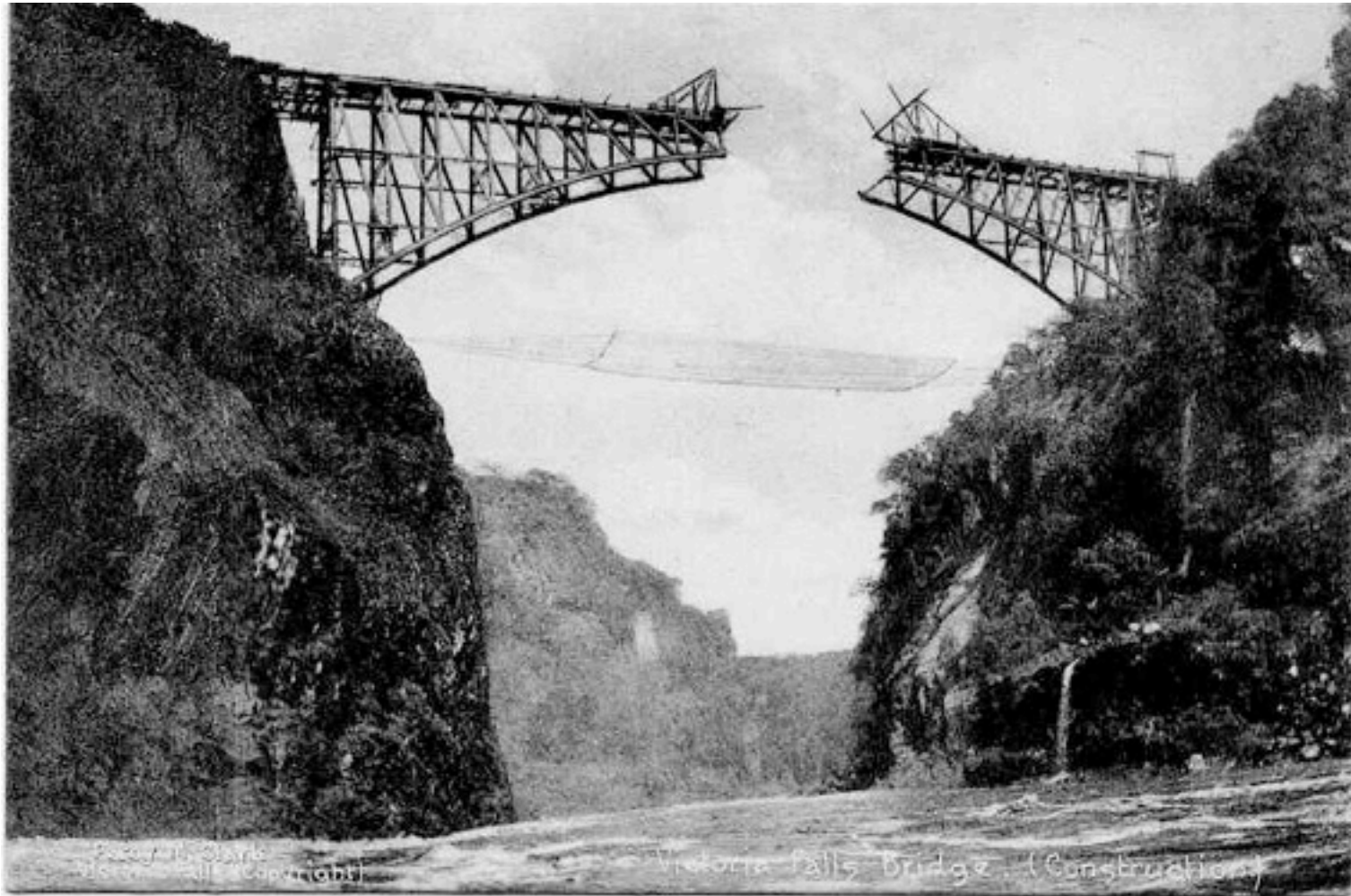


Tanulás az idegrendszerben



Structure – Dynamics – Implementation – Algorithm – Computation - Function

Funkcióvezérelt modellezés

- Abból indulunk ki, hogy milyen feladatot valósít meg a rendszer
- Horace Barlow: "A wing would be a most mystifying structure if one did not know that birds flew."
- Teleologikus jellegű megközelítés

