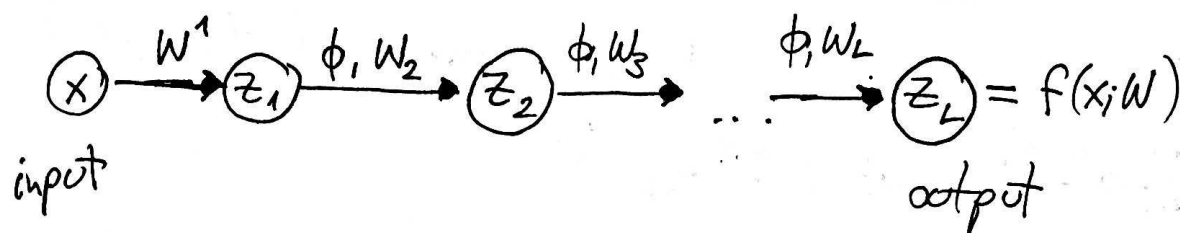


Machine learning

Title: Learning and Memory in the Brain
(Les Gauthier 2025)

• Feedforward network with $L-1$ hidden layers $z_1, \dots, z_{L-1}$

Output $f(x; W) = z^L(x)$, $z^L = W^L \phi(z^{L-1}) \in \mathbb{R}^{N_L}$
 $z^1 = W^1 x$ $\nearrow$ nonlinearity
 input $x \in \mathbb{R}^{N_0}$ $\nearrow$ Weights $\{W^1, \dots, W^L\}$
 $W^l \in \mathbb{R}^{N_l \times N_{l-1}}$



• training

~~Objective~~ Objective function $C(y, f(x; W))$: How much does the network output $f(x; W)$ for given input x differ from desired output y .

E.g. $C(y, f(x; W)) = \|y - f(x; W)\|^2$ (MSE)

~~Gradient descent~~

Given training data $\{x^\mu, y^\mu\}_\mu$, loss $\mathcal{L}(W) := \sum_\mu C(y^\mu, f(x^\mu; W))$

Gradient descent

$W^l \leftarrow W^l - \eta \frac{\partial \mathcal{L}(W)}{\partial W^l}$
 $\uparrow$
 learning rate

•/ Backpropagation of the error : efficient recursive algo. to compute $\frac{\partial C}{\partial w^l}$ for all $l=1, \dots, L$

Needed : $\frac{\partial C}{\partial w_{ij}^l} = \frac{\partial C}{\partial z_{ij}^l} \frac{\partial z_{ij}^l}{\partial w_{ij}^l}$ (sum over repeated indices $i, j, l, \dots$)

$$= \frac{\partial C}{\partial z_{ij}^l} w_{i_{l-1} i}^l \phi'(z_{i_{l-1} i}^{l-1}) \dots w_{i_{l+1} i}^{l+1} \phi'(z_{i_{l+1} i}^l) \phi(z_{ij}^{l-1})$$

recursion : $\equiv \delta_i^l$ (error signal at layer l)

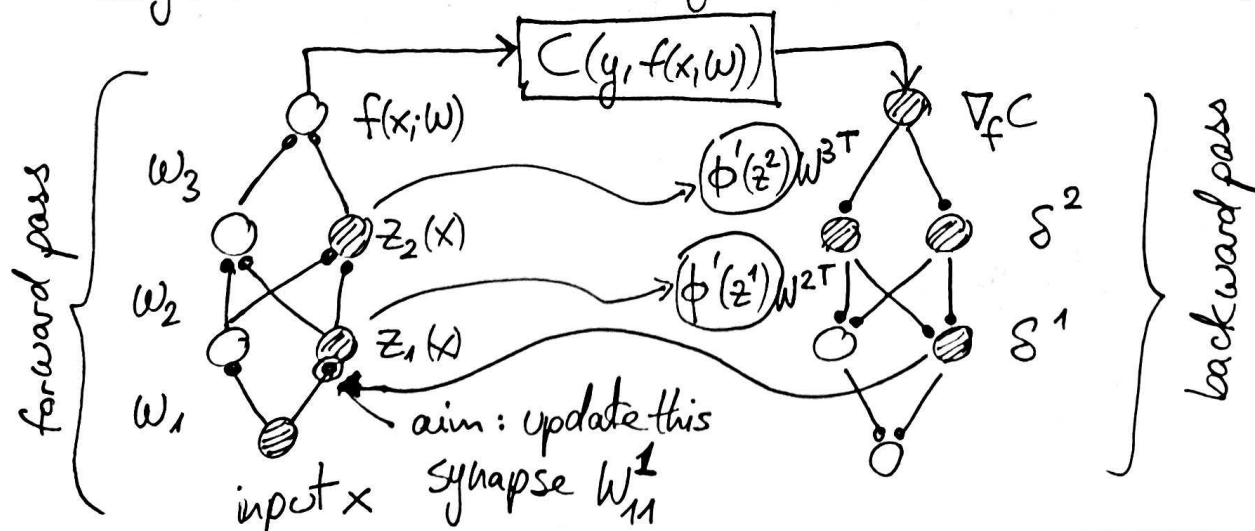
$$\delta_i^{l-1} = \sum_k \delta_k^l w_{ki}^l \phi'(z_i^{l-1}) \text{ with } \delta^l \in \mathbb{R}^{N_l}$$

