

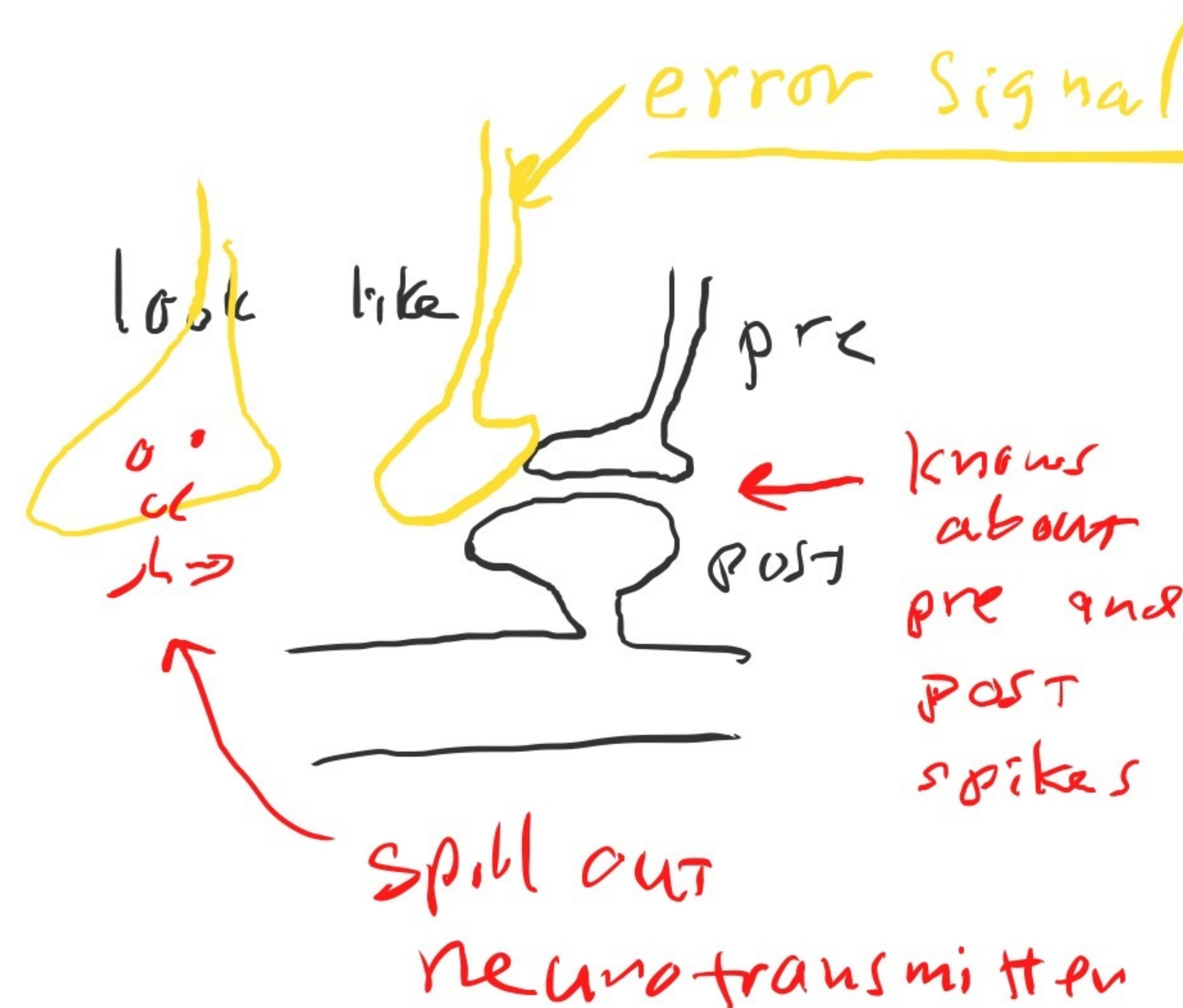
$$\dot{V}_i = \phi \left(\sum_j W_{ij} V_j \right) - V_i$$

biologically plausible learning rules

$$\Delta W_{ij} = f(V_i, V_j)$$

$\nwarrow$ pre $\swarrow$ post

is this a powerful enough learning rule to do anything useful?



more correct:

$$y = \phi(\underline{w} \cdot \underline{x})$$

$$y = \underline{w} \cdot \underline{x}$$

$$\Delta \underline{w} = \eta y (\underline{x} - y \underline{w})$$

→ Oja's rule

$$= \eta \underline{w} \cdot \underline{x} (\underline{x} - \underline{w} \cdot \underline{x} \underline{w}) = \eta \underline{w} \cdot \underline{x} \underline{x} - \eta \underline{w} \underline{w} \cdot \underline{x} \underline{x} \cdot \underline{w}$$

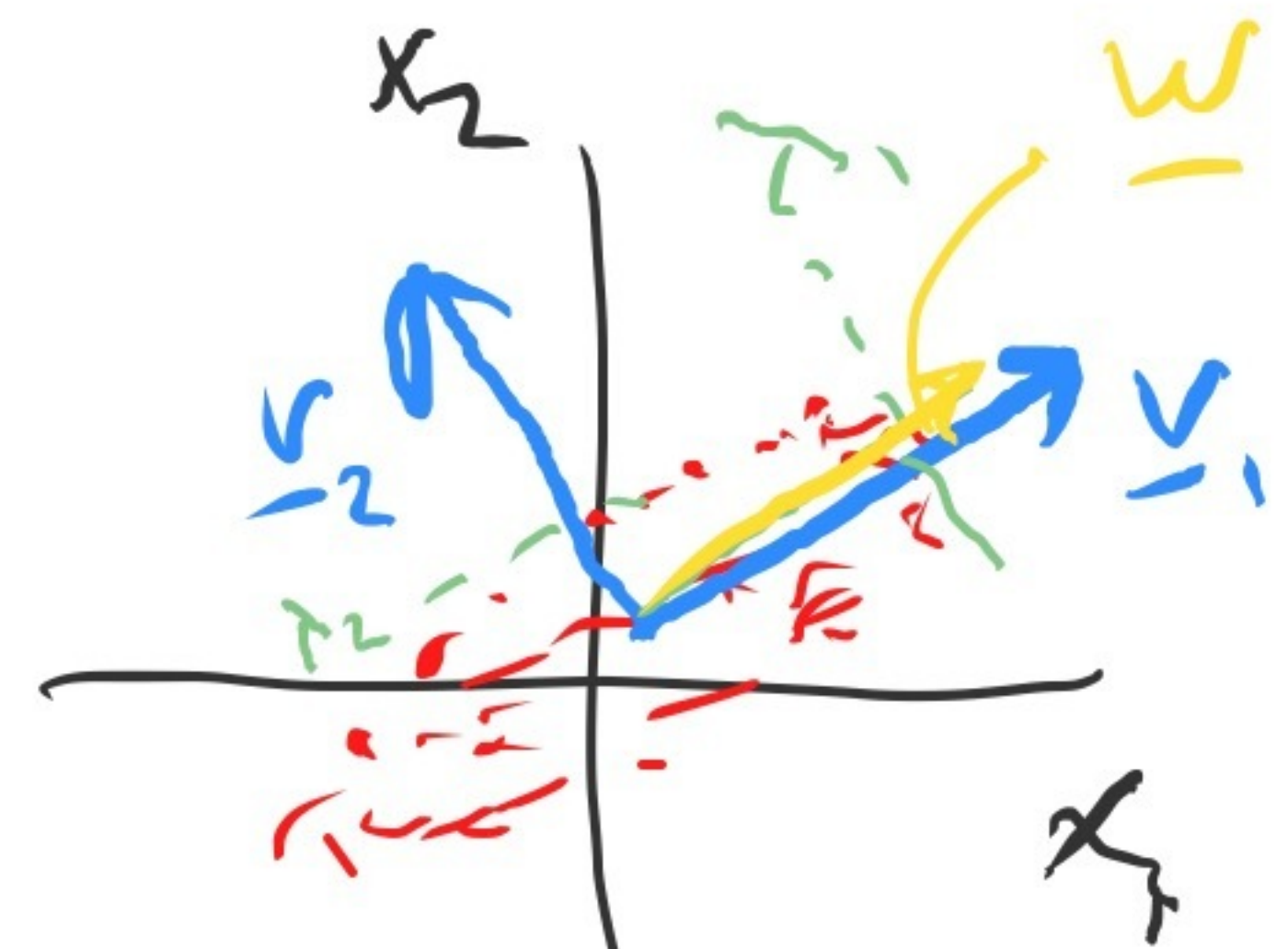
$$\eta \dot{\underline{w}} = \eta \left[\underline{w} \cdot \langle \underline{x} \underline{x} \rangle - \underline{w} \underline{w} \cdot \langle \underline{x} \underline{x} \rangle \cdot \underline{w} \right]$$

