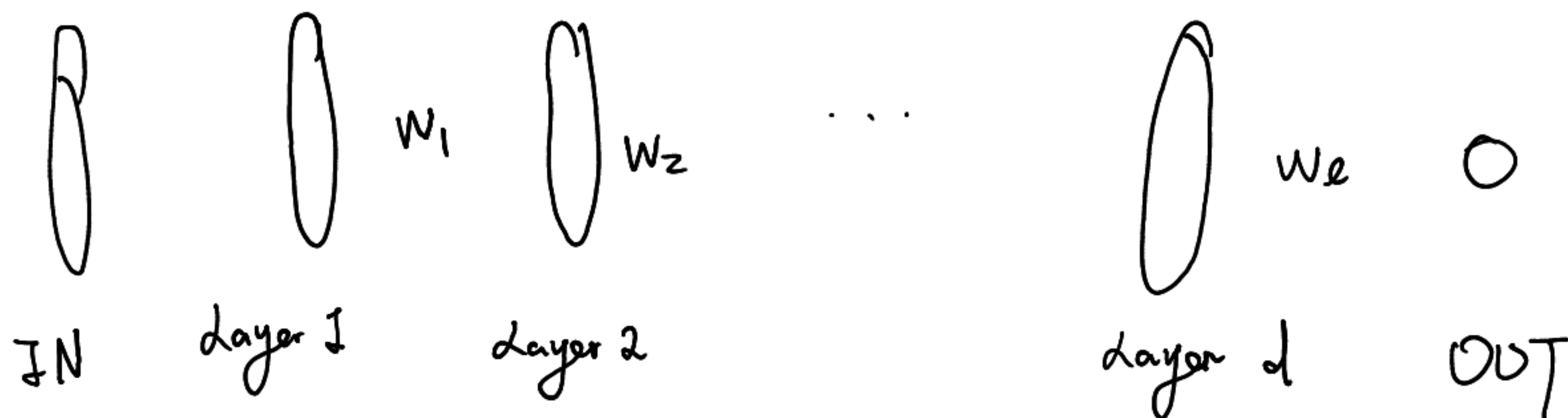


## Lecture Plan ;

- \* Background (Recap from last lecture)
- \* Warmup : Multi-layer linear feedforward neural network
- \* The general case :
  - Forward pass
  - Backward pass (with a first attempt)
- \* Extensions :
  - Acyclic computation graphs
  - Weight sharing (e.g. Conv Nets)
  - Cyclic computation graphs (RNN)
  - Hessian-vector product

★ Background . A natural algorithm for training a neural network is gradient descent . How do we run the gradient descent algorithm efficiently ?

Warmup : linear neural network .



$$\begin{aligned} \text{loss}(x, y) &= \left( w_d w_{d-1} \dots w_1 x - y \right)^2 \\ &= \lambda(w_1, w_2, \dots, w_d) \end{aligned}$$

Gradients .

$$\begin{aligned} &\frac{\partial \lambda(w_1, \dots, w_d)}{\partial w_i} \\ &\quad \bullet \lambda(w_1, \dots, w_{i-1}, w_i + \delta, w_{i+1}, \dots, w_d) \\ &\quad \quad - \lambda(w_1, \dots, w_{i-1}, w_i, w_{i+1}, \dots, w_d) \\ &= (w_d \dots (w_i + \delta) \dots x - y)^2 - (w_d \dots (w_i) \dots x - y)^2 \\ &= 2(w_d \dots w_{i+1} \cdot \delta \cdot w_{i-1} \dots w_1 x - y) \cdot (w_d \dots w_1 x - y) \\ &\quad + O(\delta^2) \end{aligned}$$

$\uparrow$   
 $(1)$

Matrix calculus. (Taylor expansion of matrix functionals) (Expand a bit more)

$$d(\dots, w_i + \delta, \dots) - d(\dots, w_i, \dots)$$

$$= \left\langle \frac{\partial d(w_1, \dots, w_d)}{\partial w_i}, \delta \right\rangle \quad (2)$$

Remark: inner product  $\langle a, b \rangle = \sum_{i=1}^n a_i b_i$ , for two vectors  $a, b \in \mathbb{R}^n$ , and

$$\langle X, Y \rangle = \sum_{i=1}^{n_1} \sum_{j=1}^{n_2} X_{i,j} Y_{i,j}, \text{ for two matrices } X, Y \in \mathbb{R}^{n_1 \times n_2}.$$

So, our goal is to simplify Eq(1) to a form like Eq(2).