or $\delta^{l-1} = \text{diag}[\phi'(z_i^{l-1})] (W^l)^T \delta^l$ with $\delta^L = \nabla_f C$

$$\frac{\partial C}{\partial w^l} = \delta^l \phi(z^{l-1})^T \text{ (and } \frac{\partial C}{\partial w^1} = \delta^1 x^T \text{)}$$

•/ In the Brain: How could a synapse w_{ij}^l receive the error signal δ_i^l which determines how it should be updated to optimize the objective function $C(y, f(x, w))$?

E.g. x : visual stimulus, y : desired motor output (e.g. move your arm)



Several issues:

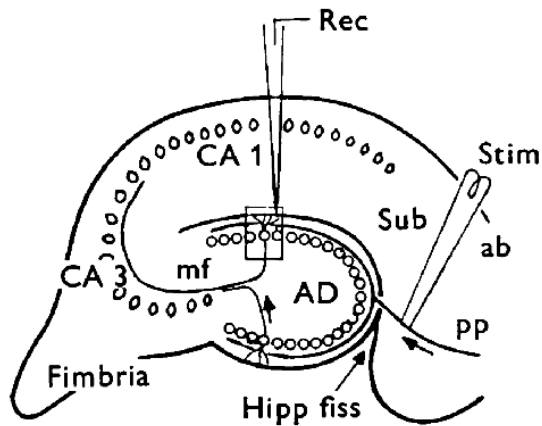
- no nonlinearity in the backward pass:

$$z^l = w^l \phi(z^{l-1}) \quad (\text{forward})$$

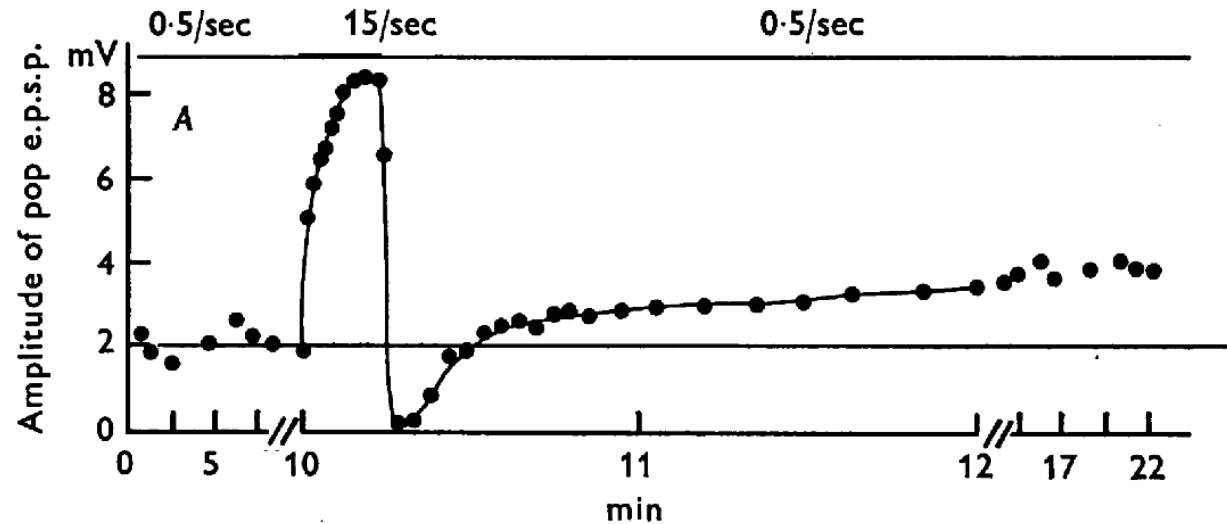
$$\text{but } \delta^{l-1} = \phi'(z^{l-1}) w^{lT} \delta^l \quad (\text{backward})$$

- Either forward and backward pass happen in the same network, then synapses need to be symmetric $w_{ij}^l = w_{ji}^l$, but:
 - that's usually not the case and
 - the network would need to precisely time and alternate forward and backward passes
- Or the backward pass happens in a separate network with exactly the same (but reversed) synapses. But then it would also need knowledge of $\phi'(z^{l-1})$ and still time the alternating forward/backward passes
- (Not so clear what exactly the output in the brain would ~~be~~ ^{be})

