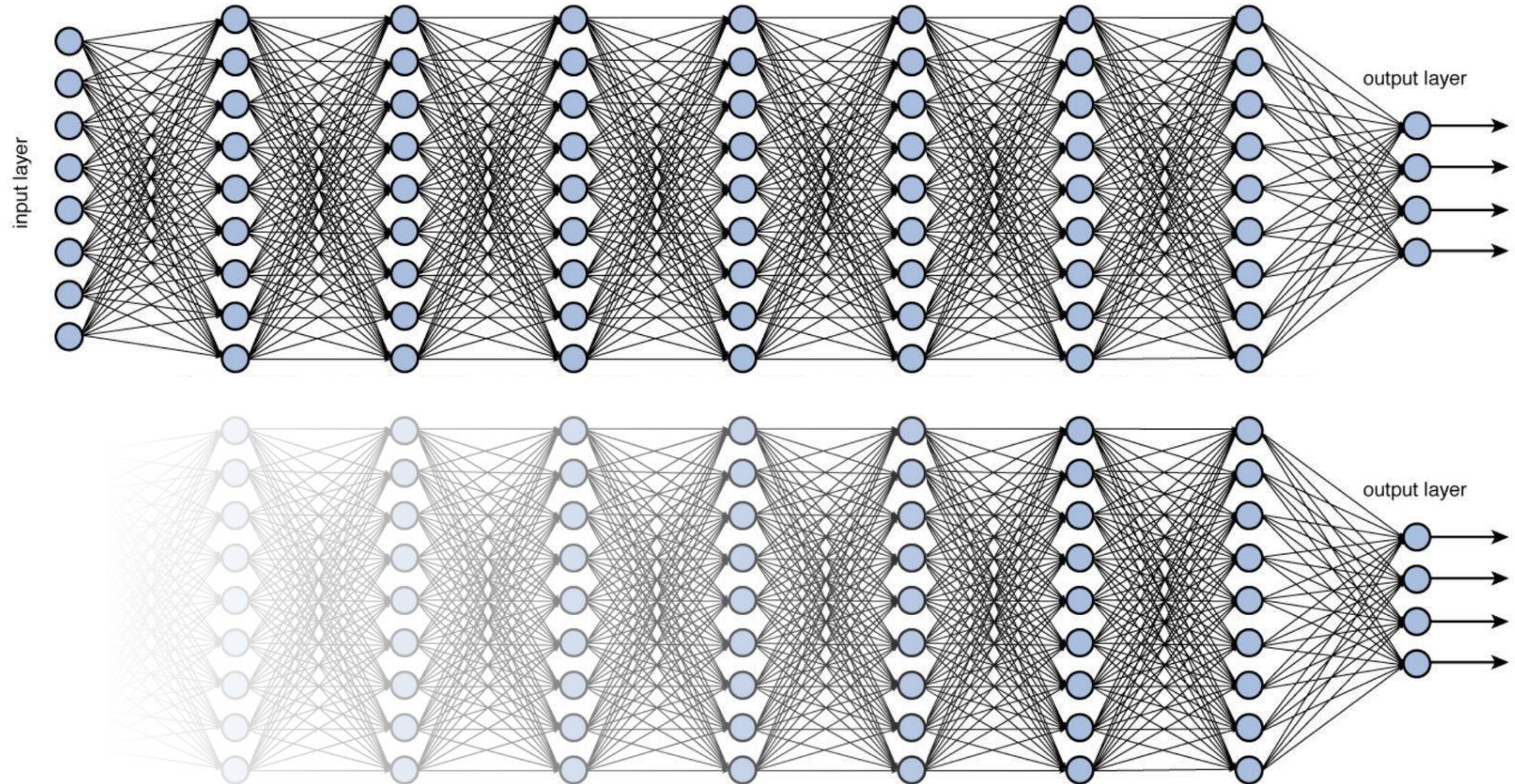
$$\min_{\theta} L(\theta)$$