$$w_d \dots w_{i+1} \delta w_i \dots w_1 x$$

$$= \text{Tr} [w_d \dots w_{i+1} \delta w_i \dots w_1 x]$$

$$= \text{Tr} [w_i \dots w_1 x w_d \dots w_{i+1} \delta] \quad (\text{Tr}[ABC] = \text{Tr}[CAB])$$

$$= \langle (w_i \dots w_1 x w_d \dots w_{i+1})^T, \delta \rangle$$

$$= \langle (w_d \dots w_{i+1})^T (w_i \dots w_1 x)^T, \delta \rangle$$

$$\Rightarrow \text{Eq}(2) = d(w_d \dots w_1 x - y) \cdot (\langle (w_d \dots w_{i+1})^T (w_i \dots w_1 x)^T, \delta \rangle - y)$$

$$\text{Hence } \frac{\partial \mathcal{L}}{\partial W_i} = 2(W_d \dots W_1 x - y) \cdot (W_d \dots W_{i+1})^T \cdot (W_i \dots W_1 x)^T. \quad (3)$$

How do we implement the above gradient computation efficiently?

Forward Pass.

IN:  $x$

Output of layer 1:  $W_1 \cdot x \leftarrow a_1$

$\vdots$

Output of layer  $i$ :  $W_i \dots W_1 x \leftarrow a_i$

Notice that  $a_i$  appears in Eq (3)!

Backward Pass.

OUT:  $y$

Gradient wrt to quadratic loss function:

$$g_d = 2(W_d \dots W_1 x - y)$$

Gradient wrt to the input to the  $d$ -th layer:

$$\begin{aligned} g_{d-1} &= \frac{\partial (W_d \cdot z_d - y)^2}{\partial z_d} = 2W_d^T (W_d \cdot z_d - y) \\ &= W_d^T \cdot g_d \end{aligned}$$

Gradient wrt to the input the  $(d-1)$ -th layer:

$$\begin{aligned}
 J_{d-2} &= \frac{\partial (W_d \cdot W_{d-1} \cdot z_{d-1} - y)^2}{\partial z_{d-1}} \\
 &= 2 W_{d-1}^T W_d^T (W_d \cdot W_{d-1} \cdot z_{d-1} - y) \\
 &= W_{d-1}^T \cdot J_{d-1} \\
 &\vdots
 \end{aligned}$$

$$J_i = W_{i+1}^T \cdot J_{i+1}$$

Put together.

Observe that  $\frac{\partial d}{\partial W_i} = J_i \cdot a_i^T !$

Therefore, our full algorithm is:

Backpropagation for linear neural net

IN :  $(x_1, y_1) \dots (x_m, y_m)$

Req :  $(W_1^{(t)}, \dots, W_d^{(t)}), \eta$

OUT :  $(W_1^{(t+1)}, \dots, W_d^{(t+1)})$

Forward pass :