Long Term Potentiation (LTP)



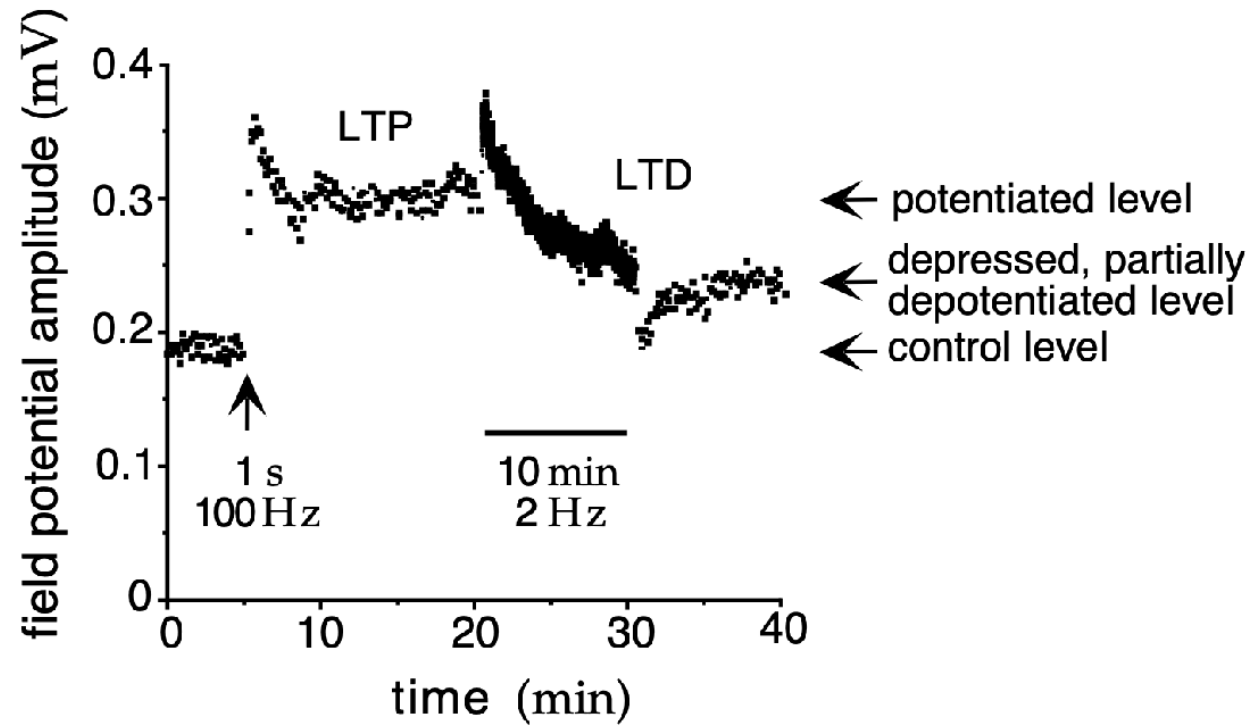
Perforant path fibres (pp) in hippocampus of rabbits anaesthetized with urethane.



Increasing the stimulation frequency from 5/sec to 15/sec during a short period (~2min) increases the neuron's response for a long period

Bliss and Lomo (1973)