DISCESA DEL GRADIENTE

$$\theta \leftarrow \theta - \eta \nabla_{\theta} L(\theta)$$

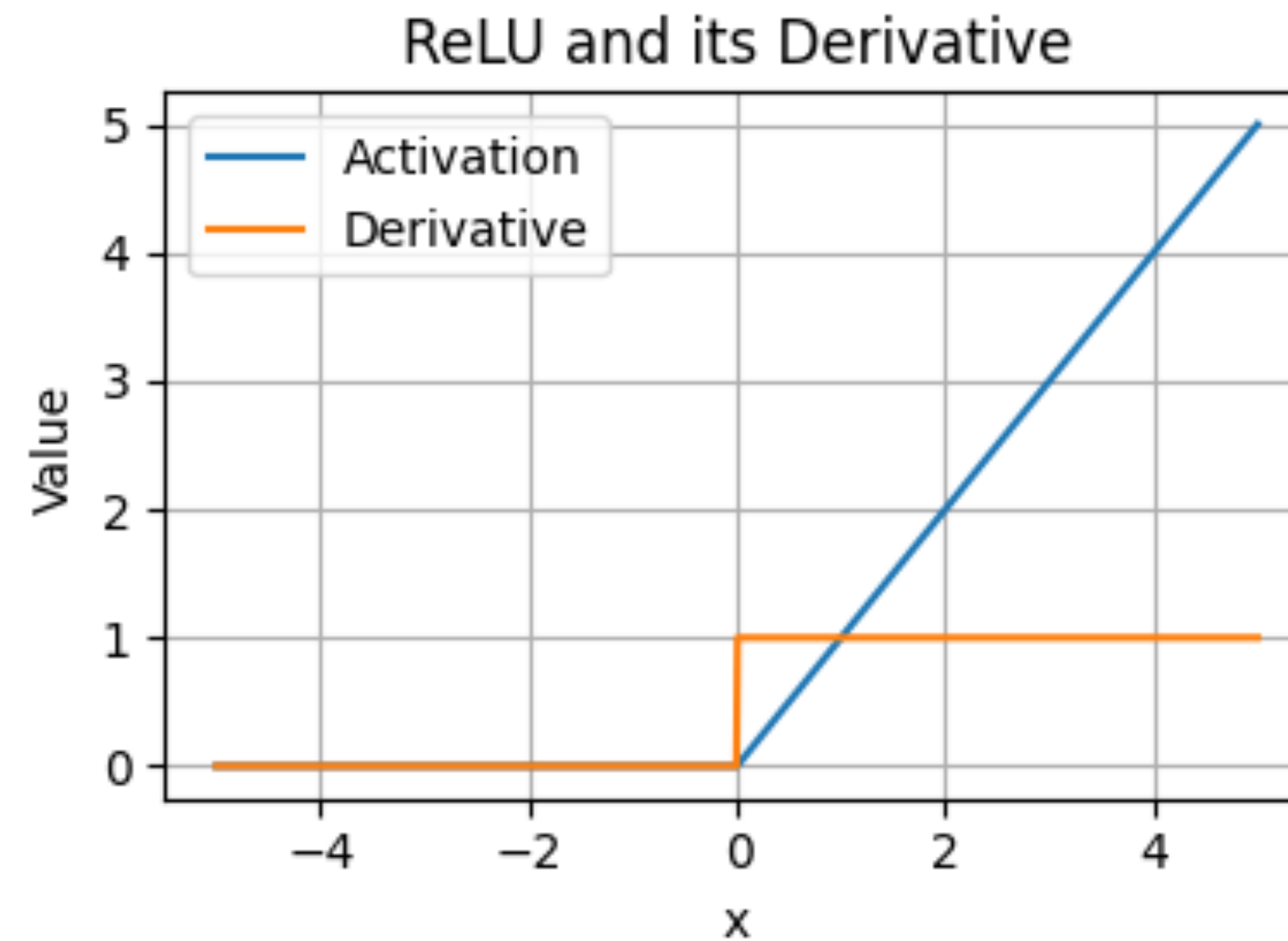
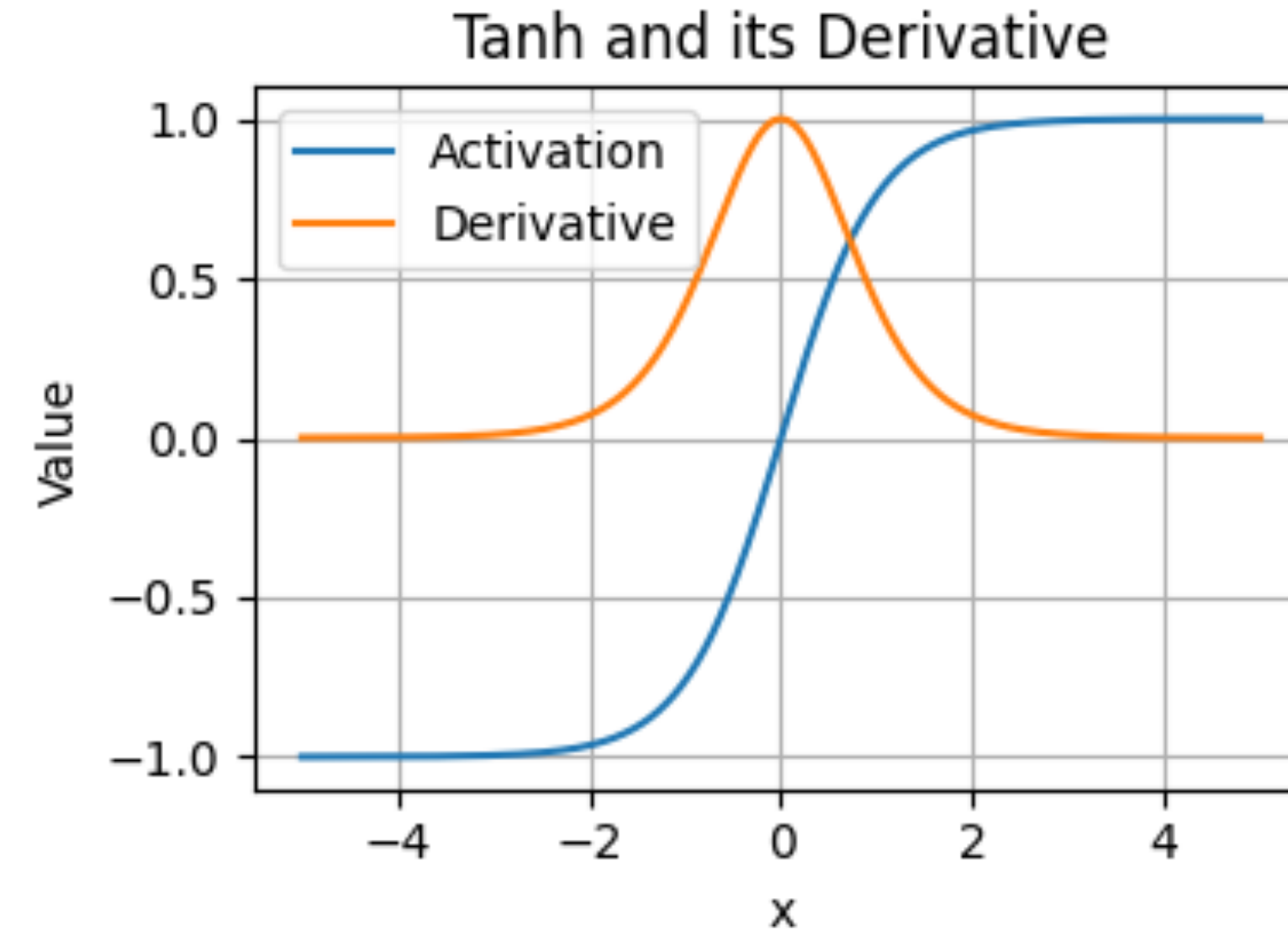
