

In this course we will need to learn the weights of the neural network. As we have discussed in lecture that we will be using gradient descent to learn the network weights.

$$W^{(K+1)} = W^{(K)} - \epsilon \nabla_W \mathcal{L}$$

In the gradient descent step we will need to compute $\nabla_W \mathcal{L}$. In a neural network the weights in the earlier layers are connected to the loss function through a composition of functions, and in such cases computing the gradients require repeated application of the chain rule.

In this discussion we will go over a computational mechanism that will enable us to compute the gradients in an efficient and modular manner.

Chain rule revisited:

Let's consider the following regularized linear classification model

$$z = wx + b$$

$$y = \sigma(z)$$

$$\mathcal{L} = \frac{1}{2} (y - t)^2$$

$$R = \frac{1}{2} w^2$$

$$\mathcal{L}_{reg} = \mathcal{L} + \lambda R$$

we will learn the parameters w and b using gradient descent. Let's use chain rule to compute the partial derivatives $\frac{\partial \mathcal{L}_{\text{reg}}}{\partial w}$, $\frac{\partial \mathcal{L}_{\text{reg}}}{\partial b}$

$$\mathcal{L}_{\text{reg}} = \frac{1}{2} [\sigma(wx+b) - t]^2 + \frac{\lambda}{2} w^2$$

$$\frac{\partial \mathcal{L}_{\text{reg}}}{\partial w} = \frac{1}{2} \frac{\partial}{\partial w} [\sigma(wx+b) - t]^2 + \frac{\lambda}{2} \frac{\partial}{\partial w} w^2$$

$$= [\sigma(wx+b) - t] \frac{\partial}{\partial w} [\sigma(wx+b) - t] + \lambda w$$

$$= [\sigma(wx+b) - t] \sigma'(wx+b) \frac{\partial}{\partial w} (wx+b) + \lambda w$$

$$= [\sigma(wx+b) - t] \sigma'(wx+b) x + \lambda w$$

$$\begin{aligned}
\frac{\partial \mathcal{L}_{\text{reg}}}{\partial b} &= \frac{\partial}{\partial b} \left[\frac{1}{2} (\sigma(\omega x + b) - t)^2 + \frac{\lambda}{2} \omega^2 \right] \\
&= \frac{1}{2} \frac{\partial}{\partial b} [\sigma(\omega x + b) - t]^2 + \frac{\lambda}{2} \frac{\partial}{\partial b} \omega^2 \\
&= [\sigma(\omega x + b) - t] \frac{\partial}{\partial b} [\sigma(\omega x + b) - t] + 0 \\
&= [\sigma(\omega x + b) - t] \sigma'(\omega x + b) \frac{\partial}{\partial b} (\omega x + b) \\
&= [\sigma(\omega x + b) - t] \sigma'(\omega x + b)
\end{aligned}$$

The above computation involves

- (a) Lots of copies of terms from one line to next
- (b) Lots of redundant work. For instance, the first 3 steps in the two derivations above are nearly identical.

The idea behind backprop is to share repeated computations wherever possible.

Computational graphs and backprop:

Computational graphs help to visualize and contextualize how we aim to backpropagate derivatives. It is the mechanism by which deep learning libraries compute gradients.

A computational graph is a directed acyclic graph where each node in the graph denotes a mathematical operation.