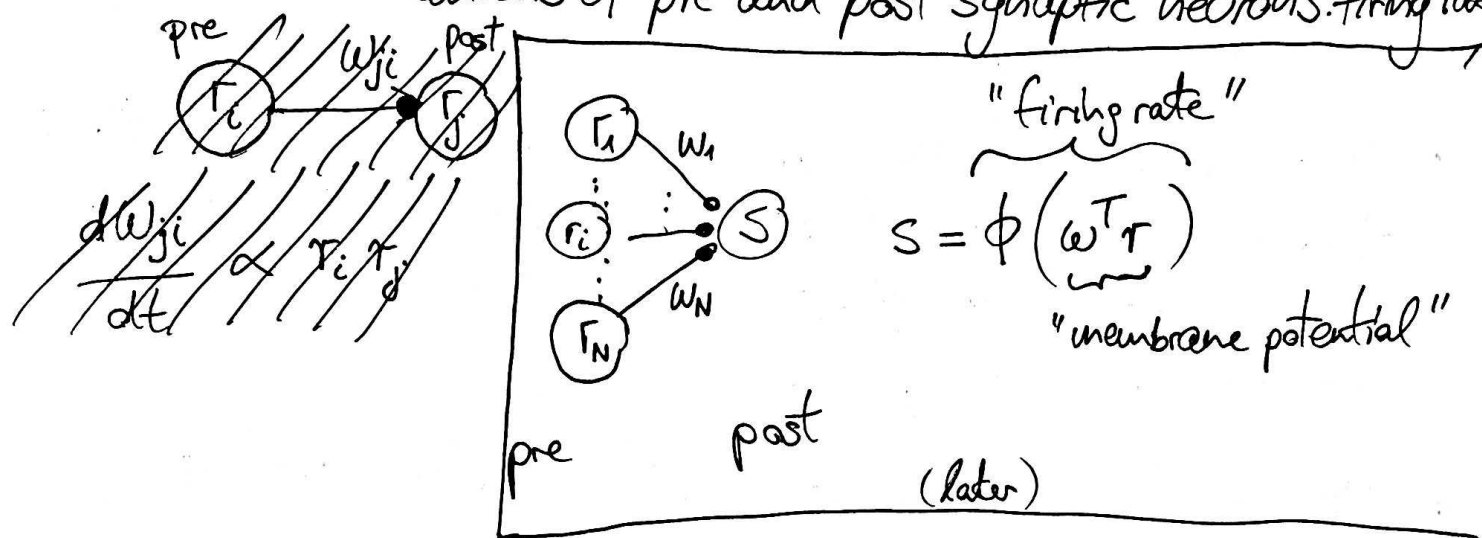
Long Term Depression



③ Neuroscience

- Hebb (1949) conjectures that if input from neuron A often contributes to the firing of neuron B, then the synapse $A \rightarrow B$ is strengthened $\textcircled{A} \rightarrow \textcircled{B}$
- Experimental evidence: Bliss & Lomo (1973) and many studies afterwards

• Theoretical models: Hebbian learning/plasticity based on correlations of pre and post synaptic neurons. firing rates



$$\frac{d w_i}{dt} \propto r_i S \quad (\text{basic hebbian learning})$$

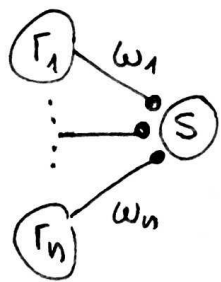
• Problems

- Weights grow unconstrained ~~$\frac{d w}{dt} = \frac{(r^T S + \bar{w})}{\|r\|^2 + \bar{w}}$~~ (normalization)
- Symmetric weights per default $\frac{d w_{ij}}{dt} \propto r_i r_j \rightarrow f(r_i) g(r_j)$

What can this do?

- Unsupervised: How to optimize an objective function without access to any error signal?
- Can this do interesting things? (Yes $\rightarrow$ Hopfield) ④

- Oja's modification of Hebb's rule: Normalize weights + small updates γ



$$w(t+1) = \frac{w(t) + \gamma r(t) s(t)}{\|w(t) + \gamma r(t) s(t)\|} \approx w(t) + \gamma (\underbrace{sr}_{\text{"learning"}} - \underbrace{ws^2}_{\text{"forgetting"}})$$

here: lin. approx

$$\text{where } s(t) = \phi(r^T w) \approx r^T w$$

↑ transfer function: Membrane pot $\rightarrow$ Firing rate

Take away: Weights w converge to be the eigenvector of

$$C := \langle r r^T \rangle_t \quad (\text{correlation matrix of the input to neuron } s)$$

Covariance

with largest eigenvalue, i.e. w becomes the 1st PC.
(principle component).

Proof:

Note that $\|w(t)\| = 1 \quad \forall t$ (Def. of the update rule)

$$w(t+1) \approx w(t) + \gamma \left. \frac{\partial w(t+1)}{\partial \gamma} \right|_0$$

$$\begin{aligned} \left. \frac{\partial w(t+1)}{\partial \gamma} \right|_0 &= sr - \frac{w}{\|w\|} \left. \frac{\partial \|w + \gamma rs\|}{\partial \gamma} \right|_0 \\ &= sr - w \frac{(w + \gamma rs)^T rs}{\|w + \gamma rs\|} \Big|_{\gamma=0} = sr - w \underbrace{w^T r s}_s \end{aligned}$$

