

ML & Optimization Lecture 18

- Nonconvex optimization w/ SGD
- Backpropagation for general FFNs, and MLPs
- Example of 2-layer NN for Fashion MNIST

FFN: L layers (0 input layer, L output layer)

n_l is the number of neurons on layer l $w^l \in \mathbb{R}^{n_l \times n_{l-1}}$

σ activation function

$a^l \in \mathbb{R}^{n_l}$ preactivations, $a^l = w^l o^{l-1} + b^l$

$o^l \in \mathbb{R}^{n_l}$ activation, $o^l = \sigma(a^l)$
 $= \sigma(w^l o^{l-1} + b^l)$ $b^l \in \mathbb{R}^{n_l}$

output of NN given an input $x \in \mathbb{R}^{n_0}$ is

$$o^L(x) = o^L(w^L \sigma(\dots \sigma(w^1 x + b^1) \dots) + b^L)$$

Last time: NNs are parametrized DAG of computation
Once the architecture is fixed, we learn the parameters of
the NN via (for FFNNs)

$$\mathcal{W}_* = \{ \omega_*^l, b_*^l \} = \underset{\substack{\omega^l \in \mathbb{R}^{n_l \times n_{l-1}} \\ b^l \in \mathbb{R}^{n_l}}}{\text{argmin}} f(\mathcal{W})$$

$$\text{where } f(\mathcal{W}) = \frac{1}{n} \sum_{i=1}^n \ell(\phi(x_i), y_i) + \lambda \sum_{l=1}^L R_l(\omega^l, b^l)$$

where R_l is a layer-specific regularizer, e.g. a common
choice is weight decay on ω^l :

$$R_l(\omega^l, b^l) = \|\omega^l\|_F^2 = \sum_{i \in [n_l]} \sum_{j \in [n_{l-1}]} (\omega^l)_{ij}^2$$

which encourages small weights

This is in general a non-convex optim problem, so guarantees of the forms:

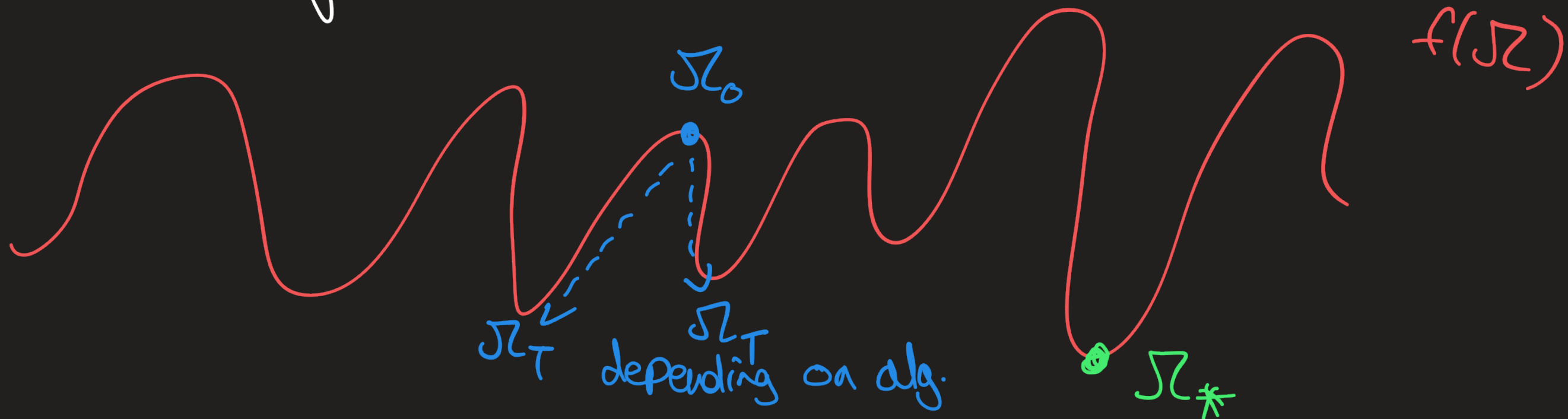
$$(1) f(\mathcal{Z}_T) - f^* \leq \epsilon$$