$$T \dot{\underline{w}} = \gamma \left[\underline{C} \cdot \underline{w} - \underline{w} \cdot \underline{C} \cdot \underline{w} \underline{w} \right]$$

$$\underline{w} = \sum_k a_k \underline{v}_k$$

$$\underline{C} \cdot \underline{v}_k = \lambda_k \underline{v}_k$$

$$\underline{v}_k \cdot \underline{v}_l = \delta_{kl}$$



$$T \sum_k \dot{a}_k \underline{v}_k = \gamma \left[\sum_k a_k \lambda_k \underline{v}_k - \left[\sum_l a_l^2 \lambda_l \right] \sum_k a_k \underline{v}_k \right]$$

$$T \dot{a}_k = \gamma \left[\lambda_k - \sum_l a_l^2 \lambda_l \right] a_k$$

Assume all eigenvalues are different

$$\lambda_k = \sum_l a_l^2 \lambda_l \quad \forall k$$

$$a_k \geq 1$$

$$a_{l \neq k} = 0$$

typically,
largest
eigenvalue
wins

$$\sum_{l,m} a_l \underline{v}_l \cdot \underline{C} \cdot a_m \underline{v}_m$$

$$= \sum_{l,m} a_l a_m \overbrace{\underline{v}_l \cdot \underline{v}_m}^{\delta_{lm}} \lambda_m$$

$$= \sum_l a_l^2 \lambda_l$$

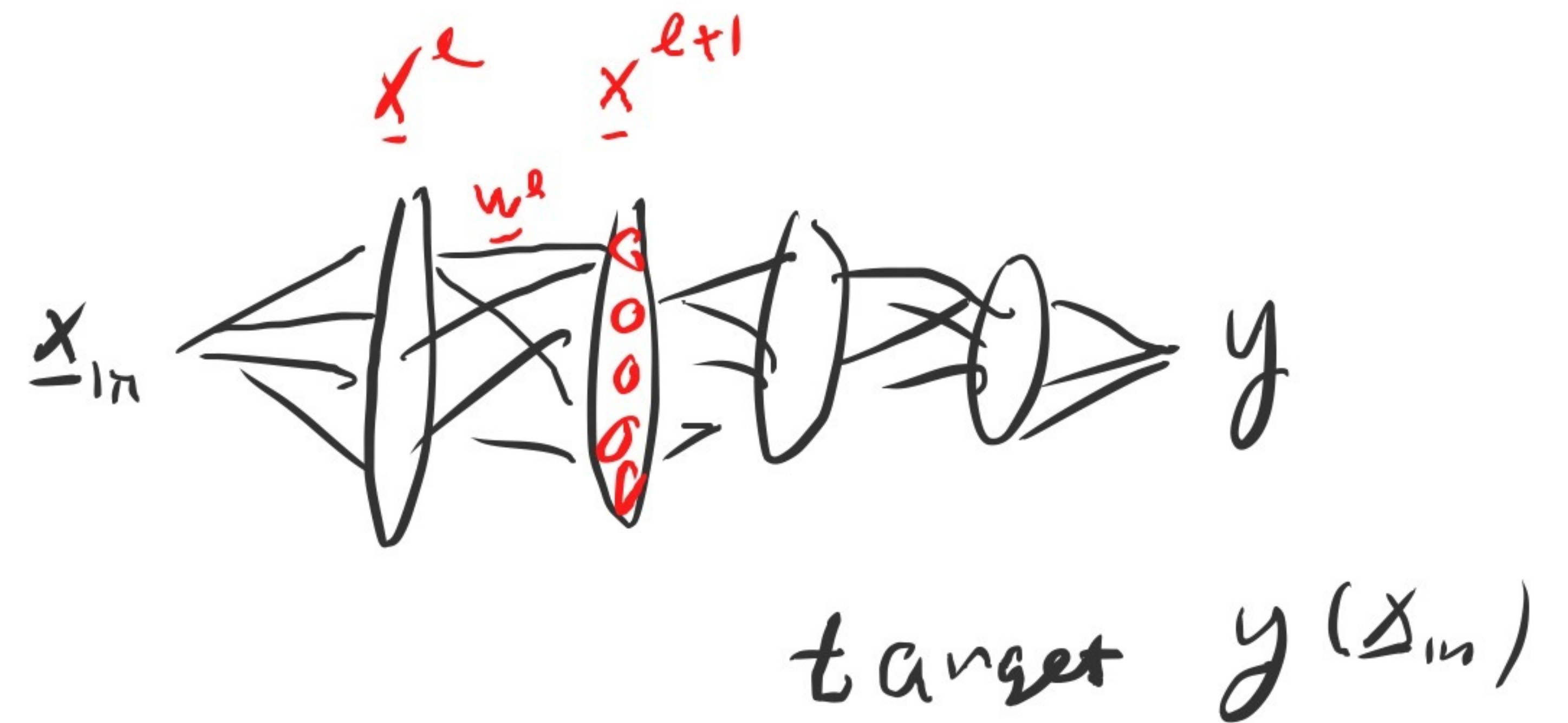
$$\underline{w} = \underline{v}_k$$

learning in real networks

1. vanilla deep networks

$$\underline{x}^{l+1} = \phi(\underline{h}^l)$$

$$\underline{h}^l = \underline{w}^l \cdot \underline{x}^l$$



$$\mathcal{L}(\underline{x}_{in}, y(\underline{x}_{in}))$$

$$\Delta \underline{w}^l = -\gamma \frac{\partial \mathcal{L}}{\partial \underline{w}^l}$$

$$\frac{\partial \mathcal{L}}{\partial \underline{w}^l} = \frac{\partial \mathcal{L}}{\partial \underline{h}^l} \cdot \frac{\partial \underline{h}^l}{\partial \underline{w}^l}$$

$$= \frac{\partial \mathcal{L}}{\partial \underline{h}^l} \underline{x}^l$$

$$= \delta^l \underline{x}^l$$

$$\left\| \frac{\partial \mathcal{L}}{\partial \underline{w}_i^l} = \frac{\partial \mathcal{L}}{\partial h_i^l} x_i^l \right.$$