Let's draw the computational graph for our running example,

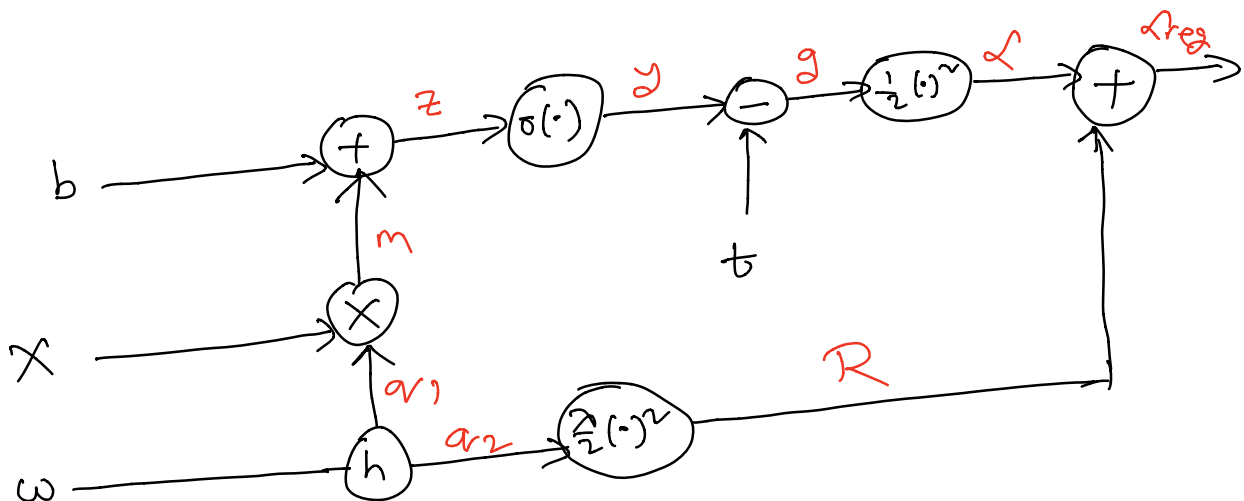
$$z = wx + b$$

$$y = \sigma(z)$$

$$\mathcal{L} = \frac{1}{2} (y - t)^2$$

$$\mathcal{R} = \frac{1}{2} \omega^2$$

$$\mathcal{L}_{reg} = \mathcal{L} + \lambda \mathcal{R}$$



The backprop algorithm consists of two steps:

- (i) In forward Pass we take the input X and the initial values of w and b and propagate in the forward direction through the network to compute the quantities in red.
- (ii) In backward Pass we compute the derivatives starting at the output and propagating it backward to the input.

So,

$$\mathcal{L}_{\text{reg}} = \mathcal{L} + R$$

$$\frac{\partial \mathcal{L}_{\text{reg}}}{\partial \mathcal{L}} = 1, \quad \frac{\partial \mathcal{L}_{\text{reg}}}{\partial R} = 1$$

$$\mathcal{L} = \frac{1}{2} g^2$$

$$\frac{\partial \mathcal{L}}{\partial g} = g$$

$$\begin{aligned} \frac{\partial \mathcal{L}_{\text{reg}}}{\partial g} &= \frac{\partial \mathcal{L}}{\partial g} \frac{\partial \mathcal{L}_{\text{reg}}}{\partial \mathcal{L}} \\ &= g \end{aligned}$$

$$R = \frac{\lambda}{2} a_2^2$$

$$\frac{\partial R}{\partial a_2} = \lambda a_2$$

$$\begin{aligned} \frac{\partial \mathcal{L}_{\text{reg}}}{\partial a_2} &= \frac{\partial R}{\partial a_2} \frac{\partial \mathcal{L}_{\text{reg}}}{\partial R} \\ &= \lambda a_2 \end{aligned}$$

$$g = y - t$$

$$\frac{\partial g}{\partial y} = 1$$

$$\frac{\partial \mathcal{L}_{reg}}{\partial y} = \frac{\partial g}{\partial y} \frac{\partial \mathcal{L}_{reg}}{\partial g}$$

$$= g$$

$$y = \sigma(z)$$

$$\frac{\partial y}{\partial z} = \sigma'(z)$$

$$\frac{\partial \mathcal{L}_{reg}}{\partial z} = \frac{\partial y}{\partial z} \frac{\partial \mathcal{L}_{reg}}{\partial y}$$

$$= \sigma'(z) g$$

$$z = b + m$$

$$\frac{\partial z}{\partial b} = 1, \quad \frac{\partial z}{\partial m} = 1$$

$$\frac{\partial \mathcal{L}_{reg}}{\partial b} = \frac{\partial z}{\partial b} \frac{\partial \mathcal{L}_{reg}}{\partial z}$$