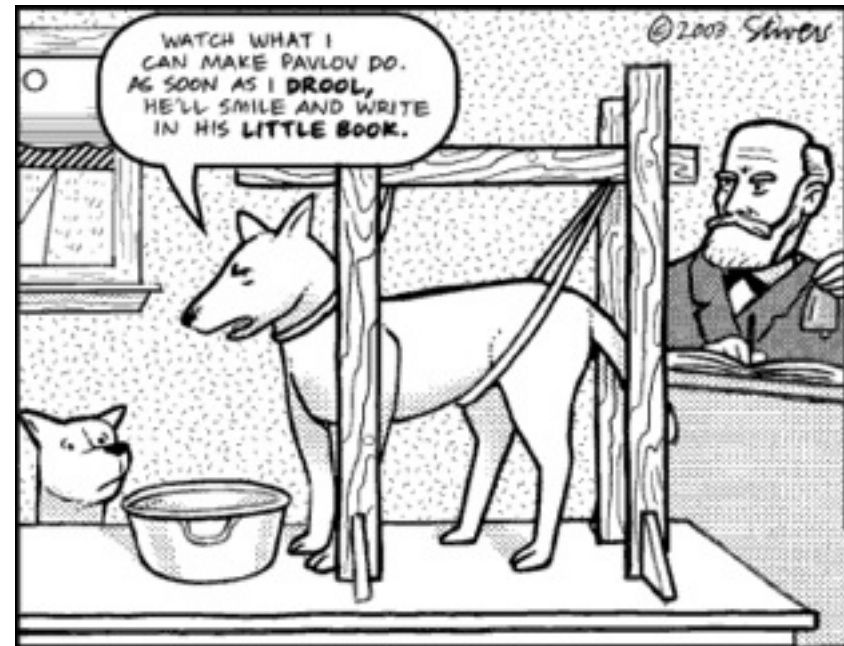
Az idegrendszer funkciója

- Rodolfo Llinás: az egyetlen funkció a mozgatus
- Adaptív kontroll
- Helmholtz: tudattalan következtetés



Tanulás pszichológiai szinten

- Classical conditioning



- Hebb ötlete:
"Ha az A sejt axonja elég közel van a B sejthez, és ismétlődően vagy folyamatosan hozzájárul annak tüzeléséhez, akkor valamely, egyik vagy mindkét sejtre jellemző növekedési folyamat vagy metabolikus változás következménye az lesz, hogy az A sejt hatékonysága a B sejt tüzeléséhez való hozzájárulás szempontjából megnő."

Kölcsönhatás az MI-vel

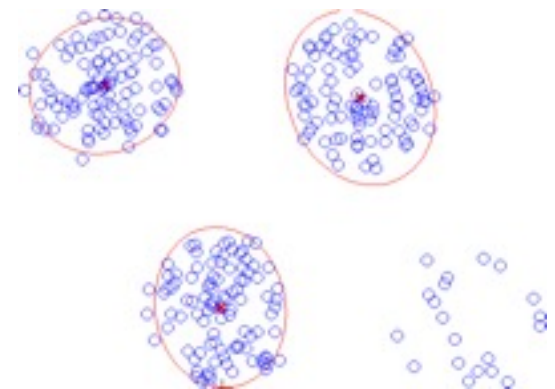
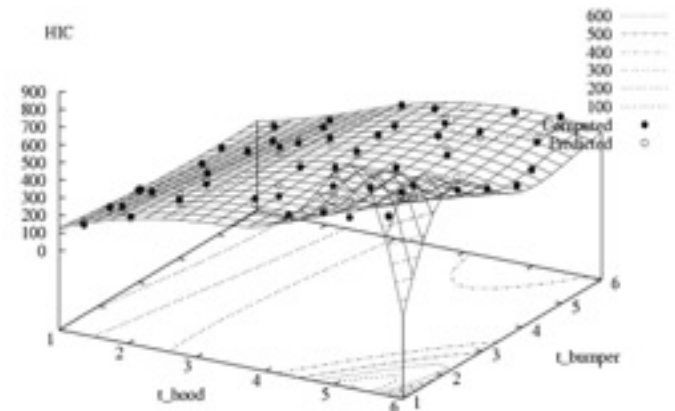
- Az MI azóta machine learning lett (plusz data mining, stb.)
- Eleinte az idegtudomány inspirálta az MI-t
 - mesterséges neuronhálózatok
- Az információáramlás jelenleg főként fordított
 - ML módszerek az idegi funkciók modellezésére

A tanulás problémája matematikailag

- Modell paramétereinek hangolása adatok alapján
- Kettős dinamika
 - Változók (bemenet-kimenet leképezés) - gyors
 - Paraméterek - lassú
- Memória és tanulás különbsége
 - Memória használatánál a bemenetre egy konkrét kimenetet szeretnék kapni a reprezentáció megváltoztatása nélkül
 - Tanulásnál minden bemenetet felhasználok arra, hogy finomítsam a reprezentációt, miközben kimenetet is generálok