$$\frac{\partial \mathcal{L}}{\partial \underline{h}^l} = \frac{\partial \mathcal{L}}{\partial \underline{h}^{l+1}} \cdot \frac{\partial \underline{h}^{l+1}}{\partial \underline{h}^l}$$

$$\underline{h}^{l+1} = \underline{w}^{l+1} \cdot \underline{x}^{l+1}$$

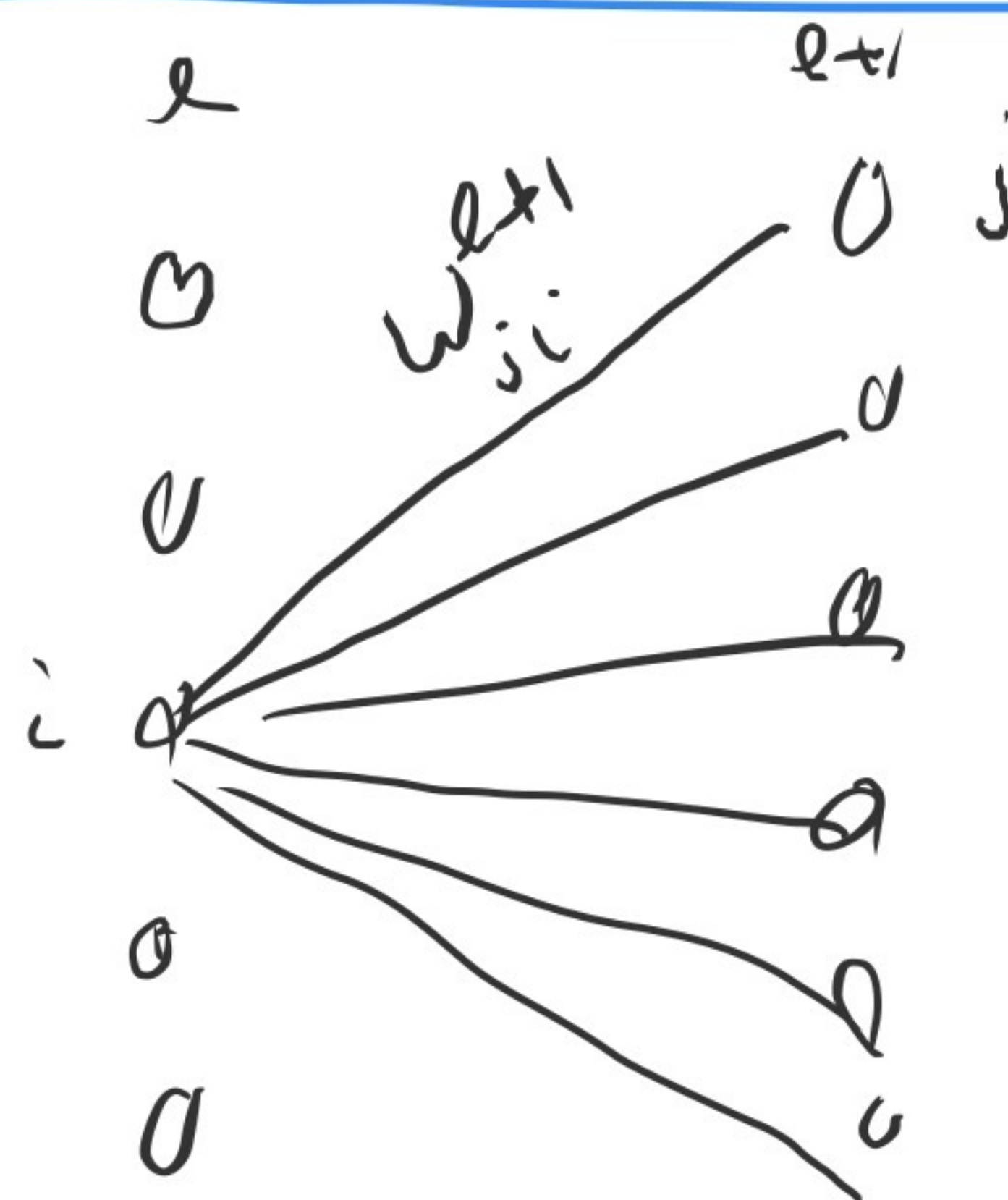
$$= \underline{w}^{l+1} \cdot \phi(\underline{h}^l)$$

$$\frac{\partial \underline{h}^{l+1}}{\partial \underline{h}^l} = \underline{w}^{l+1} \odot \phi'(\underline{h}^l)$$

$$\delta^l = \frac{\partial \mathcal{L}}{\partial \underline{h}^l}$$

$$\delta^l = \delta^{l+1} \cdot \underline{w}^{l+1} \odot \phi'(\underline{h}^l)$$

$$\delta_i^l = \sum_j \delta_j^{l+1} w_{ji}^{l+1} \phi'(h_i^l)$$



$$\underline{\delta}^l = \underline{\delta}^{l+1} \cdot \underline{w}^{l+1} \odot \phi'(\underline{h}^l)$$

↑
outgoing weights

① Tim Lillicrap: replace $\underline{w}^{l+1}$ with a random matrix $\underline{A}$
(feedback alignment)

- works well on simple problems
- doesn't work on hard problems

② upgrade: train $\underline{A}$

$\underline{x}^l$		$\underline{w}^l$	$\underline{x}^{l+1}$
0	$\xrightarrow{\quad}$		0
0	$\xleftarrow{\quad}$		0
0		$\underline{A}^l$	0
0		-	0

training:

$$\Delta \underline{A}^l = \eta \langle \underline{x}^l \underline{x}^{l+1} \rangle$$

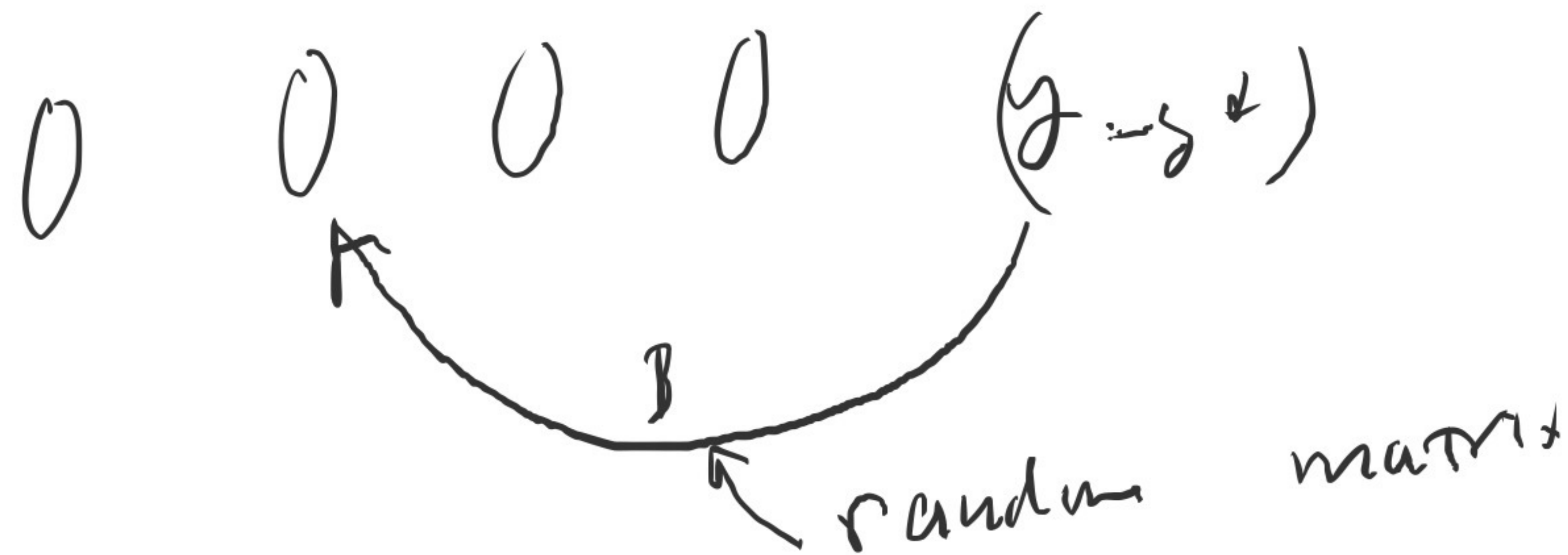
during training: add
white noise to $\underline{x}^l$

