



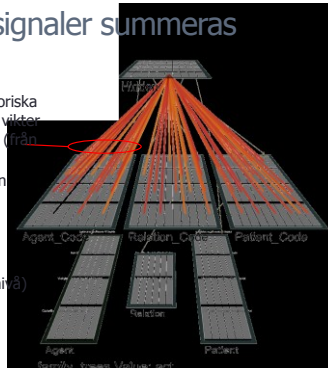
Inlärnin utan övervakning

Översikt

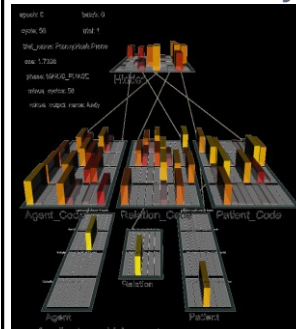
- Biologiska mekanismer bakom inlärnin
- Inlärnin utan övervakning
 - Hebbiansk modellinlärnin
- Självorganisering
 - Arbetsfördelning mellan noder i ett lager som utvecklas utan övervakning

Inkommande signaler summeras

- 'net' för varje nod
 - Summan av alla excitatoriska signaler * motsvarande vikter inom samma projektion (från samma lager)
 - Skalning av projektionen
 - abs och rel wt_scale
 - Summan av alla projektioner
 - Plus nodens bias (individuellt känslighetsnivå)



k aktiva noder i varje lager



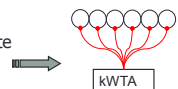
- När alla 'net' beräknats
- Beräkna lämplig inhibition för lagret
- För de noder som hamnat över tröskeln
 - Beräkna deras (ut)aktivering
 - $XX1(V_m - \Theta) \rightarrow act$
 - Representerar spik-frekvens

Inhibering

- Neural aktivering som minskar sannolikheten för att mottagande neuronerna ska avfira
 - Signalsubstansen GABA-A el. GABA-B
- Ungefär 15 % av de kortikala neuronerna är inhibitoriska
 - Verkar (i huvudsak) lokalt, dvs. skickar signal till kringliggande neuron i samma hjärnområde

Inhibitorisk tävlan (kWTA)

- Biologiskt sett:
 - Visst antal inhibitoriska noder
 - Verkar lateralt (åt sidan) inom samma lager
 - Flera uppdateringar innan inhibitionen landar på rätt nivå
- Av beräknings och praktiska skäl:
 - Inhibitoriska noder implementeras inte
 - Istället, k-Winners-Take-All (kWTA)
 - kWTA utför samma funktion som laterala inhibitoriska synapser
 - Inhiberar på rätt nivå för varje cykel



Biologiska mekanismer bakom inlärning

Två typer av inlärning

- Inlärning utan övervakning **Idag**
 - Modellinlärning
 - Skapa en modell av omgivningen
 - Kategorisera/strukturera input på basis av statistiska mönster
 - Baserad på endast input
 - Unsupervised learning, Hebbian learning, competitive learning
- Inlärning med övervakare **Nästa föreläsning**
 - Uppgiftsinlärning
 - Producera *önskad* output givet ett visst input
 - Baserad på input och (med viss tidsförskjutning) korrekt output
 - Supervised learning, error-driven learning

Donald Hebbs postulat (1949)

The Organization of Behavior
A NEUROPSYCHOLOGICAL THEORY
D. O. HEBB
McGill University
1949
New York · JOHN WILEY & SONS, Inc.
London · CHAPMAN & HALL, Limited

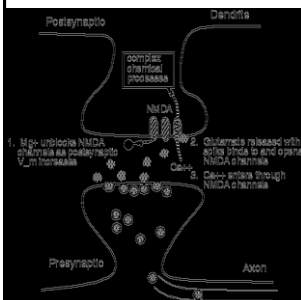
A NEUROPHYSIOLOGICAL POSTULATE
Let us assume then that the persistence or repetition of a reverberatory activity (or "trace") tends to induce lasting cellular changes that add to its stability. The assumption can be precisely stated as follows: *When an axon of cell A is near enough to excite a cell B and repeatedly or persistently takes part in firing it, some growth process or metabolic change takes place in one or both cells such that A's efficiency, as one of the cells firing B, is increased.*
The most obvious and I believe much the most probable suggestion concerning the way in which one cell could become more capable of firing another is that synaptic knobs develop and increase the area of contact between the afferent axon and efferent soma. ("Soma" refers to dendrites and body, or all of

* See p. 229 for a further discussion of this point and an elaboration of the assumption made concerning the nature of memory.