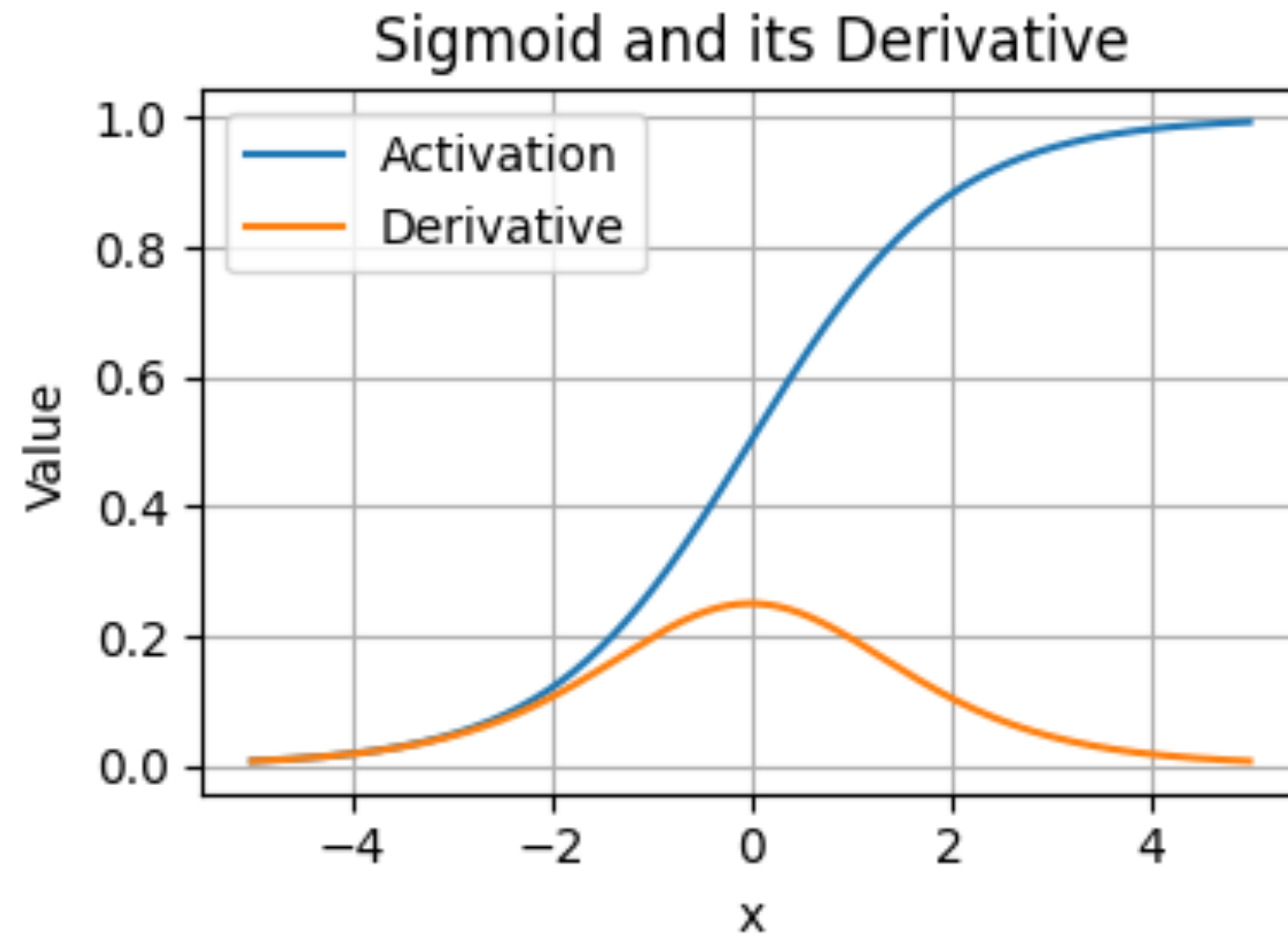
VANISHING GRADIENT PROBLEM