$$(2) \|\mathcal{Z}_T - \mathcal{Z}_*\|_F^2 \leq \epsilon$$

convergence of func values to global min (ϵ -suboptimal)
convergence of parameters to global min

$$\sum_{\ell \in [L]} \left[\|\omega^\ell - \omega_*^\ell\|_F^2 + \|b^\ell - b_*^\ell\|_2^2 \right]$$

do not hold, because in case (2) we can only hope to converge to local minimizers, and in case (1) there is no guarantee to a global minimum value.



In the case of general nonconvex optim, we guarantee convergence to approximate stationary points, i.e. parameters $\mathcal{J}$ satisfying

$$\|\nabla f(\mathcal{J})\|_F^2 = \sum_{\ell \in [L]} \left(\|\nabla_{\mathbf{w}^\ell} f(\mathcal{J})\|_F^2 + \|\nabla_{\mathbf{b}^\ell} f(\mathcal{J})\|_2^2 \right) \leq \varepsilon.$$

We call such an $\mathcal{J}$ an ε -approximate stationary point. Intuitively, this suggests $\mathcal{J}$ is near an extremum of f (could be max, min, or saddle pt). For certain classes of functions, we can quantify this intuition.

We use first-order algs, especially (adaptive) stochastic gradient descent, to train NNs. Luckily we can prove that first-order algs converge to ϵ -approx stat. pts on smooth non-convex functions.

Ex

If f is L -smooth, GD w/ stepsize $\alpha = \frac{1}{L}$ satisfy

$$\frac{1}{T} \sum_{i=1}^T \|\nabla f(x_{i-1})\|_2^2 \leq \frac{2LR}{T}$$

where $R = f(x_0) - f^*$

Note that this implies

$$\min_{i \in [T]} \|\nabla f(x_{i-1})\|_2^2 \leq \frac{1}{T} \sum_{i=1}^T \|\nabla f(x_{i-1})\|_2^2 \leq \varepsilon$$

for $T = O(\frac{1}{\varepsilon})$, so with that many iterates we visit at least one ε -stat pt.

It also implies that if $\bar{i}$ randomly choose an iterate x from $x_0, \dots, x_{T-1}$

$$\mathbb{E} \|\nabla f(x)\|_2^2 \leq \varepsilon$$

so x is an ε -approx stat. point in expectation

Proof

↓ b/c of L -smoothness of f

$$f(x_t) \leq f(x_{t-1}) + \langle \nabla f(x_{t-1}), x_t - x_{t-1} \rangle + \frac{L}{2} \|x_t - x_{t-1}\|_2^2$$

recall $x_t = x_{t-1} - \alpha \nabla f(x_{t-1})$, so

$$= f(x_{t-1}) - \alpha \|\nabla f(x_{t-1})\|_2^2 + \frac{\alpha^2 L}{2} \|\nabla f(x_{t-1})\|_2^2$$

$$= f(x_{t-1}) + \alpha \left(\frac{\alpha L}{2} - 1 \right) \|\nabla f(x_{t-1})\|_2^2$$

since $\alpha = \frac{1}{L}$, we have

$$f(x_t) - f(x_{t-1}) \leq -\frac{1}{2L} \|\nabla f(x_{t-1})\|_2^2$$

or equivalently

$$\frac{1}{2L} \|\nabla f(x_{t-1})\|_2^2 \leq f(x_{t-1}) - f(x_t)$$

now sum these inequalities over $t=1, \dots, T$:

$$\frac{1}{2L} \sum_{t=1}^T \|\nabla f(x_{t-1})\|_2^2 \leq \sum_{t=1}^T \{f(x_{t-1}) - f(x_t)\}$$