Synapsens effektivitet

- Effektivitet
 - Sändarens förmåga att förändra membran-potentialen (V_m) hos mottagaren
 - Motsv. kopplingens vikt i ett artificiellt nätverk
- Vad som avgör synapsens effektivitet
 - Mängden av signalsubstans som utsöndras från den presynaptiska terminalen (= andel kanaler som öppnas hos mottagaren, g_e)
 - Antalet postsynaptiska kanaler, g_e

Long-Term Potentiation (LTP)



- Ökning av V_m hos mottagaren
- Mg^{2+} öppnar NMDA-kanaler
- Släpper in Ca^{2+} hos mottagaren

Long-Term Depression (LTD)

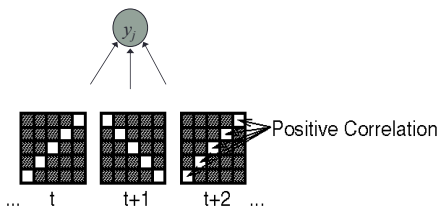
- Long-Term Depression
 - Mekanismerna inte lika väl kända
 - LTD: En svag postsynaptisk aktivitet → mindre effektiv öppning av NMDA kanalerna → lägre koncentration av kalciumjoner Ca^{2+} → komplexa kemiska processer → synapsens effektivitet minskar

Hebbs regel

Mål med modellinläring

- Modellinläring = forma en intern modell av yttre verkligheten på basis av regelbundenheter i input
- Nätet gör detta utan lärare
- Extraherar återkommande mönster i indata
 - Olika delar av input som brukar förekomma tillsammans
 - Särdrag

Korrelation



- Lär sig att pixlar hänger ihop på två sätt i indata

Hebbs regel: LTP

$$\Delta w_{ij} = \varepsilon x_i y_j$$

$$w_{ij}(t+1) = w_{ij}(t) + \Delta w_{ij}$$

ε = inlärningsfaktor (**lrate**)

x_i = presynaptisk nads aktivering

y_j = postsynaptisk nads aktivering

w_{ij} = vikten från nod x_i till nod y_j



Hebbs regel: konsekvenser

$$\Delta w_{ij} = \varepsilon x_i y_j$$

- Viktförändringen:
 - Om $\varepsilon = 0 \rightarrow \Delta w_{ij} = 0$ (ingen viktförändring)
 - Om $\varepsilon = 1 \rightarrow \Delta w_{ij} >> 0$ (stor viktförändring)
 - Vikten kan endast öka
 - Vikten ökar endast om x_i och y_j aktiva samtidigt (> 0)
- Biologiskt plausibel då information som används för att uppdatera vikten är lokal (berör synapsen)

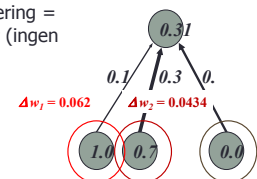
Hebbs regel: exempel

- Anta att mottagarens aktivering = viktad summa av insignaler (ingen bias, skalning, etc.)

$$\begin{aligned}
 &0.1 \times 1.0 + \\
 &0.3 \times 0.7 + \\
 &0.2 \times 0.0 = 0.31
 \end{aligned}$$

- Resp. vikt uppdateras:

$$\begin{aligned}
 \Delta w_1 &= \varepsilon \times x_i \times y_j \\
 \Delta w_1 &= 0.2 \times 1.0 \times 0.31 = 0.062 \\
 \Delta w_2 &= 0.2 \times 0.7 \times 0.31 = 0.0434 \\
 \Delta w_3 &= 0.2 \times 0.0 \times 0.31 = 0
 \end{aligned}$$



Principal Component Analysis (PCA)

- Vad lär sig ett nät vars vikter uppdateras med Hebbs regel?
 - Att detektera mest frekventa sambandet (korrelationen) mellan varje bit i indata
 - T.ex. vilken konstellation en pixel *oftast* deltar i
 - Den **principiella komponenten** bland alla korrelationer
- Notera att det är $x_1, \dots, x_n$ via sina resp. vikter som tillsammans bestämmer y 's aktivering
 - Korrelationen mellan $x_1, \dots, x_n$ bestäms indirekt via y

Problem att fixa

- Hebbs regel kan endast öka vikterna
- Vad göra när flera mottagnoder?
 - Ska alla noder lära sig samma korrelation, dvs. samma särdrag?

Conditional PCA (CPCA)

$$\Delta w_{ij} = \varepsilon (y_j x_i - y_j w_{ij})$$

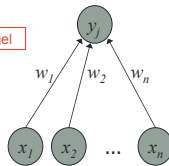
Skilnad mot Hebbs regel

ε = inlärningsfaktor (**rate**)

x_i = sändarnodens aktivering

y_j = mottagnodens aktivering

w_{ij} = vikten mellan noderna x_i och y_j



CPCA

- "Conditional" (villkorlig) innebär här att inkommande vikter till y ska återspegla sannolikheten för att nodens aktivering orsakades av x_i
 - Individuell vikt w_i återspeglar sannolikhet att x_i aktiv när y aktiv
- Hebbs ursprungliga regel tar inte hänsyn till detta, då den saknar en viktnedskärningskomponent (LTD)