$$= \sigma'(z) g = \sigma'(wx + b) [\sigma(wx + b) - t]$$

$$\begin{aligned}\frac{\partial \mathcal{L}_{\text{reg}}}{\partial m} &= \frac{\partial z}{\partial m} \frac{\partial \mathcal{L}_{\text{reg}}}{\partial z} \\ &= \sigma'(z) g\end{aligned}$$

$$m = q_1 x$$

$$\frac{\partial m}{\partial q_1} = x \quad \frac{\partial m}{\partial x} = q_1$$

$$\begin{aligned}\frac{\partial \mathcal{L}_{\text{reg}}}{\partial x} &= \frac{\partial m}{\partial x} \frac{\partial \mathcal{L}_{\text{reg}}}{\partial m} \\ &= q_1 \sigma'(z) g\end{aligned}$$

$$\begin{aligned}\frac{\partial \mathcal{L}_{\text{reg}}}{\partial q_1} &= \frac{\partial m}{\partial q_1} \frac{\partial \mathcal{L}_{\text{reg}}}{\partial m} \\ &= x \sigma'(z) g\end{aligned}$$

now by law of total derivatives,

$$\frac{\partial \mathcal{L}_{\text{reg}}}{\partial \omega} = \frac{\partial a_1}{\partial \omega} \frac{\partial \mathcal{L}_{\text{reg}}}{\partial a_1} + \frac{\partial a_2}{\partial \omega} \frac{\partial \mathcal{L}_{\text{reg}}}{\partial a_2}$$

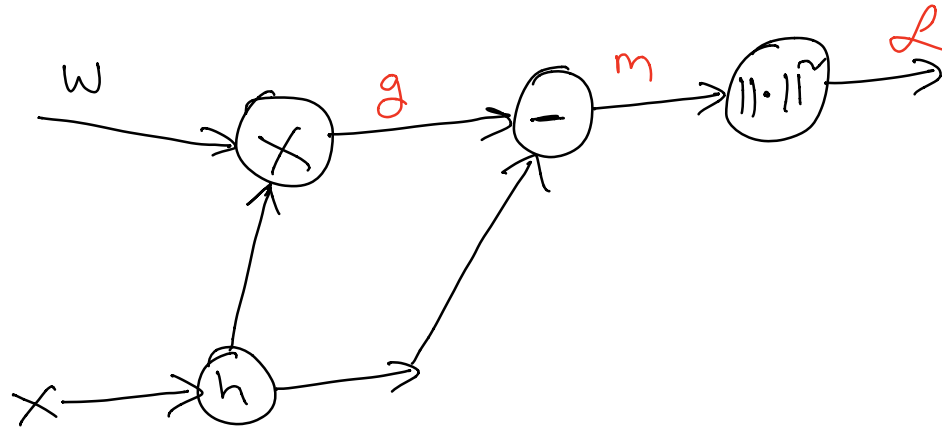
$$\frac{\partial \mathcal{L}_{\text{reg}}}{\partial \omega} = \frac{\partial \mathcal{L}_{\text{reg}}}{\partial a_1} + \frac{\partial \mathcal{L}_{\text{reg}}}{\partial a_2}$$

$$= x \sigma'(z)g + \lambda a_2$$

$$= x \sigma'(\omega x + b) [\sigma(\omega x + b) - t] + \lambda \omega$$

Practice problems:

1



In the above computational graph,

$$h(x) = Ix; \quad I \in \mathbb{R}^{n \times n}$$

$$g = wx; \quad g \in \mathbb{R}^n$$

$$m = g - x; \quad m \in \mathbb{R}^n$$

So,

$$\mathcal{L} = \|m\|_2^2$$

$$\nabla_m \mathcal{L} = 2m \in \mathbb{R}^n \quad (\text{Eqn } 131)$$

$$m = g - x$$

$$\nabla_g m = I \in \mathbb{R}^{n \times n}$$

$$\begin{aligned}\nabla_g \alpha &= \nabla_g^m \nabla_m \alpha \\ &= 2m \in \mathbb{R}^n\end{aligned}$$