$$= \{f(x_0) - f(x_1)\} + \{f(x_1) - f(x_2)\} + \dots + \{f(x_{T-1}) - f(x_T)\}$$

$$= f(x_0) - f(x_T)$$

$$\leq f(x_0) - f^* = R$$

and now multiply both sides by $\frac{2L}{T}$, to get

$$\frac{1}{T} \sum_{t=1}^T \|\nabla f(x_{t-1})\|_2^2 \leq \frac{2LR}{T}$$

What about SGD? We can imagine since g_t is only an estimate of $\nabla f(x_t)$, we will pay the price in the form of slower convergence.

Assume f is L -smooth, and

$$- \mathbb{E}[g_t | x_t] = \nabla f(x_t)$$

$$- \mathbb{E} \|g_t - \nabla f(x_t)\|_2^2 \leq \sigma^2$$

and let R be as before, after T rounds of SGD

$$\left\{ \mathbb{E} \left(\frac{1}{T} \sum_{t=1}^T \|\nabla f(x_{t-1})\|_2^2 \right) \leq \frac{2LR}{T} + \sigma^2 \text{ if } \alpha = \frac{1}{2L} \right.$$

$$\left. \mathbb{E} \left(\frac{1}{T} \sum_{t=1}^T \|\nabla f(x_{t-1})\|_2^2 \right) \leq 4 \sqrt{\frac{RL\sigma^2}{T}} \text{ if } \alpha = \sqrt{\frac{R}{TL\sigma^2}} \right.$$

$$\text{and } T \geq \frac{4RL}{\sigma^2}$$

We now know SGD works in a specific quantifiable sense. In practice, we reach local minima and their quality is good.

The practical question left is: how to compute

$\nabla f(\mathcal{Z})$, where $\mathcal{Z} = \{\omega^l, b^l\}_{l=1}^L$ and

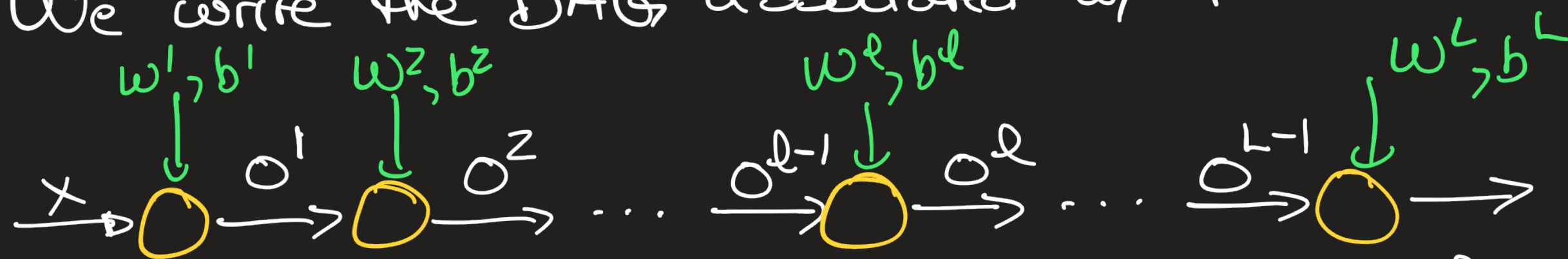
$$f(\mathcal{Z}) := \frac{1}{n} \sum_{i=1}^n \ell(\phi(x_i), y_i) + \lambda \sum_{l=1}^L R_l(\omega^l, b^l)$$