CPCA

$$\Delta w_{ij} = \varepsilon (y_j x_i - y_j w_{ij})$$

1. y_j och x_i båda aktiva ($x_i > w_{ij}$) $\rightarrow$ viktökning
 2. y_j aktiv, x_i inte aktiv ($x_i < w_{ij}$) $\rightarrow$ **viktnedskärning**
 3. y_j inte aktiv $\rightarrow$ ingen viktförändring
- Inlärningen avstannar när vikten avspeglar sändarnodens förväntade aktivering i de fall då mottagnoden är aktiv, $w_{ij} = P(x_i | y_j)$

Problem

1. Vikterna tenderar att bli låga
 - Korrelation inom *delmängd* av datamängden
 - Om särdrag utgör 30% av de särdrag som noden aktiveras av skulle den ha 0,3 i maxvikt
 - Sannolikheten att just input x är aktiv när y är aktiv blir lägre ju fler inputmönster y tar hand om
2. Mottagnoderna inte tillräckligt selektiva
 - Startar på slumpmässigt värde
 - Noderna aktiva för många olika indatamönster

CPCA utökning

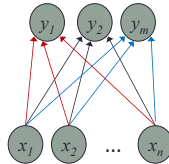
- Normalt hålls vikterna inom $[0, 1[$
 - Nu fylls inte det utrymmet längre
- **Åternormalisering**, förhindrar problem 1
 - Normalisera vikterna
 - Töjning av vikterna till intervallet $[0, 1[$
 - Proportionerlig förstoring av alla vikter
- **Kontrastförstärkning**, förhindrar problem 2
 - Applicera en sigmoidal förstärkarfunktion på vikterna för att få en större kontrast mellan svaga och starka vikter
 - Antagandet är att svaga vikter avspeglar fall då mottagarnoden blev oavsiktligt aktiverad

Flera mottagarnoder

Självorganisering

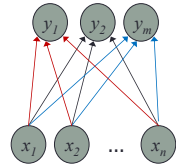
Vill ha specialiserade mottagarnoder

- Vill att olika noder ska lära sig olika särdrag
 - Vill/kan inte tala om vilka som ska lära sig vilka särdrag



Självorganiserande nät

- Starta med små slumpvikter
 - Olika noder har då olika chans att bli aktiverade av samma input
- Noderna får tävla (kWTA)
 - Vinnarna förblir aktiva och får uppdatera sina vikter
 - Självförstärkande effekt: efter viktuppdateringen har vinnarna ännu lättare att vinna nästa gång denna input visas



Inhibitorisk tävlan (kWTA)

- kWTA låter de k mest aktiva noderna överleva
 - k = antal noder el. en procentsats
- Fördelen med kWTA och avg-kWTA:
 - Möjliggör mjuk tävling, då flera noder kan vinna
- Resulterar i att:
 1. Även noder med svag aktivitet tillåts lära sig lite
 - Dvs. flera noder kommer att utveckla känslighet för viss input
 2. Distribuerade representationer erhålls i resp. lager
 - Dvs. flera noder samverkar till att representera ett visst särdrag i input

Självorganisering genom CPCA + kWTA

- CPCA + kWTA $\rightarrow$ självgruppering av noder som reagerar på samma särdrag, dvs. samma kategori av input
- Antalet kategorier styrs genom att
 - Öka antalet inputnoder (öka detaljnivån i input)
 - Litet antal pixlar $\rightarrow$ 2 kat: runda el. avlånga
 - Stort antal pixlar $\rightarrow$ 10 kat: olika siffror
 - Ändra antalet dolda noder + olika värden på k i kWTA
 - Fler eller färre aktiveringsmönster kan bildas

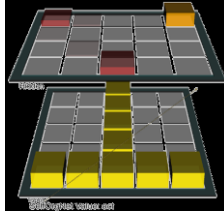
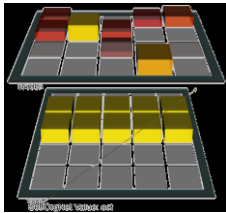
Lagrets storlek och kWTA...

$k = 9 \rightarrow 2042975$ repr.

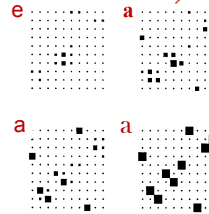
Fler representationer

$k = 2 \rightarrow 25 \cdot 24 / 1 \cdot 2 = 300$ repr.

Mindre risk för överlapp



Distribuerad representation



Har aldrig sett input,
men kommer att känna
igen den som a

- Aktiveringsmönster över många enheter motsvarar
 - objekt
 - begrepp, etc.
- Ger robusthet
- Ger generaliserbarhet

Överlapp

- Ibland vill man inte ha överlapp
 - Överlapp kan ge **interferens**
 - T.ex. vill vi kunna skilja minnet av olika julfiranden genom åren
- Ibland är överlapp bra
 - Input som nätet aldrig sett tidigare, men som liknar inlärdd input kan delvis aktivera inlärdd output
 - Nätet kan avge respons som är någorlunda adekvat
 - Förmågan att ge adekvat respons för ej tränad input kallas **generalisering**