compute  $a_1, a_2, \dots, a_d$

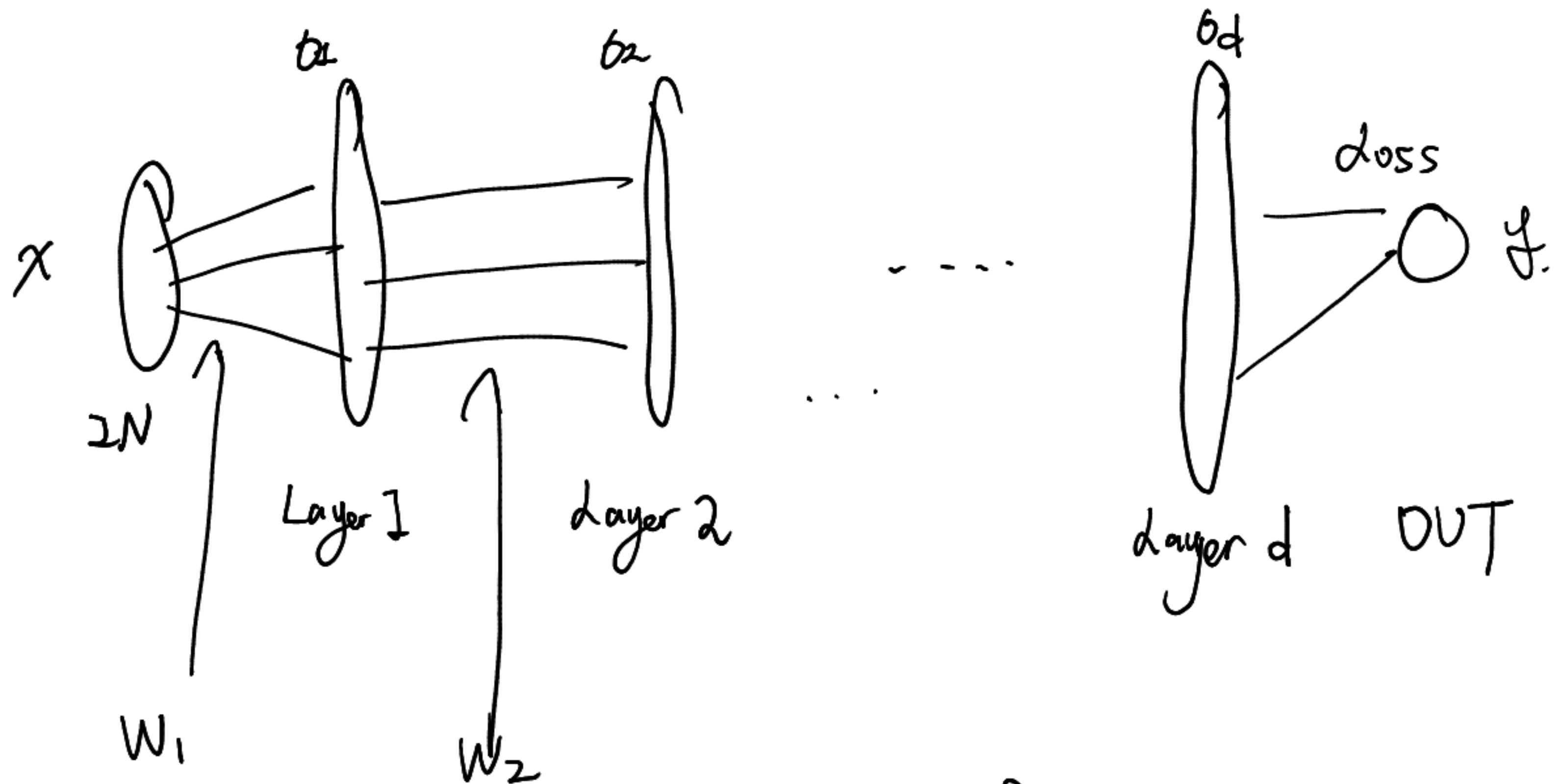
Backward pass :

compute  $J_d, J_{d-1}, \dots, J_1$

update  $W_i^{(t+1)} \leftarrow W_i^{(t)} - \eta \cdot J_i \cdot a_i^T$

The general case. Non-linear activation, arbitrary loss function, bias in activation function.