$$\Delta \underline{A}^l = \eta \langle \underline{x}^l \underline{x}^l \cdot \underline{w}^{l+1T} \rangle - \epsilon \underline{A}^l$$

$$= \eta \langle \underline{x}^l \underline{x}^l \rangle \cdot \underline{w}^{l+1T} - \epsilon \underline{A}^l$$

$$\rightarrow \underline{w}^{l+1T} = \underline{I}$$

Direct feedback alignment



$$\Delta \underline{\underline{w}}^L = \underline{\underline{B}}^L (y - y^*)$$

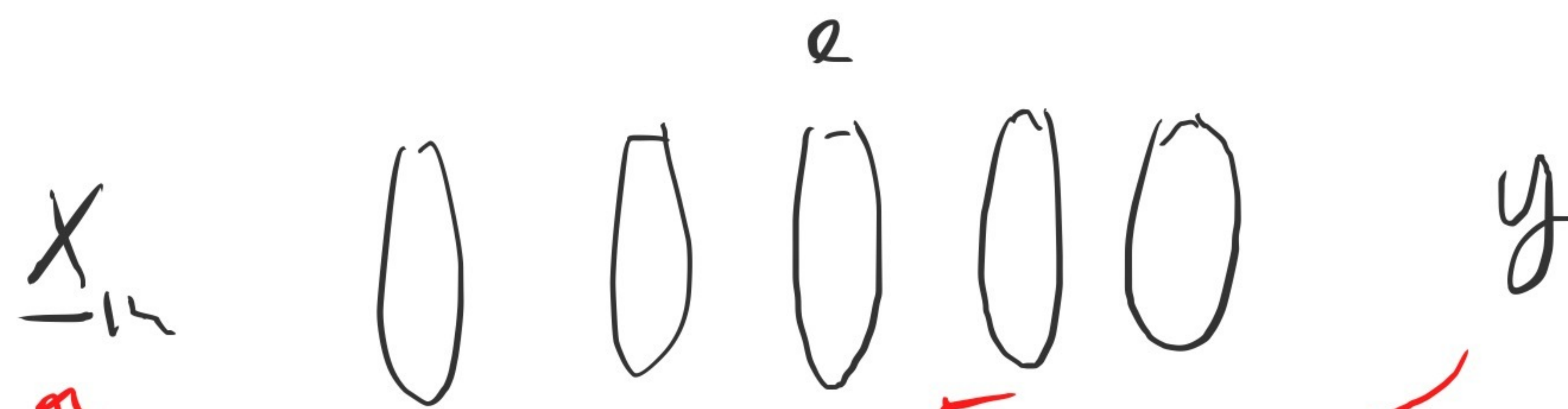
$$\underline{x}^L \quad \Delta w^L = -\gamma \underline{x}^L (y - y^*)$$

A diagram showing a vertical vector of five zeros. To its right, a weight vector w^L is represented by three red arrows pointing towards a target value y .

Works pretty well
on problems that
aren't too hard

Bottleneck approach

3-factor Hebb rule



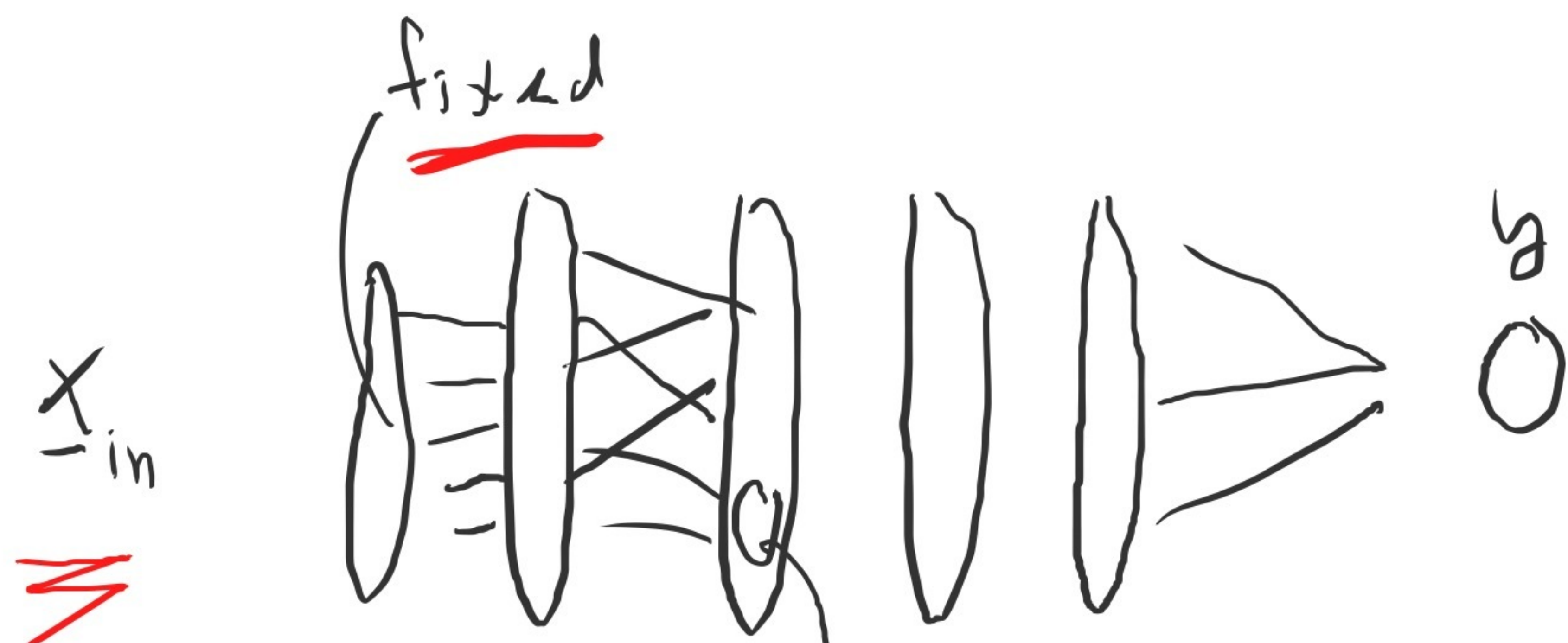
Tali Tishby: irrelevant information is getting squeezed out.