A tanulás alapvető típusai

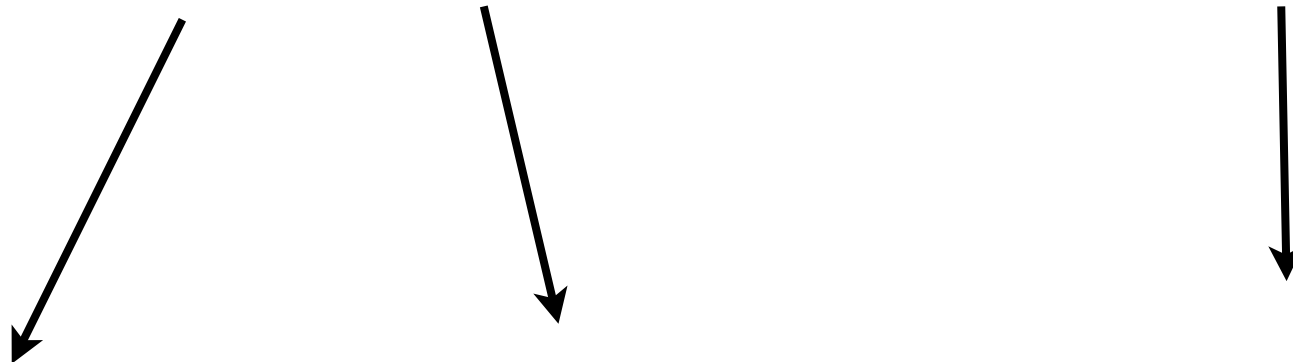
- Felügyelt
 - Az adat: bemenet-kimenet párok halmaza
 - A cél: függvényapproximáció, klasszifikáció
- Megerősítéses
 - Az adat: állapotmegfigyelések és jutalmak
 - A cél: optimális stratégia a jutalom maximalizálására
- Nem felügyelt, reprezentációs
 - Az adat: bemenetek halmaza
 - A cél: az adat optimális reprezentációjának megtalálása / magyarázó modell felírása
- Egymásba ágyazások



Lényegkiemelés

- Feature selection
 - A lényegtelen információ, zaj elhagyása
 - Nem a tanulás célja, de nagyon megkönnyíti a célfüggvény reprezentációját
- Invarianciák tanulása
 - például vizuális objektumfelismerésnél látószög, távolság, ...
 - valós és véges megfigyelésszámból következő invarianciák elkülönítése
- A kimenet és a bemenet közötti köztes reprezentáció, rejtett változók

Computational neuroscience



Az agy
számításokat
végez

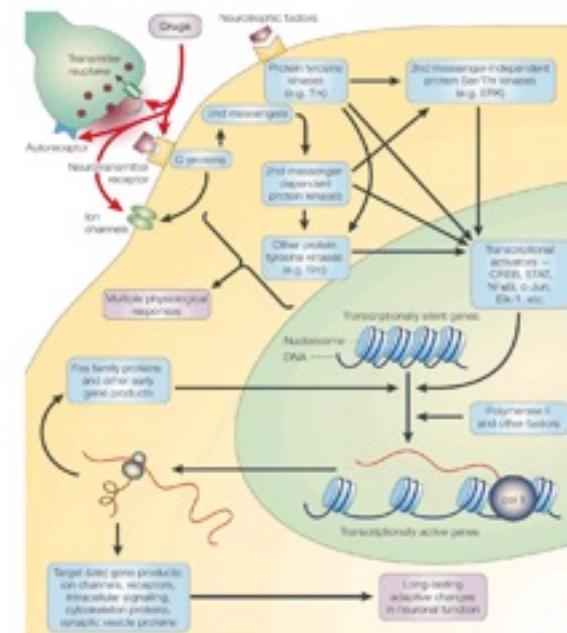
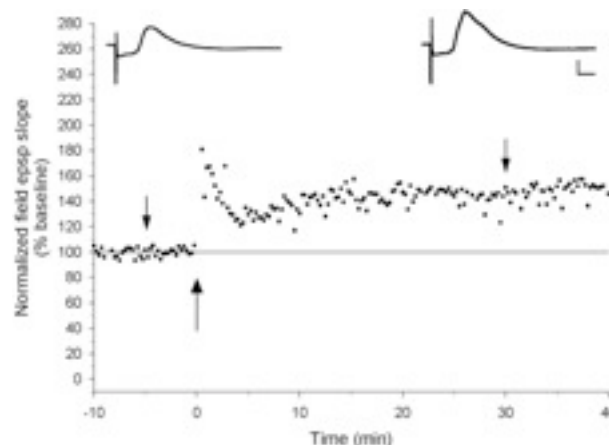
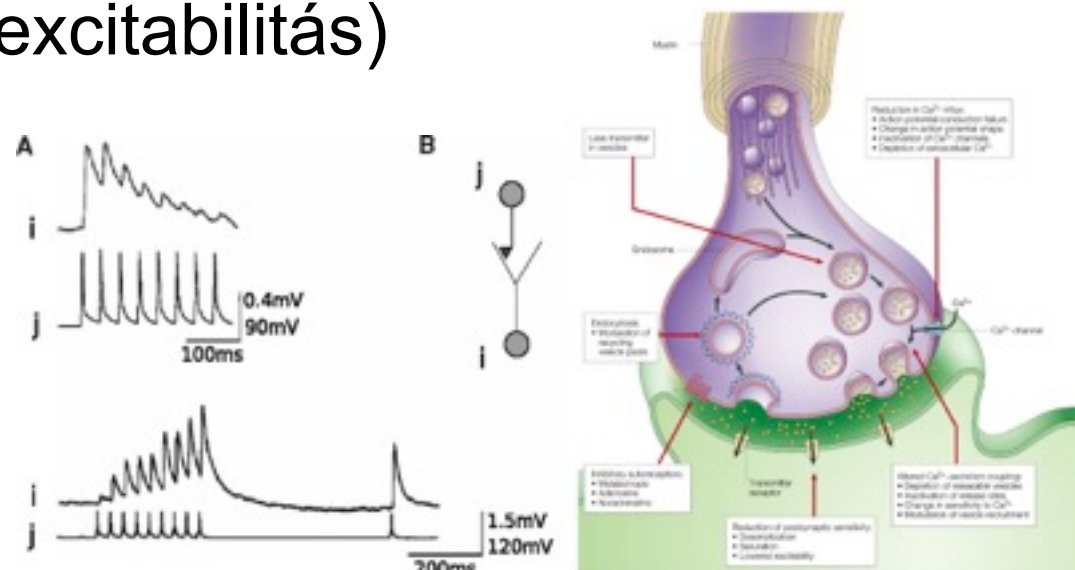
Numerikus
számításokat
végzünk az agy
modellezése
céljából