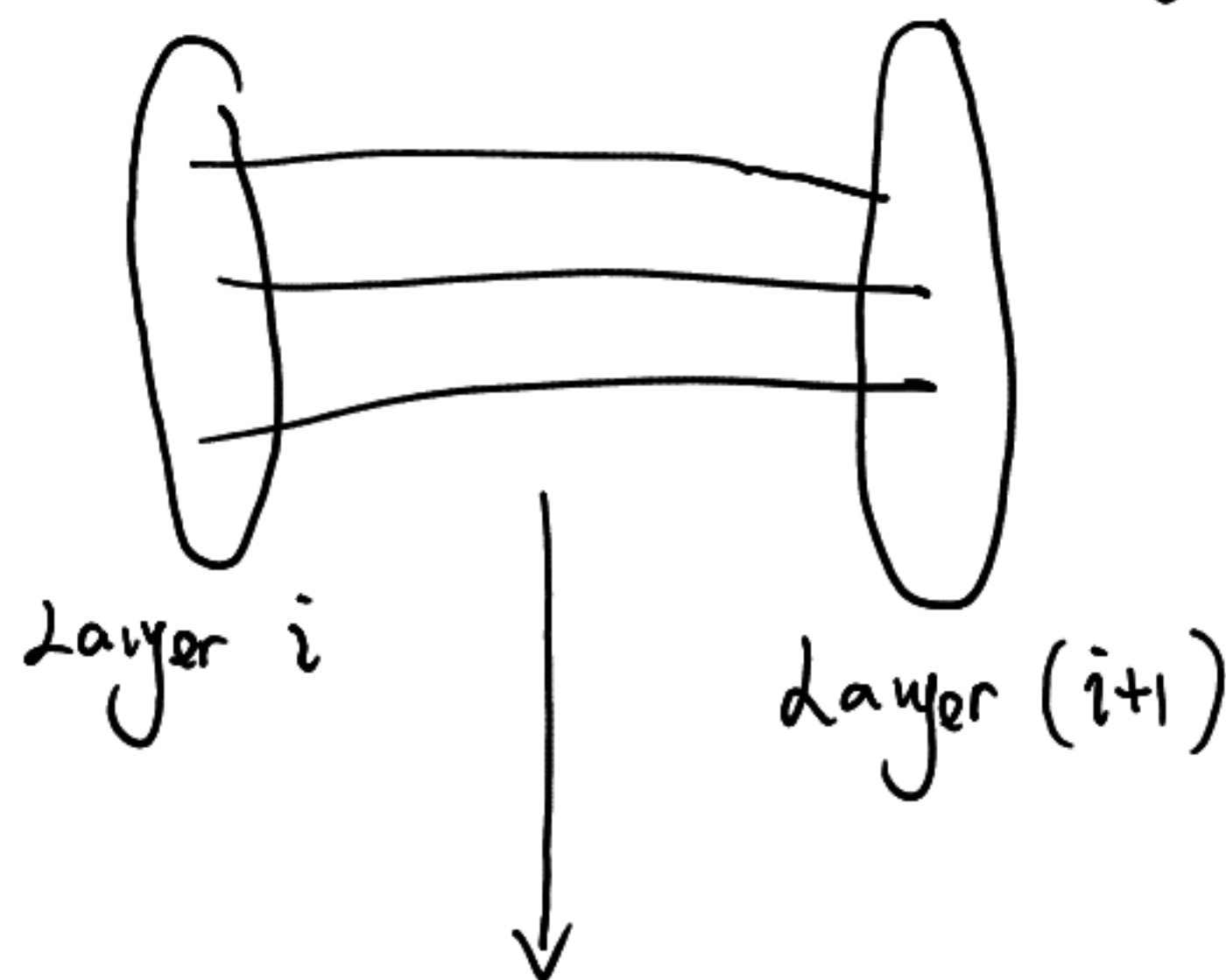
Model.



Chain rule.

We would like to calculate

$$\frac{\partial f(W_1, W_2, \dots, W_d)}{\partial W_i} ?$$



abbrev.  $\frac{\partial f}{\partial W_i}$

Weight matrix:  $W_i \in \mathbb{R}^{r_i \times r_{i+1}}$

Output of Layer  $i$ :  $a_i \in \mathbb{R}^{r_i}$

Input to Layer  $i+1$ :  $z_i \in \mathbb{R}^{r_{i+1}}$

Claim 1.  $\frac{\partial f}{\partial W_i} = \frac{\partial f}{\partial z_i} \cdot a_i^T$

Proof. Chain rule:  $\frac{\partial f}{\partial w_i} = \frac{\partial f}{\partial z_i} \cdot \frac{\partial z_i}{\partial w_i} = \frac{\partial f}{\partial z_i} \cdot a_i^T$ .

By Def.  $z_i = w_i \cdot a_i = \langle a_i^T, w_i \rangle$ , hence  $\frac{\partial z_i}{\partial w_i} = a_i^T$ .

Claim 2.  $\frac{\partial f}{\partial z_i} = \left( w_{i+1}^T \cdot \frac{\partial f}{\partial z_{i+1}} \right) \odot b_{i+1}'(z_i)$ .

Notation  $\odot$  for two vectors  $x, y \in \mathbb{R}^n$ , their Kronecker product

$$x \odot y = (x_1 y_1, x_2 y_2, \dots, x_n y_n) \in \mathbb{R}^n.$$

• for a function  $b: \mathbb{R} \rightarrow \mathbb{R}$ , we use the vector notation  $b(x) = (b(x_1), b(x_2), \dots, b(x_n))$ , and  $b'(x) = (b'(x_1), b'(x_2), \dots, b'(x_n))$ .

Proof. By def.  $\nabla f = \langle \frac{\partial f}{\partial z_i}, \nabla z_i \rangle$ , for some  $\nabla z_i$  that goes to zero.