maximize

$$I(x^e, y) - \beta I(x_{in}, x^e)$$

replace

Info w/ some other measure



x^l : goal is to predict y

$$\mathcal{L} = \frac{1}{2} \sum_{l,j} (x_j^l - y^*(x_{in}))^2$$

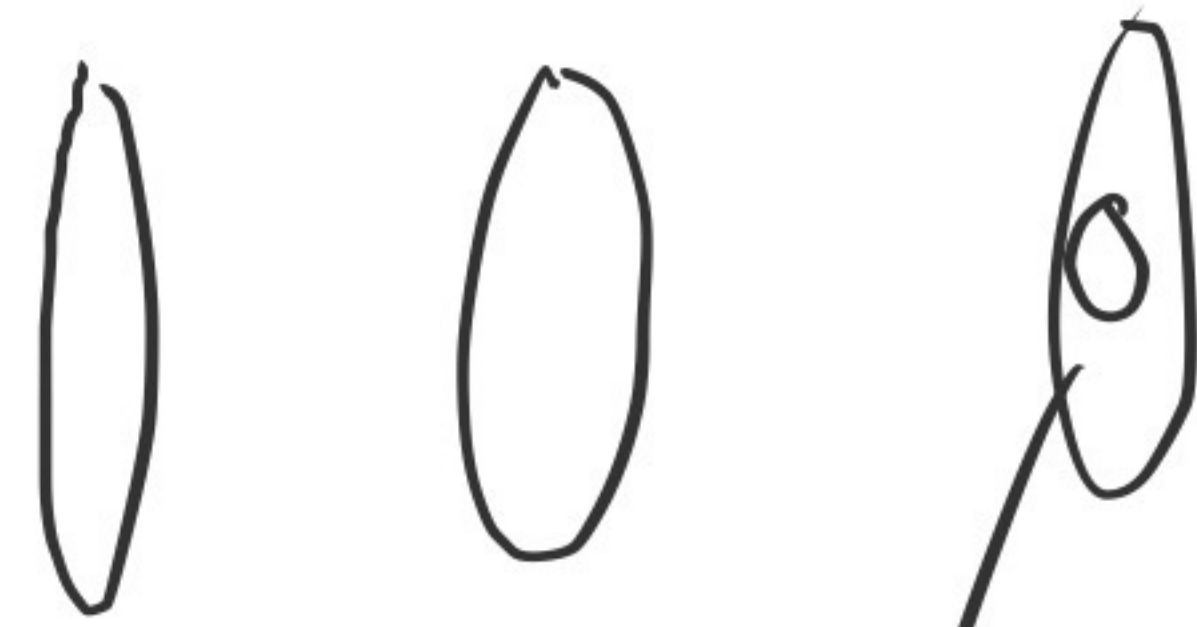
local: Hebb rule

$$= \frac{1}{2} \sum_{l,j} \left(\sum_i w_{ij}^{l-1} x_i^{l-1} - y^*(x_{in}) \right)^2$$

$$\Delta w_{ij}^{l-1} = -\eta \frac{\partial \mathcal{L}^{l-1}}{\partial w_{ij}^{l-1}}$$

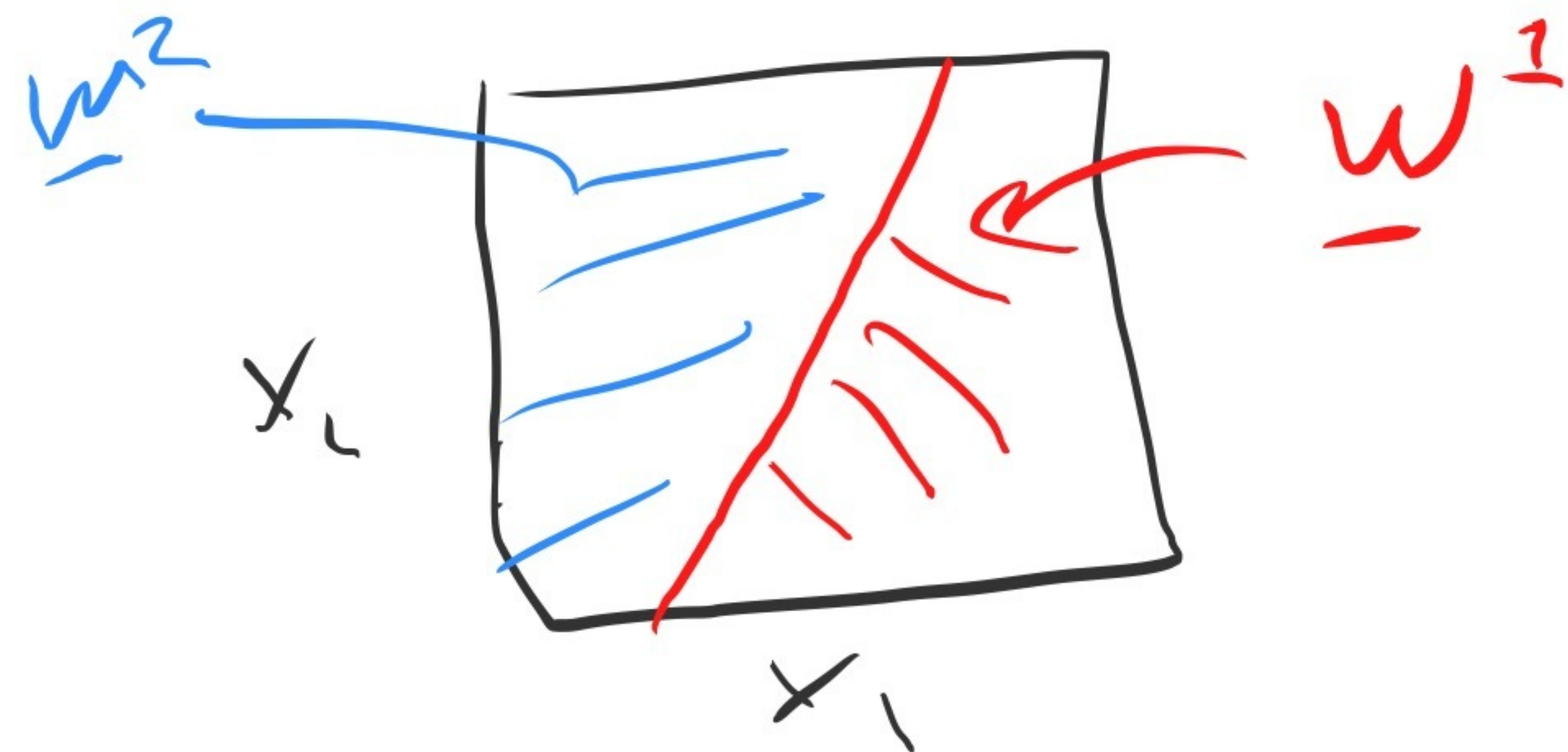
$$= -\eta \underbrace{(x_i^l - y^*)}_{\text{error (post)}} x_j^{l-1}$$