Köze van az
idegrendszerhez:

- neurokémia
- sejtek dinamikája
- neuronhálózatok
- agyterületek
- viselkedés

Idegrendszeri plaszticitás

- A plaszticitás helye: szinapszisok, posztszinaptikus sejtek tüzelési küszöbei (excitabilitás)
- Potenciáció, depresszió
- STP: kalciumdinamika, transzmitterkimerülés tartam < 1 perc
- LTP: génexpresszió (induction, expression, maintenance), NMDA magnézium-blokkja tartam > 1 perc



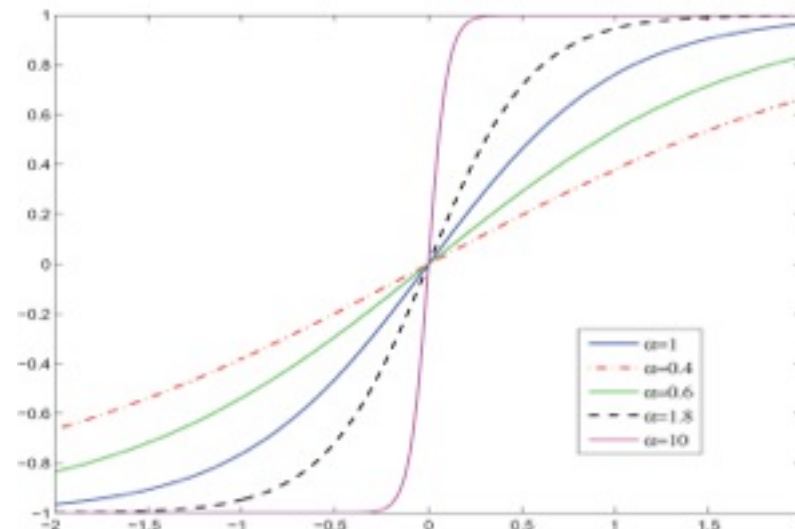
Tanulásra alkalmas neurális rendszerek

- Egyetlen sejt
- Előrecsatolt hálózat
- Rekurrens hálózat
- Ezen az órán: rátamodell



$$\nu = f(\mathbf{uw} - \theta)$$

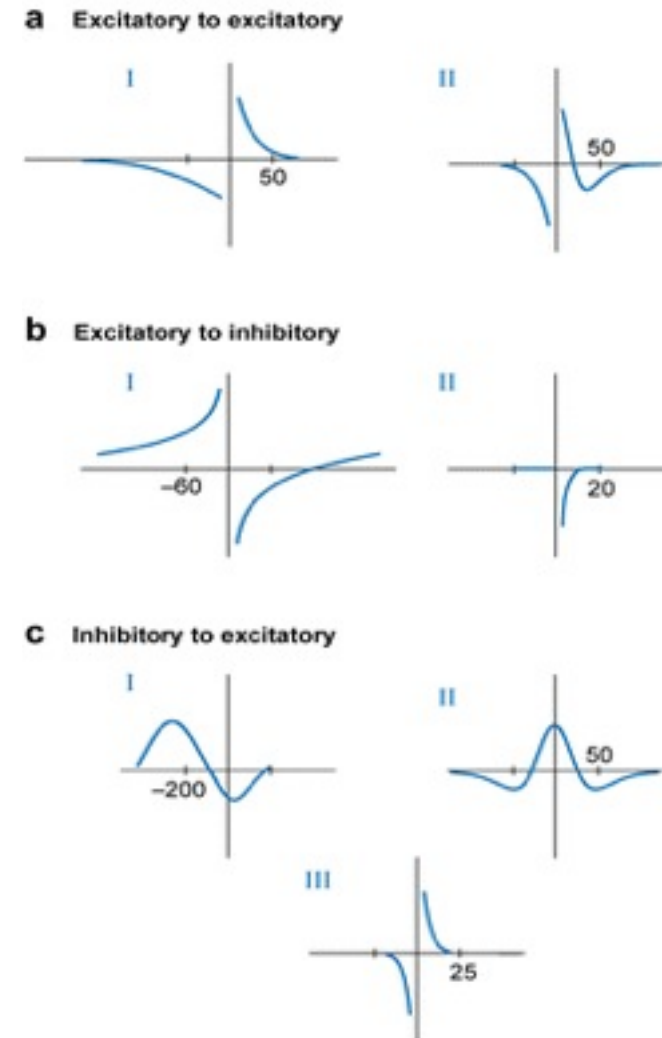
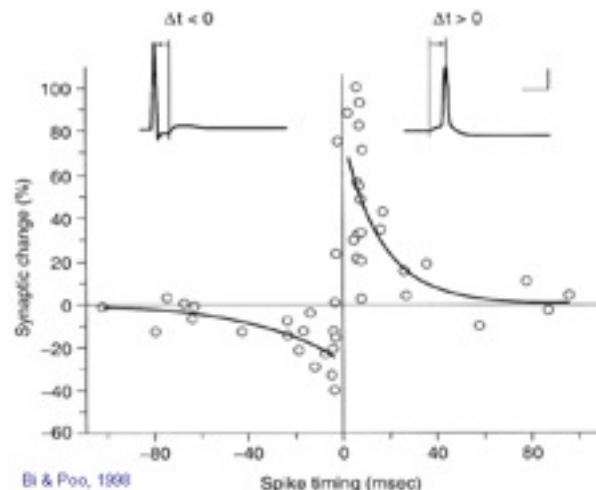
- Paraméterek: súlyok, küszöbök
- Különböző kimeneti nemlinearitások
 - Lépcső
 - Sigmoid
 - Lineáris neuron