WLOG we will consider the case $n=1$, so

$$f(\mathcal{Z}) := \ell(\phi^L(x), y) + \lambda \sum_{l=1}^L R(\omega^l, b^l)$$

Recall this is loss for a general FFN

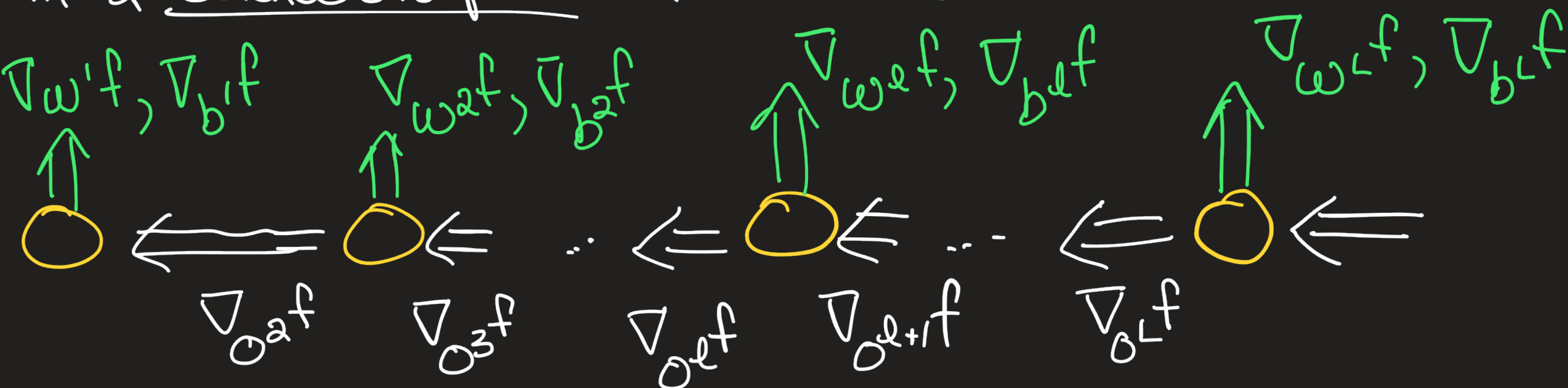
We write the DAG associated w/ f :



$$f = \ell(o^L, y) + \lambda R(\mathcal{W})$$

We see that computing the value of each layer is done in a forward pass: $o^l = \sigma^l(o^{l-1})$

Similarly, gradients w.r.t. parameters can be computed in a backward pass (chain rule):



Justification: Chain Rule

Consider $h: \mathbb{R}^n \rightarrow \mathbb{R}^p$ and $r: \mathbb{R}^p \rightarrow \mathbb{R}$

Define $\phi(z) = r(h(z))$

Think of $h: \mathbb{R}^{d-1} \rightarrow \mathbb{R}^d$ as one hidden layer, $h = o^d$
and $r: \mathbb{R}^{d+1} \rightarrow \mathbb{R}$ as $d(o^L(o^{L-1}(\dots(o^{d+1}(z))\dots))$

We claim that

$$\nabla \phi(z) = J_h(z)^T \nabla r(h(z))$$

where $J_h(z) = \left[\frac{\partial h_i}{\partial z_j} \right]_{\substack{i=1, \dots, p \\ j=1, \dots, n}} \in \mathbb{R}^{p \times n}$

is the Jacobian of h at z .

Notice $J_h(z) = \begin{bmatrix} \nabla h_1(z)^T \\ \vdots \\ \nabla h_p(z)^T \end{bmatrix}.$

Also it is a fact that for any function g , the gradient is the only vector that satisfies

$$g(x+\Delta) = g(x) + \langle \nabla g(x), \Delta \rangle + \text{h.o.t.}$$

So to compute $\nabla \phi(z)$, we expand

$$\phi(z + \Delta) = r(h(z + \Delta)) = r\left(\begin{bmatrix} h_1(z + \Delta) \\ \vdots \\ h_p(z + \Delta) \end{bmatrix}\right)$$

$$= r\left(\begin{bmatrix} h_1(z) + \langle \nabla h_1(z), \Delta \rangle + \text{h.o.t.} \\ \vdots \\ h_p(z) + \langle \nabla h_p(z), \Delta \rangle + \text{h.o.t.} \end{bmatrix}\right)$$