$$g = Wx$$

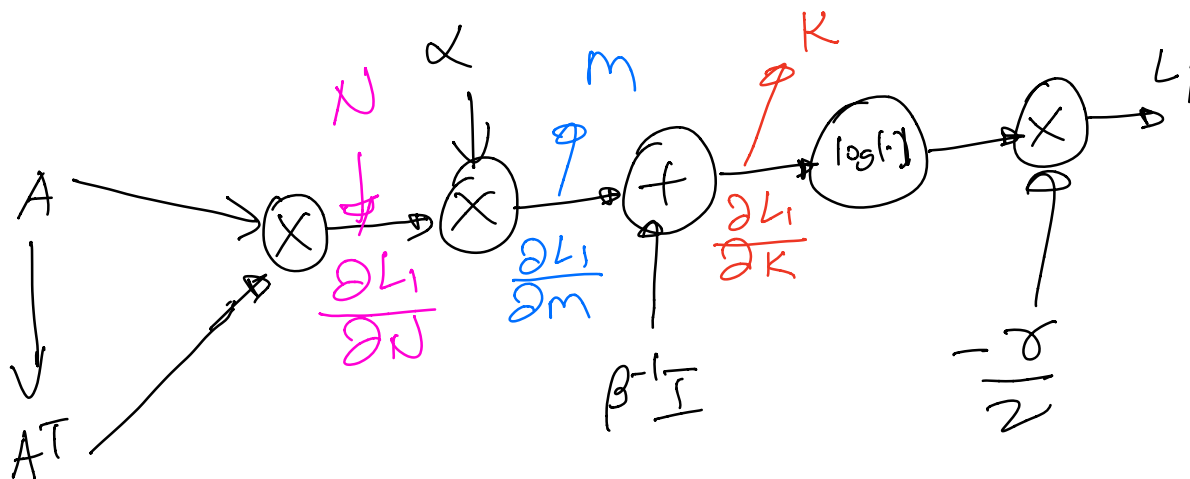
$$\nabla_W g = \nabla_W (Wx) \in \mathbb{R}^{n \times n \times m}$$

$$\begin{aligned}\nabla_W \alpha &= \nabla_W g \nabla_g \alpha \\ &= \nabla_g \alpha x^T \\ &= 2(Wx - x)x^T\end{aligned}$$

$$\therefore \nabla_W \alpha = 2(Wx - x)x^T$$

$\frac{2}{}$

a) The computational graph for α_1 is drawn below:



From the computational graph,

$$L_1 = -\frac{\gamma}{2} \log |K| = -\frac{\gamma}{2} \log (\det K)$$

Using matrix cookbook

$$\frac{\partial L_1}{\partial K} = -\frac{\gamma}{2} (K^{-1})^T = -\frac{\gamma}{2} (K^T)^{-1}$$

(Eaⁿ 57)

Since $(+)$ operator distributes the gradient, so

$$\frac{\partial L_1}{\partial m} = -\frac{\sigma}{2} (K^T)^{-1}$$

Since $(\times)$ operator switches the gradient, so

$$\frac{\partial L_1}{\partial N} = -\frac{\alpha \sigma}{2} (K^T)^{-1}$$

Now,

$$N = AA^T$$

using the hint given in the problem

$$\frac{\partial L_1}{\partial A} = \frac{\partial L_1}{\partial N} (A^T)^T$$

$$\frac{\partial L_1}{\partial A} = -\frac{\alpha \sigma}{2} (K^T)^{-1} A$$

Since A^T has the same elements as A , so taking the gradient of $\mathcal{L}_1$ w.r.t A^T we have

$$\frac{\partial \mathcal{L}_1}{\partial A^T} = A^T \frac{\partial \mathcal{L}_1}{\partial N}$$

$$= A^T \left[-\frac{\alpha \gamma}{2} (K^T)^{-1} \right]$$

$$\frac{\partial \mathcal{L}_1}{\partial A^T} = -\frac{\alpha \gamma}{2} (K^{-1} A)^T$$