A Hebb-szabály

- Timing-dependent plasticity:
 - Ha a posztzinaptikus neuron nagy frekvenciával közvetlenül a preszinaptikus után tüzel, akkor erősödik a kapcsolat
 - Az alacsony frekvenciájú tüzelés gyengíti a kapcsolatot
 - Sok más lehetőség

$$\tau \frac{dw}{dt} = v \mathbf{u}$$



Caporale N, Dan Y. 2008.
Annu. Rev. Neurosci. 31:25–46.

Stabilizált Hebb-szabályok

- Problémák a Hebb szabállyal:
 - csak nőni tudnak a súlyok
 - Nincs kompetíció a szinapszisok között – inputszelektivitás nem megvalósítható
- Egyszerű megoldás: felső korlát a súlyokra

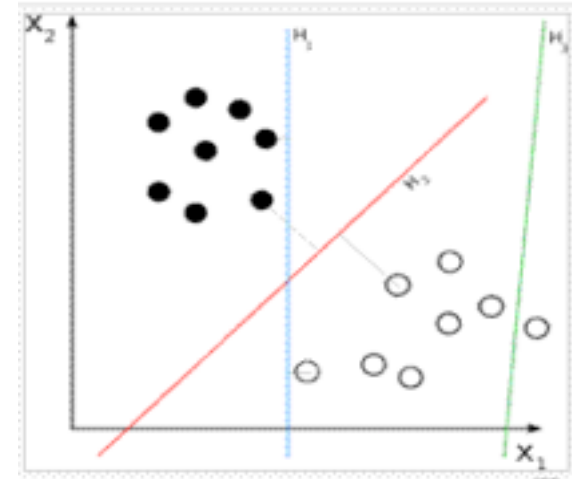
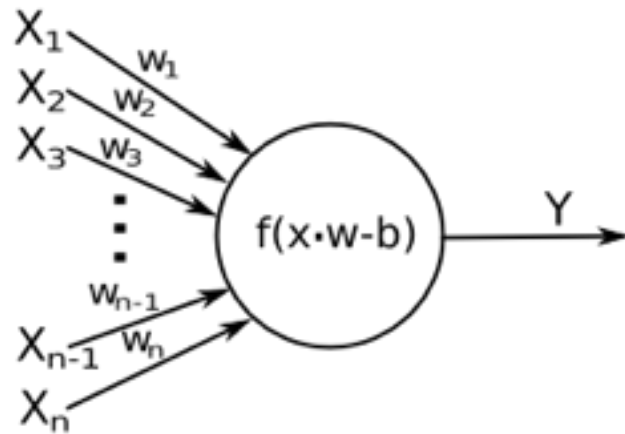
- Szinaptikus normalizáció
$$\tau_w \frac{d\mathbf{w}}{dt} = v\mathbf{u} - \frac{v(\mathbf{1} \cdot \mathbf{u})\mathbf{1}}{N_u}$$
 - Szubsztraktív normalizáció

Globális szabály, de generál egyes megfigyelt mintázatokat (Ocular dominance)

- Oja-szabály
$$\tau_w \frac{d\mathbf{w}}{dt} = v\mathbf{u} - \alpha v^2 \mathbf{u}$$