13



ACTIVATION FUNCTIONS

14



MINI-BATCH / STOCHASTIC GRADIENT DESCENT

15

FUNZIONE DI LOSS

$$L(\theta) = \frac{1}{N} \sum_{i=1}^N \ell(f_{\theta}(\mathbf{x}_i), \mathbf{y}_i)$$

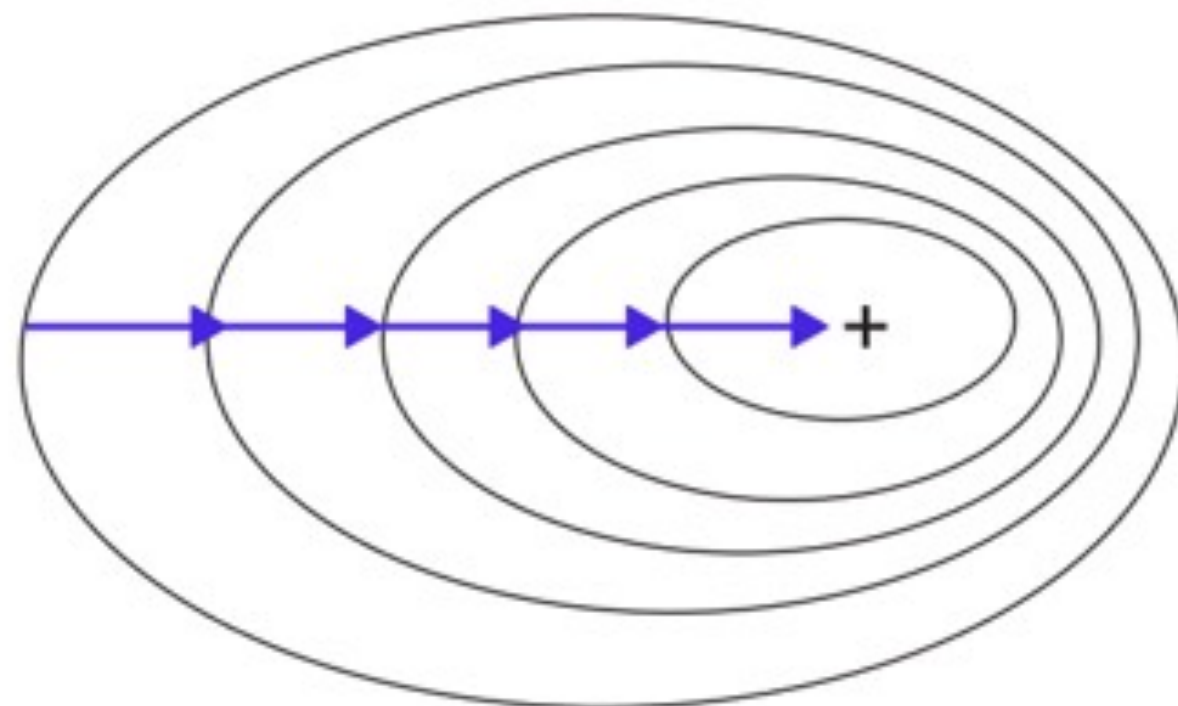
AGGIORNAMENTO SGD

$$\theta \leftarrow \theta - \eta \nabla_{\theta} \ell_i(\theta)$$

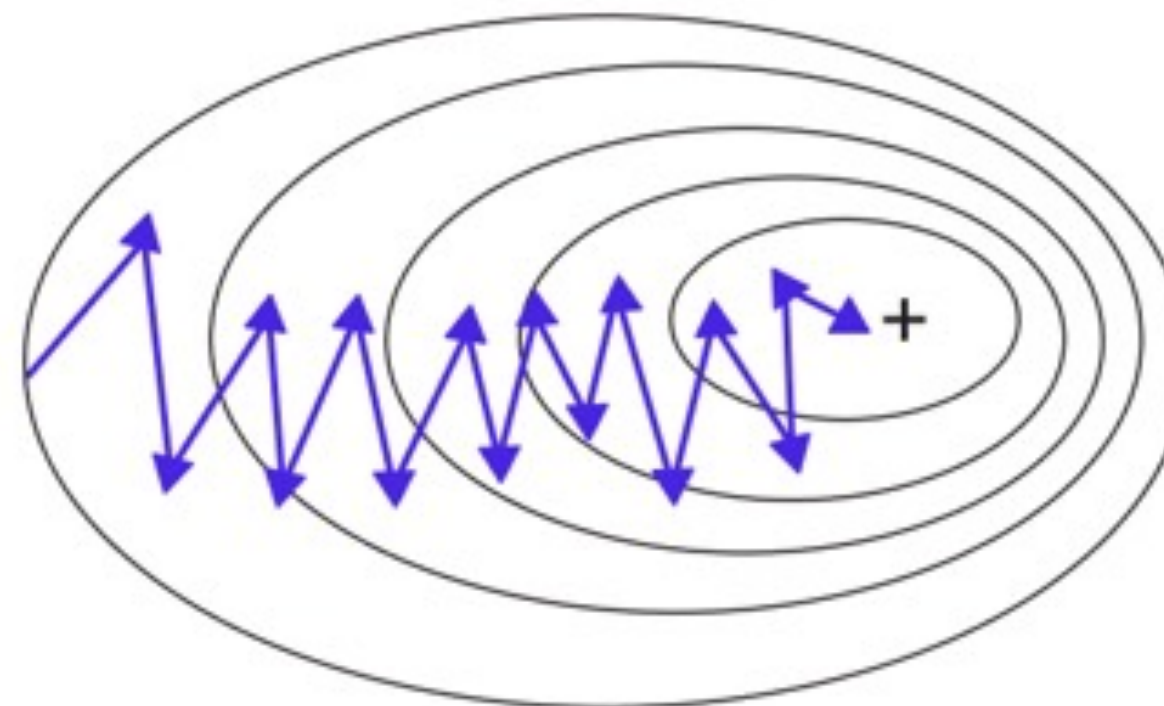
USARE $\nabla_{\theta} \ell_i(\theta)$ **COME STIMA DI** $\nabla_{\theta} L(\theta)$