pre-synaptic

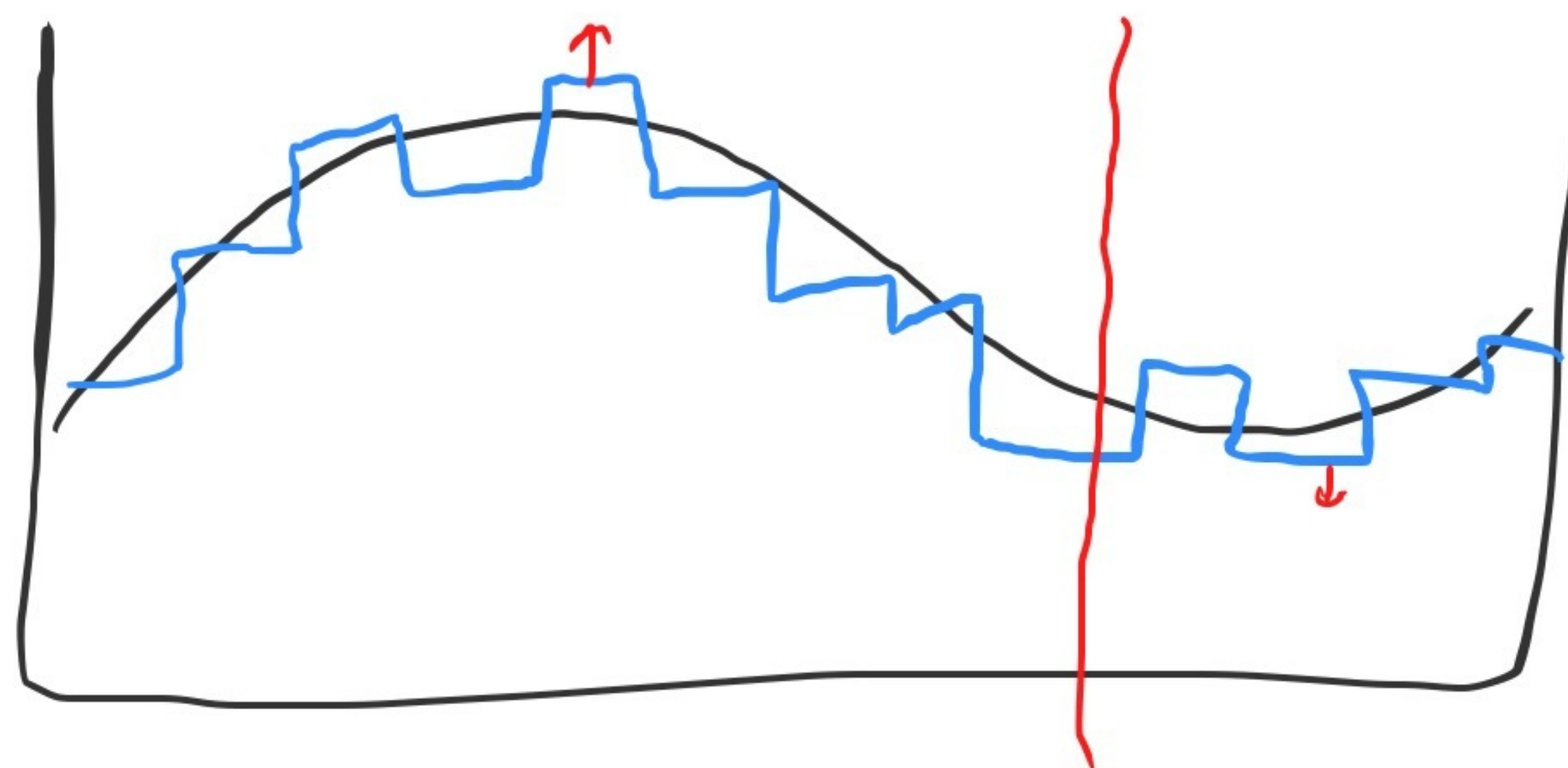
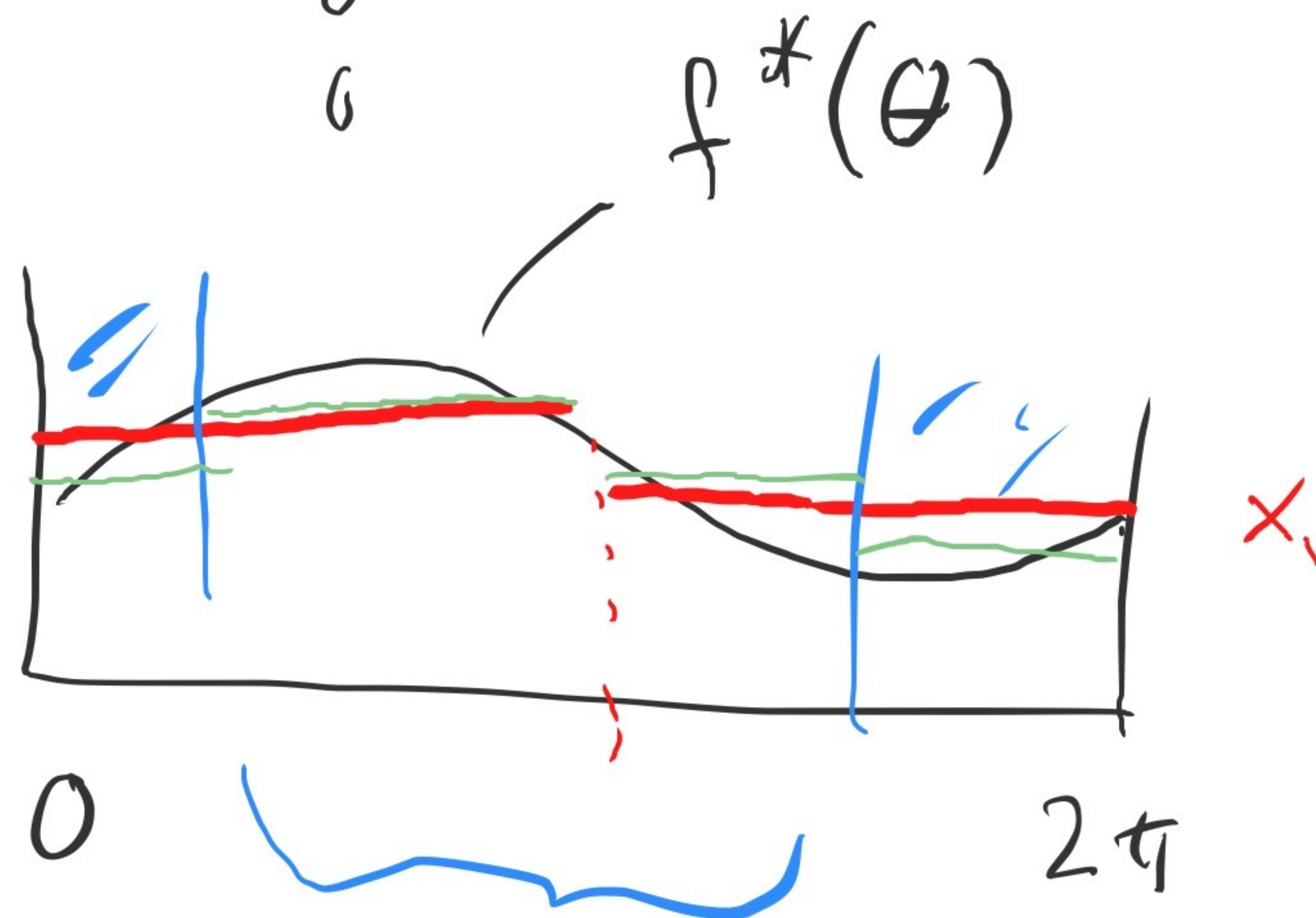