Lokális szabály, de nem generálja a megfigyelt mintázatokat

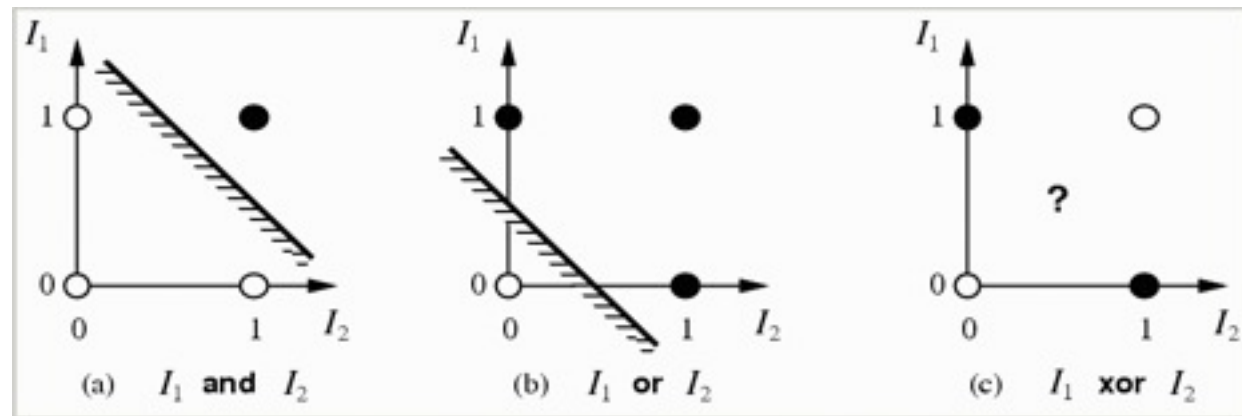
Perceptron



- Bináris neuron: lineáris szeparáció
 - Két dimenzióban a szeparációs egyenes:

$$\theta = x_1 w_1 + x_2 w_2 \longrightarrow x_2 = \frac{-w_1}{w_2} x_1 + \frac{\theta}{w_2}$$

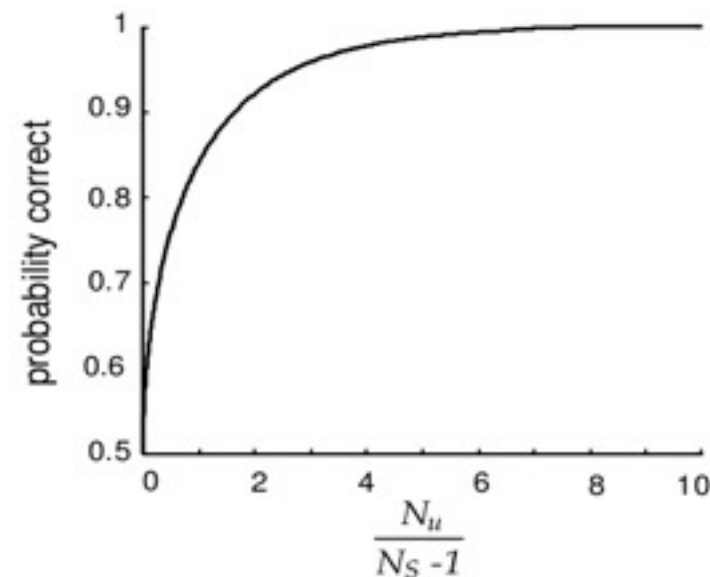
- Logikai függvények



Hebbi tanulás perceptronnal

- Felügyelt
 - A bemenetet és a kimenetet is meghatározzuk
 - Nem minden tanítóhalmazra lehet így megkonstruálni a szeparáló súlyokat
 - Hogyan lehet jobb szabályt konstruálni?

$$\tau_w \frac{d\mathbf{w}}{dt} = \frac{1}{N_s} \sum_{m=1}^{N_s} v^m \mathbf{u}^m$$



Error-correcting tanulási szabályok

- Felhasználjuk azt az információt, hogy milyen messze van a céltól a rendszer
- Rosenblatt-algoritmus – bináris neuron

$$\mathbf{w} = \mathbf{w} + \epsilon(v^m - v(\mathbf{u}^m))\mathbf{u}^m$$

- Delta-szabály

- Folytonos kimenetű neuron – gradiens-módszer

$$w_b = w_b - \epsilon \frac{\partial E}{\partial w_b} \quad E = \frac{1}{2} \sum_{m=1}^{N_s} (v^m - v(\mathbf{u}^m))^2 \quad \frac{\partial E}{\partial w_b} = - \sum_{m=1}^{N_s} (v^m - v(\mathbf{u}^m))\mathbf{u}^m$$