So, suppose we perturb  $z_i$  by  $\delta$ , we would like to know how much  $f$  changes. Note that

$$\nabla f = \langle \frac{\partial f}{\partial z_{i+1}}, \nabla z_{i+1} \rangle.$$

So we are going to track how much  $z_{i+1}$  changes.

$$\begin{aligned} z_{i+1} &= w_{i+1} \cdot a_{i+1} \\ &= w_{i+1} \cdot b_{i+1}(z_i + \delta) \\ &= w_{i+1} \cdot (b_{i+1}(z_i) + b_{i+1}'(z_i) \odot \delta) \end{aligned}$$

$$= W_{i+1} \cdot (b_{i+1}(z_i) + b'_{i+1}(z_i) \odot \delta)$$

$$= W_{i+1} \cdot b_{i+1}(z_i) + W_{i+1} \cdot (b'_{i+1}(z_i) \odot \delta)$$

$$\Rightarrow \forall z_{i+1} = W_{i+1} \cdot (b'_{i+1}(z_i) \odot \delta). \quad \text{Therefore}$$

$$\nabla f = \left\langle \frac{\partial f}{\partial z_{i+1}}, W_{i+1} \cdot (b'_{i+1}(z_i) \odot \delta) \right\rangle$$

$$\left\langle \frac{\partial f}{\partial z_i}, \delta \right\rangle = \text{Tr} \left[ \left( \frac{\partial f}{\partial z_{i+1}} \right)^T \cdot W_{i+1} \cdot (b'_{i+1}(z_i) \odot \delta) \right]$$

$$= \left\langle W_{i+1}^T \cdot \frac{\partial f}{\partial z_{i+1}}, b'_{i+1}(z_i) \odot \delta \right\rangle.$$

$$\Rightarrow \frac{\partial f}{\partial z_i} = \left( W_{i+1}^T \cdot \frac{\partial f}{\partial z_{i+1}} \right) \odot b'_{i+1}(z_i).$$

**Remark.** Suppose that  $b_{i+1}(z_i) = z_i$  (i.e. linear activation),

then  $b'_{i+1}(z_i) = 1$ , therefore

$$\frac{\partial f}{\partial z_i} = W_{i+1}^T \cdot \frac{\partial f}{\partial z_{i+1}}, \quad \text{Combining Claim 1 and 2,}$$

we have shown that

$$\frac{\partial f}{\partial w_i} = W_{i+1}^T \cdot \frac{\partial f}{\partial z_{i+1}} \cdot a_i^T, \quad \text{which recovers our earlier result!}$$

Put together, we have that

$$\frac{\partial f}{\partial w_i} = \frac{\partial f}{\partial z_i} \cdot a_i^T,$$

$$\frac{\partial f}{\partial z_i} = \left( w_{i+1}^T \cdot \frac{\partial f}{\partial z_{i+1}} \right) \cdot \delta'_{i+1}(z_i).$$

Goal. Compute  $\{a_i\}$ ,  $\left\{\frac{\partial f}{\partial z_i}\right\}$ ,  $\{\delta'_{i+1}(z_i)\}$ .

Recall that  $z_i$  is the input to the  $(i+1)$ -th layer,  
and  $a_i$  is the output of the  $i$ -th layer.

Therefore,

$$z_i = w_i \cdot a_i,$$

Forward Pass  $\{a_i\}$ .  $a_0 = x$

$$a_{i+1} = b(z_{i+1})$$

$$= b(w_i \cdot a_i).$$

Backward Pass.  $\left\{\frac{\partial f}{\partial z_i}\right\}$ .  $\frac{\partial f}{\partial z_d} = \frac{\partial \text{loss}(b_{d+1}(z_d, y^{\text{true}}))}{\partial z_d}$ .

$$\frac{\partial f}{\partial z_i} = \left( w_{i+1}^T \cdot \frac{\partial f}{\partial z_{i+1}} \right) \cdot \delta'_{i+1}(z_i).$$

Put together, we have derived the backpropagation algorithm.

---

## Backpropagation for feedforward neural networks

---

Forward Pass:

Compute  $\{a_i\}, \{z_i\}$ .

Backward Pass:

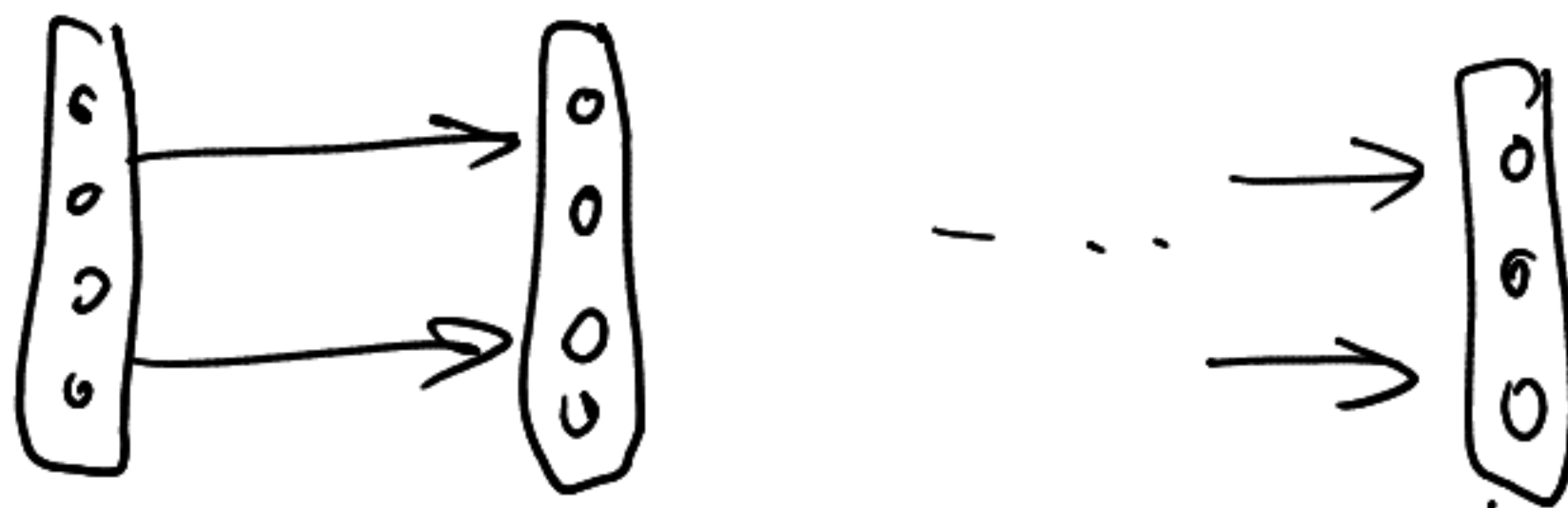
Compute  $\left\{ \frac{\partial f}{\partial z_i} \right\}$ , hence  $\left\{ \frac{\partial f}{\partial w_i} \right\}$ .

Apply gradient descent over  $\{w_i\}$ .