$$x_i^l = \sum_j w_{ij}^{l-1, c(\underline{x}_m)} x_j^{l-1}$$

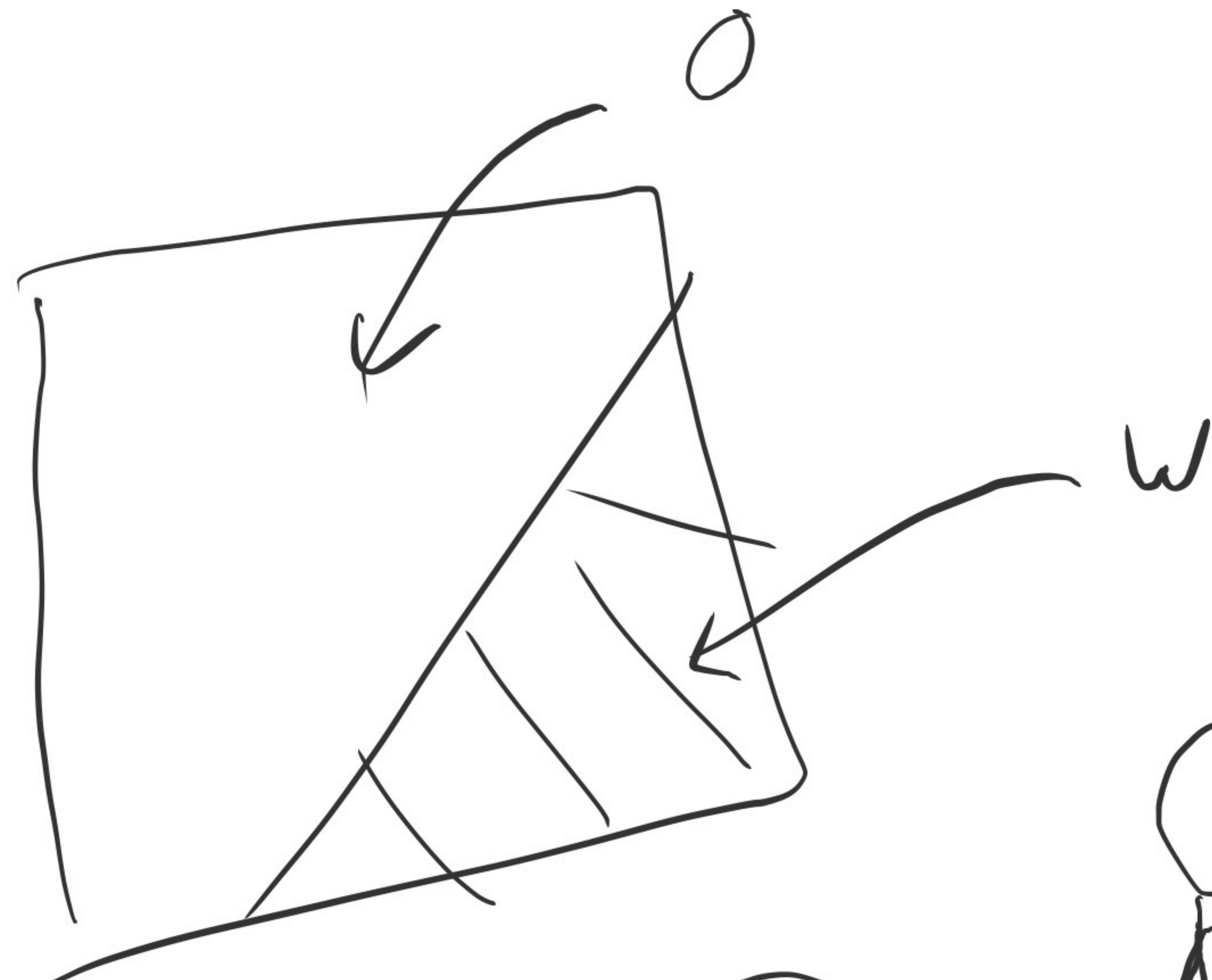
- each neuron has a bank of weights
- which one it chooses depends on input
- pre-chosen



$$\theta \quad x_{1w} \quad 1 \rightarrow 0 \quad x_1 = w_1 \cdot 1 \quad x_2 = w_2 x_1 \quad 0 \quad 0 \quad 0 \quad 0 \quad 0 \quad y$$



1. Convex
2. gets better with depth



GLN

Gated linear networks

$$\underline{r}^l = f(\underline{w}^l \cdot \underline{f}^{-1}(\underline{r}^{l-1}))$$

$$= f(\underline{w}^l \cdot \underline{w}^{l-1} \cdot \underline{f}^{-1}(\underline{r}^{l-2}))$$

each branch
is gated
on or off
(inhibits
neurons)

DGN

dendritic gated
network

Summary:

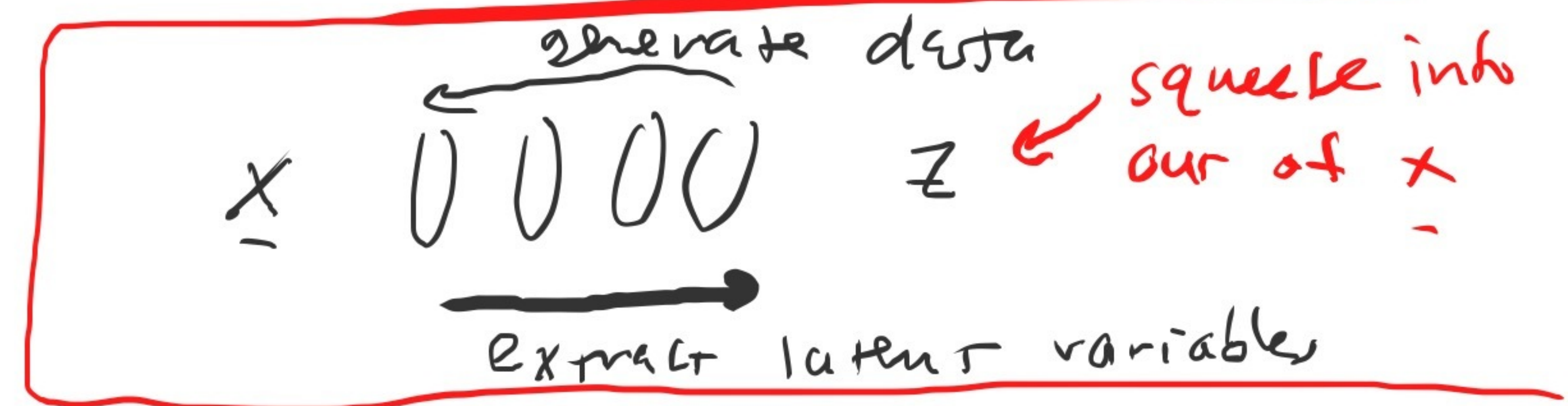
- backprop works really well
- not biologically plausible
 1. need forward and backward pass
 2. neurons need to know their outgoing weights

lots of fakes

- mainly hacks

- GLM (provably converges - convex)

- what the brain uses is unknown



the brain mainly uses unsupervised learning

$\tilde{x}$, model $P_\theta(x)$

$$\mathcal{J} = - \sum_{x \sim \tilde{x}} \log P_\theta(\underline{x})$$
$$\Delta \theta = - \gamma \frac{\partial \mathcal{J}}{\partial \theta}$$

goal of most science:
bridge levels of description

neuroscience

lower level: neural activity

upper level: behavior

goal is to understand behavior in terms of
neural activity

neurons (micro level)

$$C \frac{dV_i}{dt} = \underbrace{\left(\text{single neuron model} \right)}_{\sum_j W_{ij} (V_i - E_j) g_j(t)}$$

conductance changes due to pre-synaptic spikes

$$\Delta W_{ij} = \dots$$

behavior = $f(\underline{s})$

spikes on all the neurons (output neurons)

10^4 eqns

10-100 dimensional

understand:

solve analytically mapping from spikes to behavior

2 parts:

1. What are the equations?
(including initial conditions)

2. What's the solution?

Weather prediction

Simplified
equations

→ make up kind of reasonable
equations and cross our
fingers. (usually simpler)