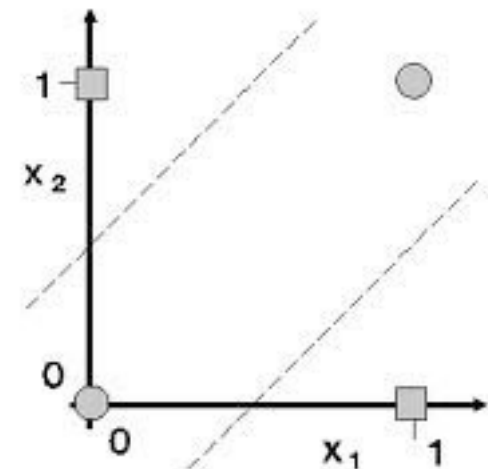
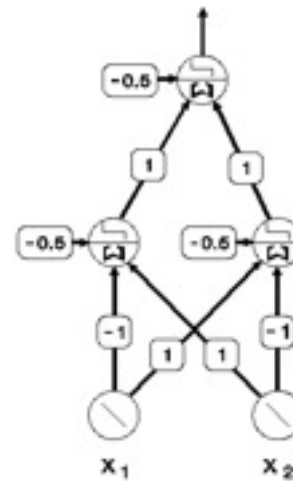
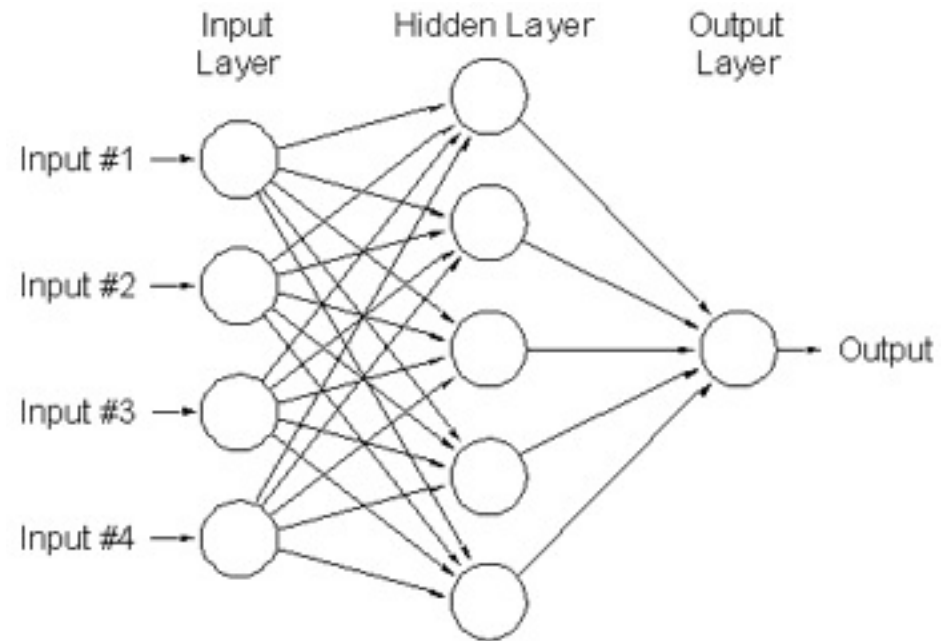
lineáris neuronra, egy pontpárra vonatkozó közelítéssel:

$$\mathbf{w} = \mathbf{w} + \epsilon(v^m - v(\mathbf{u}^m))\mathbf{u}^m$$

- Minsky-Papert 1969: a neurális rendszerek csak lineáris problémákat tudnak megoldani

Multi-Layer Perceptron

- Nemlineáris szeparáció
- regresszió
- egyenletesen sűrű l2-ben egy rejtett réteggel
- A reprezentációs képességet a rejtett réteg növelésével tudjuk növelni
- Idegrendszerben – látórendszer ...



Error backpropagation

- A Delta-szabály alkalmazása
- Első lépés: minden neuron kimenetét meghatározzuk a bemenet alapján $o = \text{sig}(\mathbf{xw})$
- A hibafüggvény parciális deriváltjai:

$$\text{sig}'(y) = \text{sig}(y)(1 - \text{sig}(y))$$

Minden neuronra: $\frac{\partial E}{\partial w_{bi}} = o_i \delta_i$

Kimeneti rétegben: $\delta_k = o_k(1 - o_k)(o_k - t_k)$

Rejtett rétegben: $\delta_j = o_j(1 - o_j) \sum w_{jq} \delta_q$

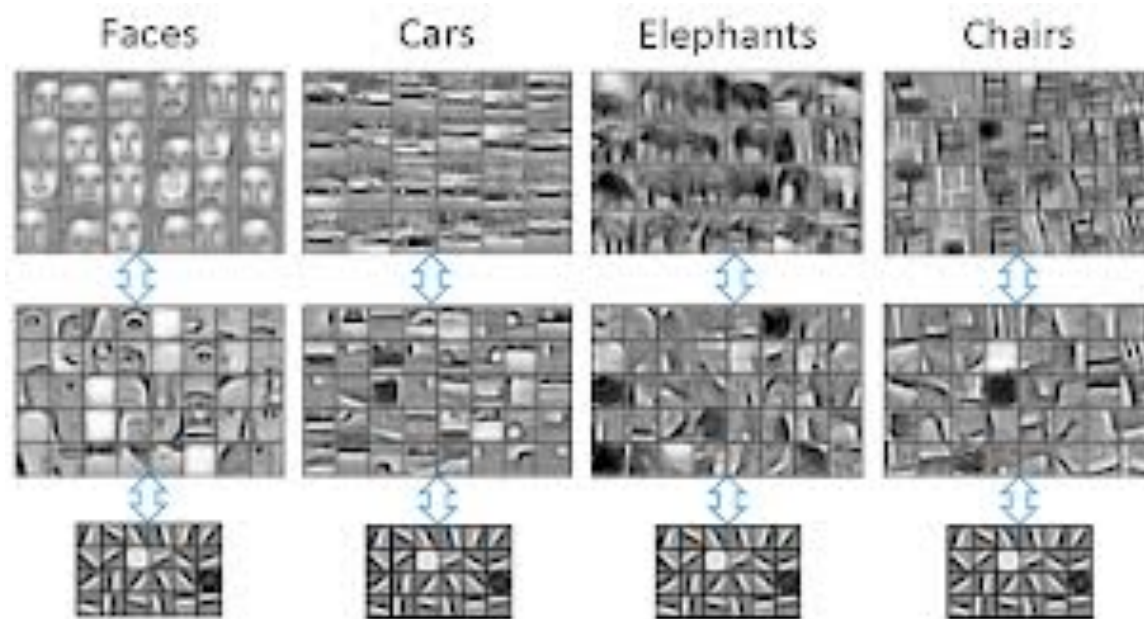
- Mivel gradiens-módszer, a hibafüggvény lokális minimumába fog konvergálni.