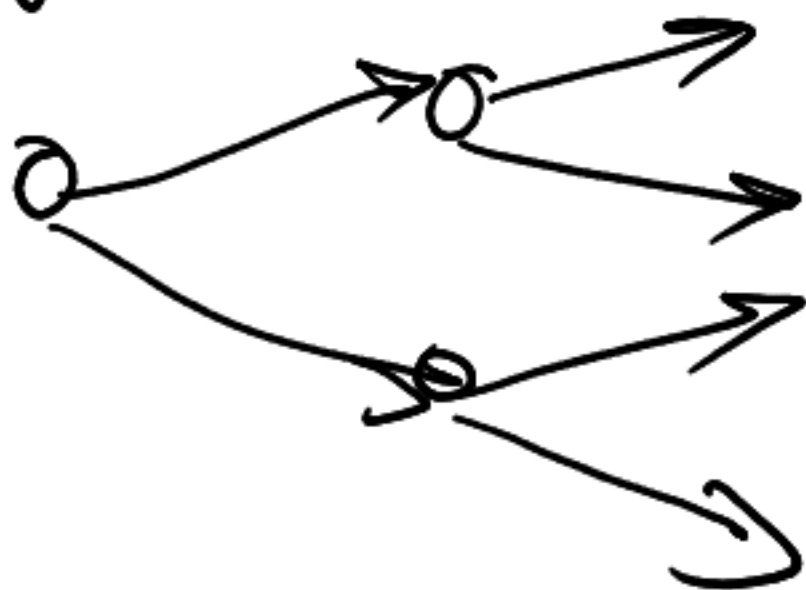
---

### Extension 5.

— Acyclic neural network

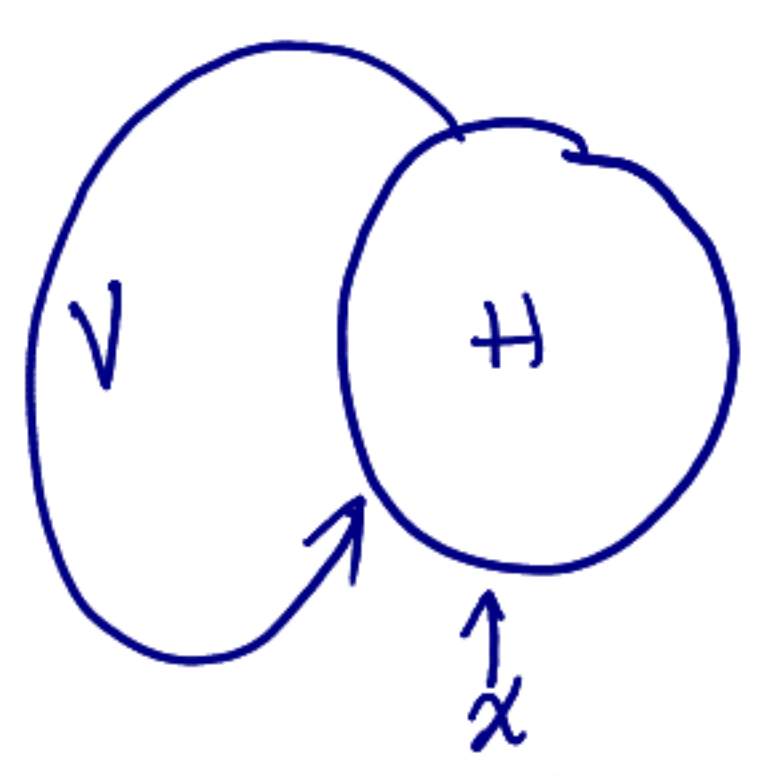


In general, the algorithm works as long as there are no cycles in the network!



What if there are cycles? e.g. RNN? yes but we may run into issues.

Exploding or vanishing gradients.

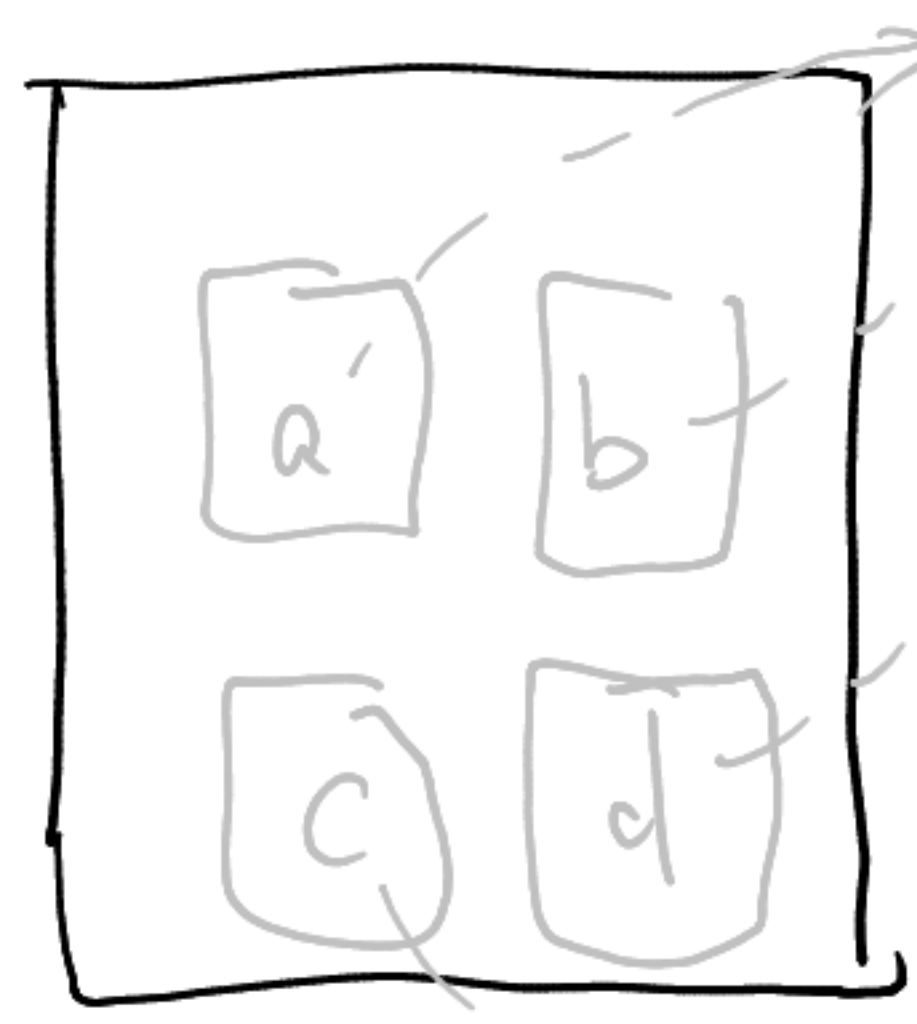


- i) if  $\|V\|_2 < 1$ , then  $\|V^i \cdot x\| < \sigma_{\max}^i(V) \cdot \|x\|$
- ii) if  $\|V\|_2 > 1$ , then  $\|V^i \cdot x\| > \sigma_{\min}^i(V) \cdot \|x\|$

Fix. Gradient clipping on long short-term memory.