$$= r\left(h(z) + \begin{bmatrix} \nabla h_1(z)^T \\ \vdots \\ \nabla h_p(z)^T \end{bmatrix} \Delta + \text{h.o.t.}\right)$$

$$= r\left(h(z) + J_h(z) \circ \Delta + \text{h.o.t.}\right)$$

$$= r(h(z)) + \langle \nabla r(h(z)), J_h(z) \circ \Delta + \text{h.o.t.} \rangle + \text{h.o.t.}$$

$$= r(h(z)) + \langle \nabla r(h(z)), J_h(z) \Delta \rangle + \text{h.o.t.}$$

$$= r(h(z)) + \langle J_h(z)^T \nabla r(h(z)), \Delta \rangle + \text{h.o.t.}$$

$$= \phi(z) + \langle J_h(z)^T \nabla r(h(z)), \Delta \rangle + \text{h.o.t.}$$

so by uniqueness of gradients,

$$\nabla_z \phi(z) = J_h(z)^T \nabla_r(h(z))$$

Returning to backward pass

First key observation: if we can compute $\Leftarrow$

$$\nabla_{o^l} f = \left[\frac{\partial f}{\partial (o^l)_i} \right]_{i=1}^{n_l}$$

then the chain rule allows us to compute

$$\nabla_{w^l} f = J_{o^l}(w^l)^T \nabla_{o^l} f + \lambda \nabla_{w^l} R_l$$

$$(*) \quad \nabla_{b^l} f = J_{o^l}(b^l)^T \nabla_{o^l} f + \lambda \nabla_{b^l} R_l$$

e.g. take $h = o^l$ which takes b^l as input

and $r = \text{id} \circ o^L \circ o^{L-1} \circ \dots \circ o^{l+1}$ which takes o^l as input

then $(*)$ follows from the chain rule

Second key observation

We can compute ∇ recursively. Since f depends on $\mathbf{o}^l$ only through $\mathbf{o}^{l+1}$, we have by the chain rule that

$$(*) \quad \nabla_{\mathbf{o}^l} f = J_{\mathbf{o}^{l+1}}(\mathbf{o}^l)^T \nabla_{\mathbf{o}^{l+1}} f$$

here $h = \mathbf{o}^{l+1}$ takes $\mathbf{o}^l$ as input

$r = l \circ \mathbf{o}^L \circ \mathbf{o}^{L-1} \circ \dots \circ \mathbf{o}^{l+2}$ takes $\mathbf{o}^{l+1}$ as input

so

$$J_h(\mathbf{o}^l)^T = J_{\mathbf{o}^{l+1}}(\mathbf{o}^l)^T$$

and

$$\nabla_{\mathbf{o}^{l+1}} r = \nabla_{\mathbf{o}^{l+1}} f$$

so

(*) follows from the chain rule

Putting these pieces together we see that we can recursively compute $\nabla_{w^L} f$, $\nabla_{b^L} f$ in a single backward pass as follows:

Backprop for FFNs

$$\text{temp} \leftarrow \nabla_{o^L} f$$

$$\nabla_{b^L} f = J_{o^L}(b^L)^T \text{temp} + \lambda \nabla_{b^L} R_L$$

$$\nabla_{w^L} f = J_{o^L}(w^L)^T \text{temp} + \lambda \nabla_{w^L} R_L$$

for $l = L-1, \dots, 1$:

$$\text{temp} \leftarrow J_{o^{l+1}}(o^l)^T \text{temp}$$

$$\nabla_{b^l} f \leftarrow J_{o^l}(b^l)^T \text{temp} + \lambda \nabla_{b^l} R_l$$

$$\nabla_{w^l} f \leftarrow J_{o^l}(w^l)^T \text{temp} + \lambda \nabla_{w^l} R_l$$