Deep vs shallow learning

Egy rejtett réteggel az előrecsatolt hálózatok tetszőleges véges energiájú függvényt tudnak közelíteni a rejtett neuronok számától függő pontossággal

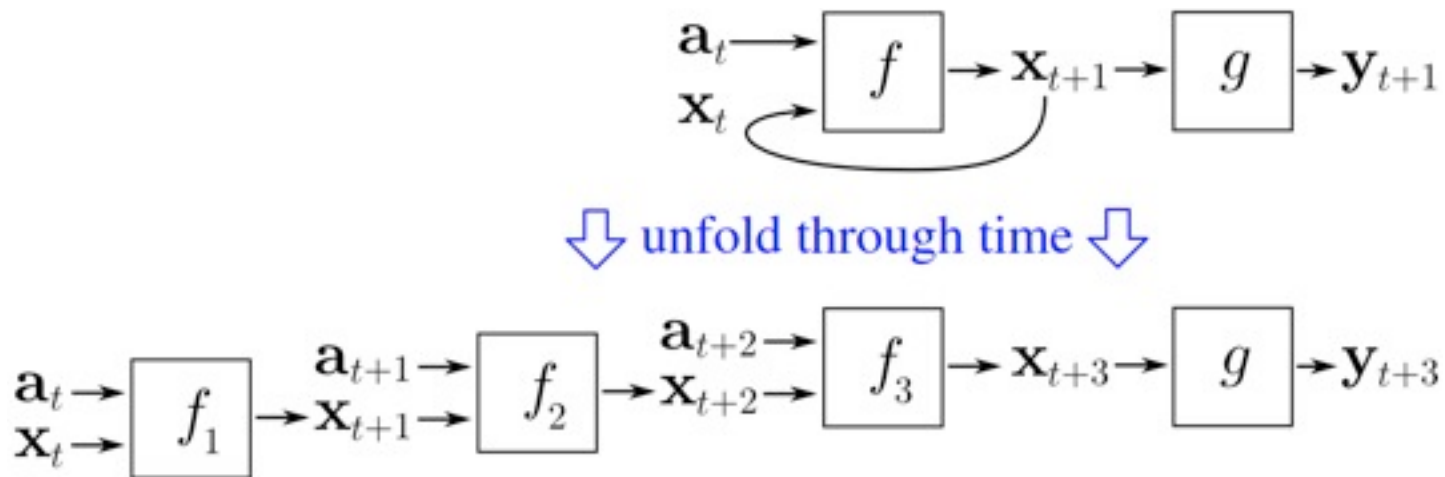
Továbbfejlesztés: Support Vector Machine

Deep learning: egymásra épülő rejtett rétegek hierarchiája



Rekurrens hálózatok

- Megengedjük a rejtett rétegbeli kapcsolatokat
- Egy végtelen rejtett réteggel rendelkező előrecsatolt hálónak felel meg
- Backpropagation through time



Sztochasztikus tanulás

Eddig: a paraméterekre és a rejtett változókra pontbecslést adunk.

Tanulható teljes eloszlás is.

Sztochasztikus neuron:

$$P(o_i = 1) = \text{sig}(\mathbf{x}_i \mathbf{w}_i)$$

A hálózat sejtjeinek aktivitása az általuk reprezentált változók együttes eloszlása szerint változik

Problémák tanulórendszerekben

- Bias-variance dilemma
 - Strukturális hiba: a modell optimális paraméterekkel is eltérhet a közelítő függvénytől (pl lineáris modellt illesztünk köbös adatra)
 - Közelítési hiba: a paraméterek pontos hangolásához végtelen tanítópontra lehet szükség

- Pontosság vs. általánosítás

- A sokparaméteres modellek jól illeszkednek, de rosszul általánosítanak: túlillesztés
- A magyarázó képességük is kisebb (lehet): Ockham borotvája