MINI-BATCH GRADIENT DESCENT

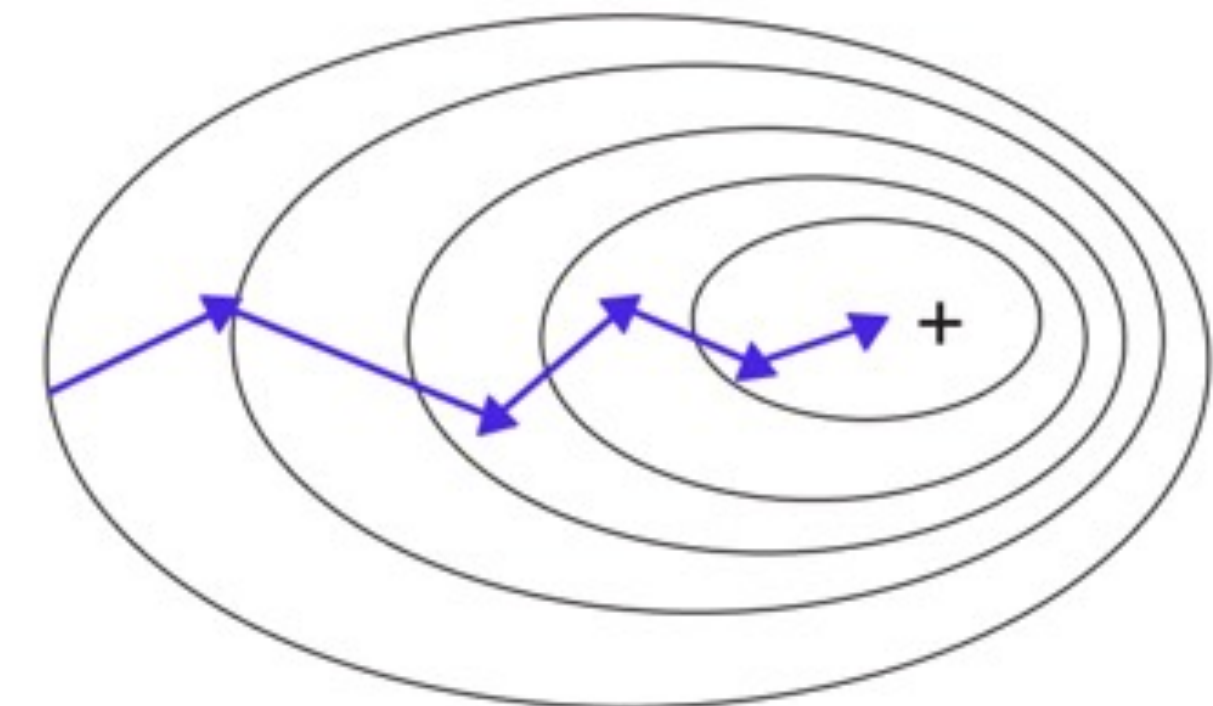
$$\nabla_{\theta} L_B(\theta) = \frac{1}{m} \sum_{i \in B} \nabla_{\theta} \ell_i(\theta)$$



Batch Gradient Descent



Stochastic Gradient Descent



Mini-Batch Gradient Descent

- Backpropagation rende possibile calcolare il gradiente in reti a più strati
- ReLU rende il gradiente stabile lungo molti strati
- Mini-batch SGD rende l'ottimizzazione scalabile su grandi dataset