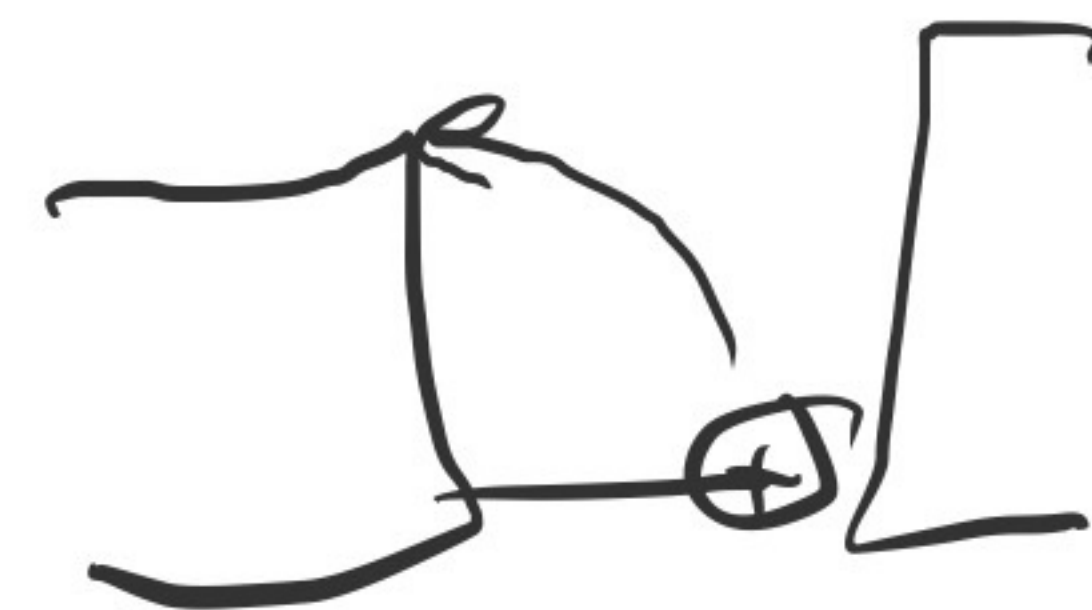
Biophysics

$$C \frac{dv}{dt} = I$$

$$I = - \sum_x g_x (v - E_x)$$

conductance
reversal potential

derive all of biophysics once you
know g_x



constant (leak)

voltage dependent (HH)

- add geometry (cable eqn.)

concentration-dependent

axon

synaptic dynamics



networks

understand
this quantity

$$\sum_j w_{ij} (V_i - \xi_j) g_j(t)$$

fast fluctuations

static piece

$$\rightarrow \sum_j \tilde{w}_{ij} V_j + \sum_j \tilde{w}_{ij} \delta V_j(t)$$

famous
nullcl



random $\Rightarrow$

not random

$$\bar{V}_i, \bar{V}_j$$

are all
that matters

w_{ij} overlaps

large sums are Gaussian

weights fixed

- Switch to learning
- we wrote down learning rules in supervised setting.
- now things get really hard

$$\frac{dw_{ij}}{dt} = f(\underline{w})$$

← averaged over activities

we don't know this one

10^{14}

↘

$$\frac{dw_{ij}}{dt} = f(w_{ij}, \underline{x})$$

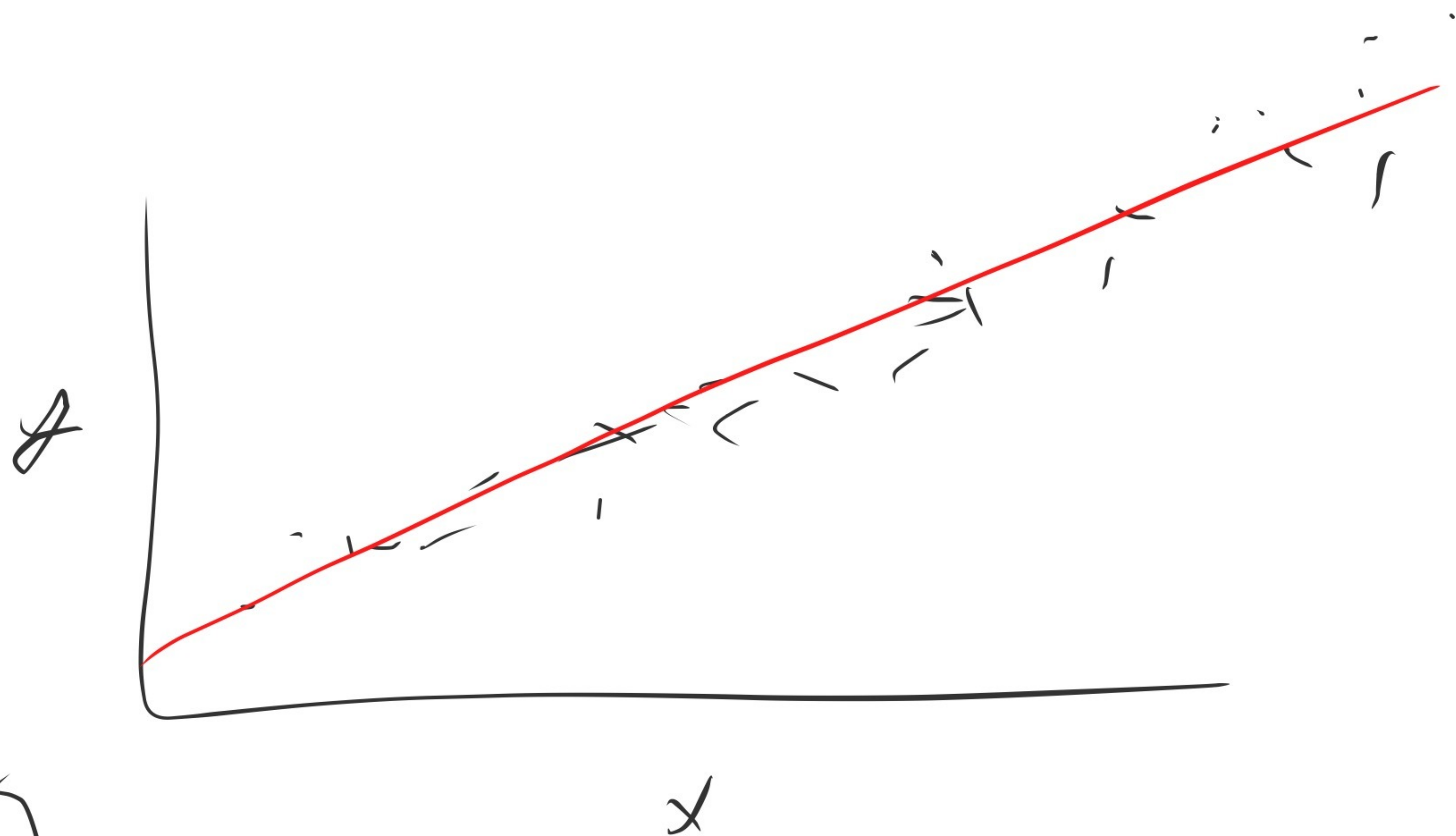
← activities of lot of neurons

$$f(w_{ij}, x_i, x_j)$$

← standard picture: pre-post

$$\frac{dx_i}{dt} = f\left(\sum_j w_{ij} x_j\right)$$

← equation highly nonlinear



if x, y, z true then



$$\frac{P}{n} = \underline{\underline{0.138}}$$