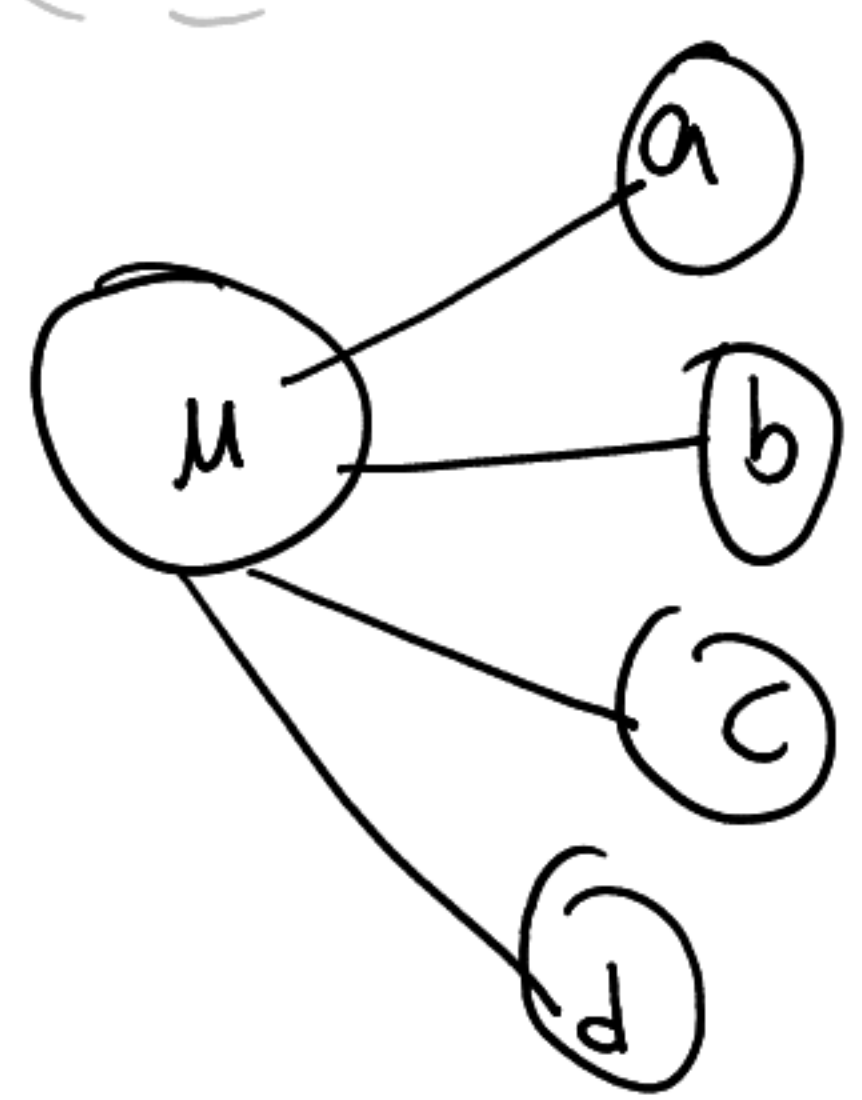
Fix. Backpropagation in time. Expand over the prev.  $k$  steps.

— Weight sharing.



multiple patches share the same weight matrix  $W$ .

Set  
 $\mu = a$   
 $\mu = b$   
 $\mu = c$   
 $\mu = d$



$$\begin{aligned} \frac{\partial f}{\partial \mu} &= \frac{\partial f}{\partial a} \cdot \frac{\partial a}{\partial \mu} + \dots \\ &= \frac{\partial f}{\partial a} + \frac{\partial f}{\partial b} \end{aligned}$$

## Second-order methods:

- Hessian :  $\frac{\partial^2 f}{\partial w_i \partial w_j}$
- Traditional use of Hessian
- Hessian in neural nets

## Conclusions

- We derived the backpropagation algorithm, which involves a forward pass and a backward pass.
- Backpropagation provides a linear-time implementation of gradient descent.
- & key idea : chain rule in calculus.
- Extensions : Acyclic networks, Weight sharing,
- Discussion : second-order methods.