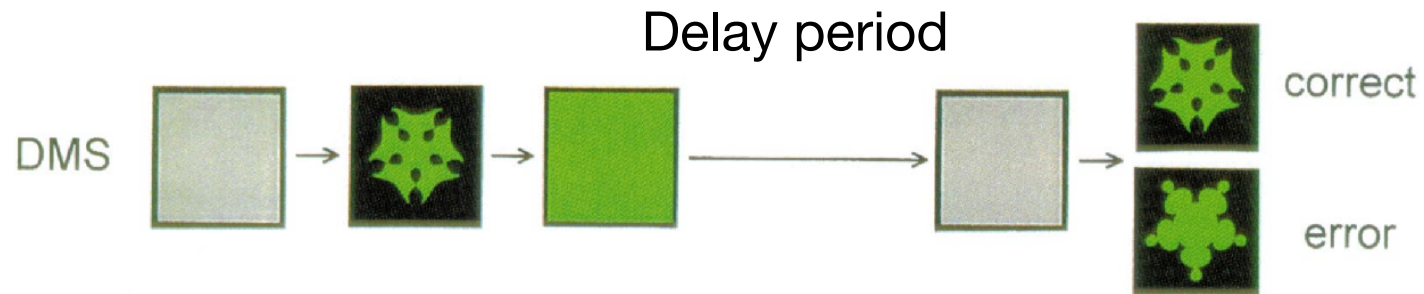
$$\Delta w(t) = \gamma (sr - ws^2) \quad \leftarrow \text{substitute } s = r^T w$$

$$\langle \Delta w(t) \rangle_t = 0 = \gamma (\langle r r^T \rangle_t w - (w^T \langle r r^T \rangle_t w) w)$$

$$\Leftrightarrow Cw = (w^T C w) w$$

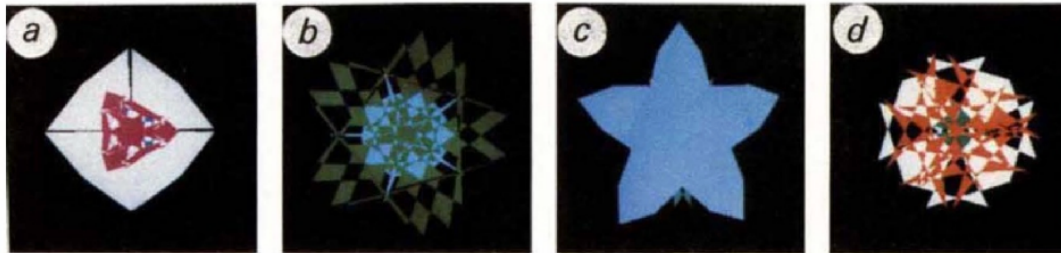
w is evec. of C with eval $w^T C w$. One can show that the dynamic converges to w st. $w^T C w$ is maximal $\square$

Delay Match to Sample Tasks

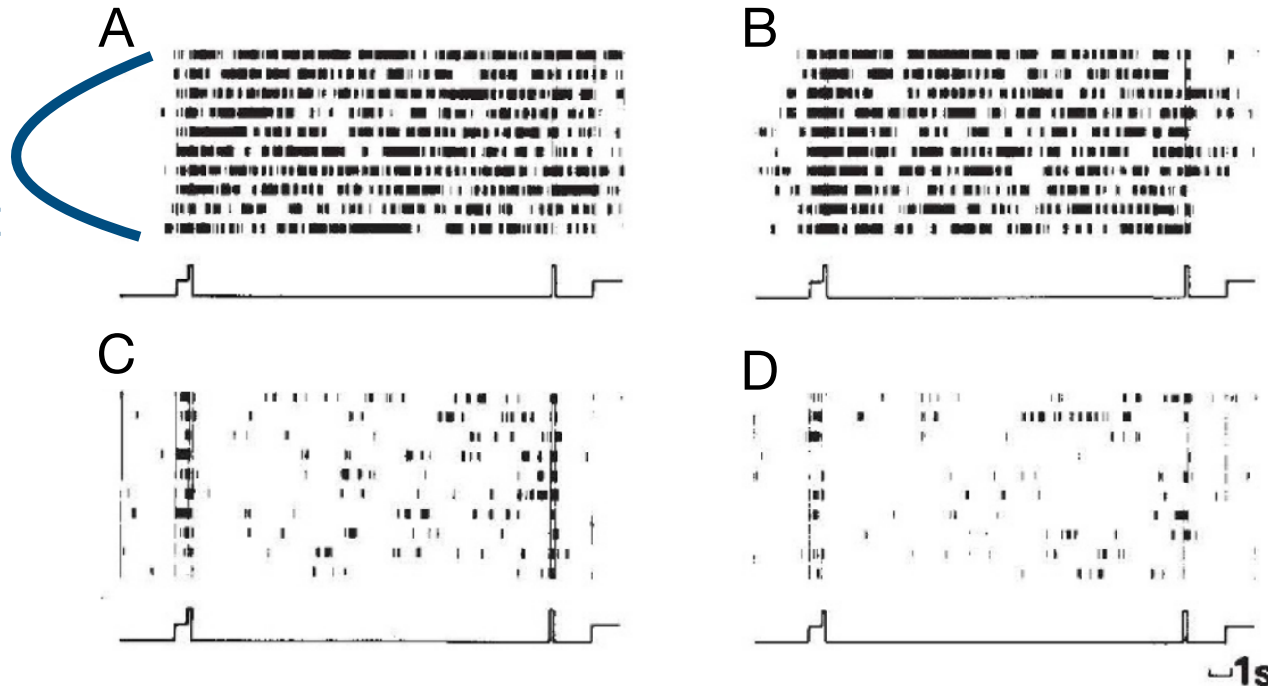


Adapted from Y Naya, K Sakai, Y Miyashita (1996)

Y. Miyashita, H.S. Chang,
 “Neuronal correlate of
 pictorial short-term
 memory in the primate
 temporal cortex”
 1988

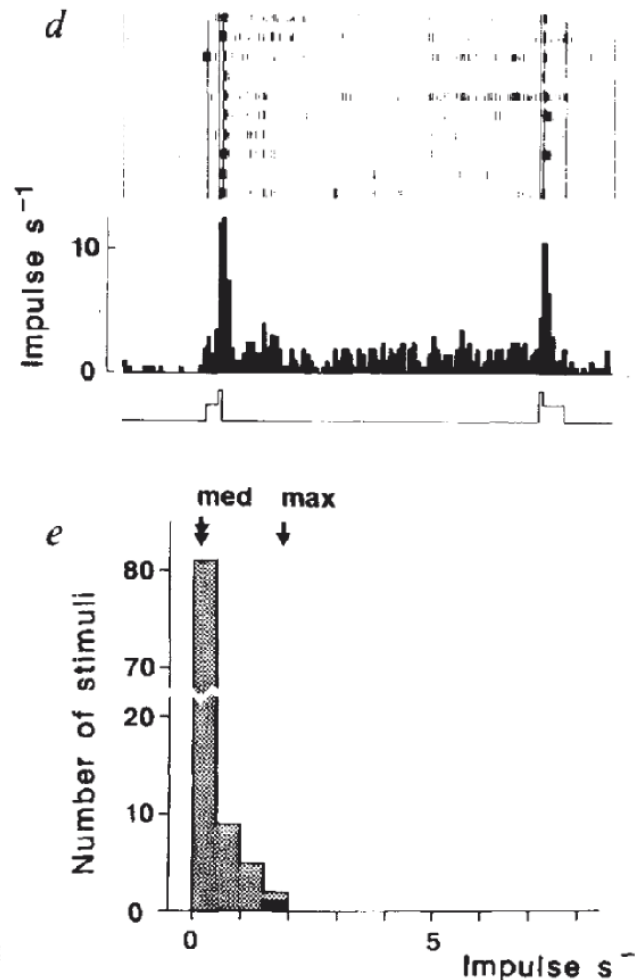
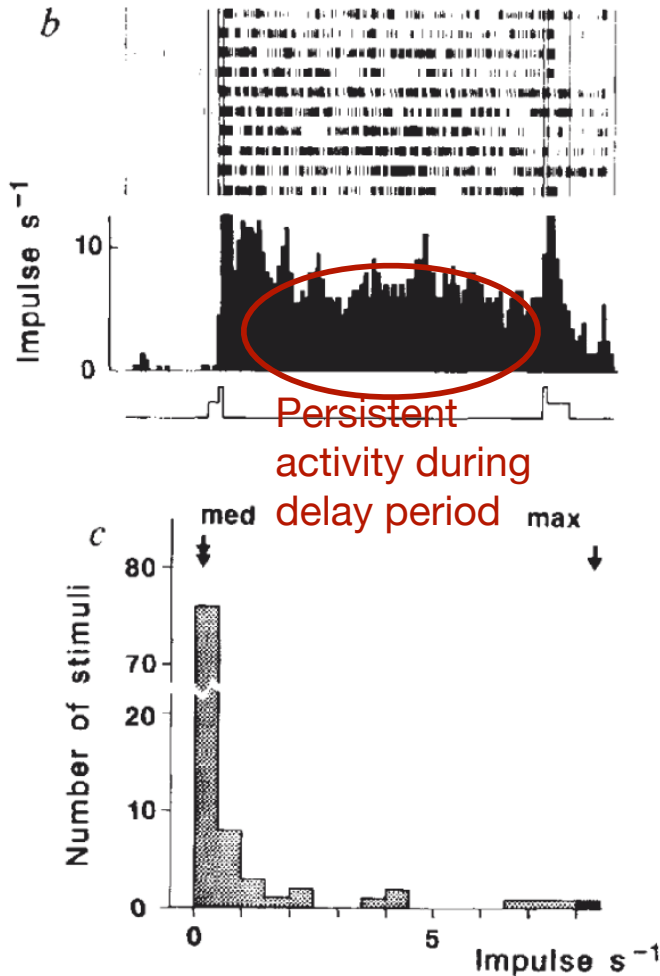


Same
 neuron,
 different
 trials



Learned image

New image



Y Miyashita, “Neuronal correlate of visual associative long-term memory in the primate temporal cortex” (1988)

(b,d) Trial averaged activity of the neuron that is most active for the given image

(c,d) Distribution of firing rates (Impulse/s) during delay period for a given neuron and 100 images

• Hopfield network (1982)

- Binary neurons $S_i(t) = \begin{cases} 1, & \text{active} \\ -1, & \text{inactive} \end{cases} \quad i=1, \dots, N$

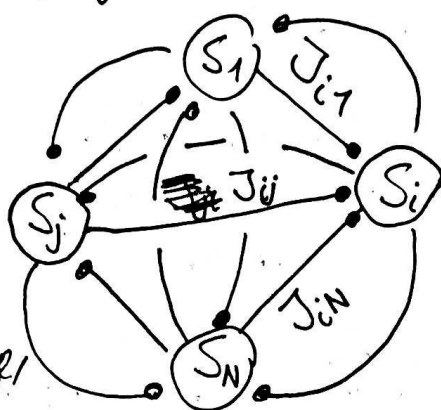
Connected with synaptic weights J_{ij} .

Evolving (synchronously or asynchr.)

$$S_i(t+1) = \text{sgn} \left(\sum_j J_{ij} S_j(t) \right)$$

- Biological interpretation

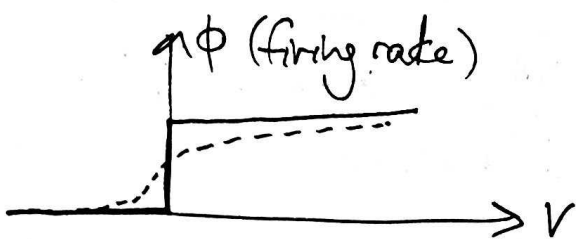
S_i : firing rate of neuron i



fully connected recurrent network

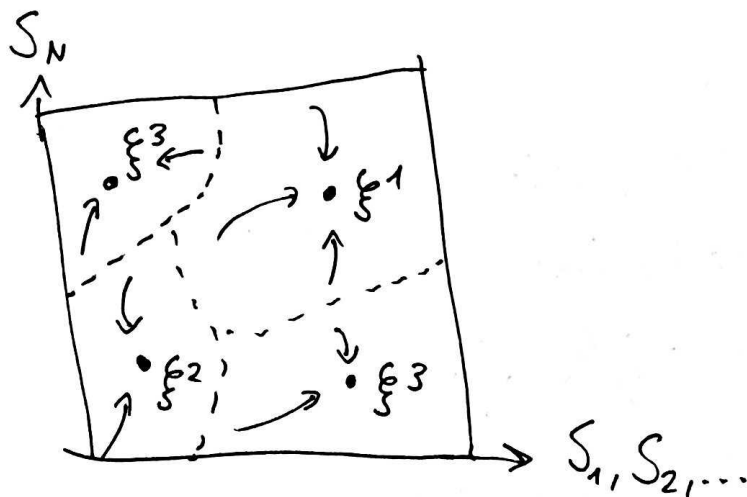
$h_i = \sum_j J_{ij} S_j$: membrane voltage, potential of neuron i

$\phi = \text{sgn}$: transfer function (potential $\rightarrow$ firing rate)



- How to store memories (here patterns of activity $\{S_i^\mu = \pm 1\}$ in the synapses J_{ij} of this network ~~that~~ ?
 ~~can be recalled in the form of persistent activity~~
 pattern index μ
 neuron i

Idea: Memories correspond to fixed points of the dynamics and can be recalled by the dynamics as soon as the network state is close enough / in the basin of attraction of the fixed point.



- Correct J corresponds to learning the patterns with Hebb's

A possible rule one by one: $J^{(0)} = 0$

$$J_{ij}^{(1)} = \frac{1}{N} \xi_i^{(1)} \xi_j^{(1)}$$

$\uparrow$ post syn. $\uparrow$ pre syn. firing rate

(the pattern to be learned is presented to the neurons of the network, i.e. they are forced to the corresp. values and the connections are updated)

$$J_{ij}^{\text{final}} = \frac{1}{N} \sum_{\mu} \xi_i^{\mu} \xi_j^{\mu} \text{ and } J_{ii} = 0$$

- Dynamics (show that it works)

Can be understood as minimizing an energy

$$E = -\frac{1}{2} \sum_{i \neq j} J_{ij} S_i S_j \stackrel{J_{ii}=0}{=} -\left(\sum_j J_{kj} S_j\right) S_k - \frac{1}{2} \sum_{(i+j) \neq k} J_{ij} S_i S_j$$

This always decreases: If neuron k is updated

$$S_k(t+1) = \text{sgn}\left(\sum_j J_{kj} S_j(t)\right) = \text{sgn}(h_k(t))$$

$$E(t) - E(t+1) = -\left(\sum_j J_{kj} S_j\right) (S_k(t) - S_k(t+1))$$

If S_k does not flip ($S_k(t+1) = S_k(t)$): $\Delta E = 0$

Otherwise: $\Delta E = E(t+1) - E(t) = -2 S_k(t+1) h_k(t) < 0$

- Patterns (if not too many) are stable: $\uparrow$ same sgn as h_k

$$\text{If } S_i(t) = \xi_i^1 \text{ then } S_i(t+1) = \text{sgn}\left(\frac{1}{N} \sum_{\mu} \sum_j \xi_i^{\mu} \xi_j^{\mu} \xi_j^1\right) \approx \text{sgn}(\xi_i^1) = \xi_i^1$$

- max number of patterns P that can be stored and perfectly retrieved with N neurons

$$\alpha := \frac{P}{N} = \frac{1}{2 \ln N}$$

Why? Study SNR

Assume $S_i(t) = \xi_i^1$. Then $S_i(t+1) = \text{sgn}(h_i(t))$

$$h_i(t) = \sum_{j \neq i} J_{ij} \xi_j^1 = \underbrace{\xi_i^1}_{\text{Signal}} + \underbrace{\frac{1}{N} \sum_{j \neq i} \sum_{\mu=2}^P \xi_i^\mu \xi_j^\mu \xi_j^1}_{= \delta_i \text{ noise}}$$

avg. over $\xi_i^\mu = \pm 1$ w.p. $\frac{1}{2}$

$$\langle \delta_i \rangle = 0$$

$$\langle \delta_i^2 \rangle \approx \frac{P}{N} = \alpha \rightarrow \text{SNR} = \frac{1}{\sqrt{\alpha}}$$

$$g \equiv \text{Prob}(\text{Noise } \delta_i < \text{signal} = 1) = \int_{-\sqrt{1/\alpha}}^{\infty} \frac{dx}{\sqrt{2\pi}} e^{-x^2/2} \quad (\delta_i \text{ is approx Gaussian with std } \sqrt{\alpha})$$

$$\left(\begin{aligned} &\text{(write } \delta_i = \sqrt{\alpha} X \text{ with } X \sim \mathcal{N}(0,1), \text{ then } \Pr(\delta < 1) = \Pr(X < \frac{1}{\sqrt{\alpha}}) \\ &= \Pr(X > -\frac{1}{\sqrt{\alpha}}) \end{aligned} \right)$$

$$g \approx 1 - \exp(-\frac{1}{2\alpha}) (\sqrt{\alpha} + O(\alpha^{3/2}))$$

$$\text{Prob}(\delta_i < 1 \forall i) = g^N \sim \exp(-N \sqrt{\alpha} \exp(-\frac{1}{2\alpha}))$$

$$\text{If } \alpha < \frac{1}{2 \ln N} : g^N \rightarrow 1$$

$$\text{If } \alpha > \frac{1}{2 \ln N} : g^N \rightarrow 0$$

- if not perfect retrieval is required: $\alpha \sim O(1)$

- Critiques of Hopfield's model

- No separation between excitatory and inhibitory neurons
(Dale's law: All synapses of a given neuron are either inhibitory or excitatory)
- Electrophys. recordings show rather a sparse coding of memories
(most $S_i^{\mu} = 0$)
- Symmetric connections are not biologically plausible
(Sompolinsky showed that $J_{ij} + \eta_{ij}$ is robust where $\eta_{ij} \sim N(0, \Delta)$ with $\Delta \sim O(1)$)
- Black-out catastrophe: All patterns are lost if one goes over the capacity. Numerous modifications have been proposed
 - Either: stop learning by some supervised signal when capacity is reached
 - Or: modify st. continual learning is possible where older patterns get forgotten as newer ones are learned.

$$\vec{w}(t) + \eta(\vec{r}(t) - \vec{w}(t)) \approx \vec{w}(t) + \eta(\vec{r}(t) - \vec{w}(t))$$

$$\vec{w}(t+1) = \frac{\vec{w}(t) + \eta \vec{r}(t)}{\|\vec{w}(t) + \eta \vec{r}(t)\|} \leftarrow \Delta \vec{w} = \eta \vec{r}(t)$$

• Normalized Hebbian rule