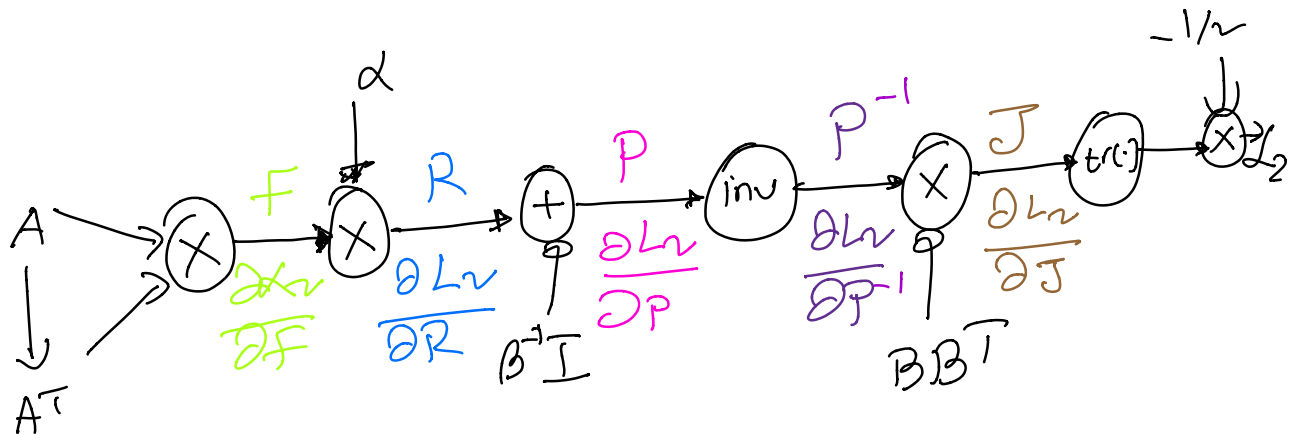
Now adding the gradient paths through A and A^T we get

$$\frac{\partial \mathcal{L}_1}{\partial A} = -\frac{\alpha \gamma}{2} (K^T)^{-1} A - \left[\frac{\alpha \gamma}{2} (K^{-1} A)^T \right]^T$$

Since K is symmetric,

$$\frac{\partial \mathcal{L}_1}{\partial A} = -\alpha \gamma K^{-1} A$$

b) The computational graph for $\hat{x}_2$ is drawn below:



From the computational graph,

$$L_2 = -\frac{1}{2} \text{tr}(J)$$

Now, from matrix cookbook

$$\frac{\partial L_2}{\partial J} = -\frac{1}{2} \mathbf{I} \quad (\text{Eqn } 99)$$

From computational graph,

$$J = P^{-1}(BB^T)$$

Hence using the hint

$$\frac{\partial L_2}{\partial P^{-1}} = -\frac{1}{2} BB^T$$

Now using the given hint

$$\frac{\partial L_2}{\partial P} = -P^{-1} \frac{\partial L_2}{\partial P^{-1}} P^{-1}$$

Since $\oplus$ operator distributes the gradient, so

$$\frac{\partial L_2}{\partial R} = -P^{-1} \frac{\partial L_2}{\partial P^{-1}} P^{-1}$$

Since $\otimes$ operator switches the gradient,

so

$$\frac{\partial \mathcal{L}_r}{\partial F} = -\alpha P^{-1} \frac{\partial \mathcal{L}_r}{\partial P^{-1}} P^{-1}$$

Now,

$$F = AA^T$$

using the given hint,

$$\frac{\partial \mathcal{L}_r}{\partial A} = -\alpha P^{-1} \frac{\partial \mathcal{L}_r}{\partial P^{-1}} P^{-1} A$$

$$\frac{\partial \mathcal{L}_r}{\partial A} = \frac{+\alpha}{2} P^{-1} BB^T P^{-1} A$$

Using the same reasoning as (a) to compute the derivative of L_2 with respect to A^T

$$\frac{\partial L_2}{\partial A^T} = -\alpha A^T P^{-1} \frac{\partial \alpha}{\partial P^{-1}} P^{-1}$$

$$= \frac{\alpha}{2} A^T P^{-1} B B^T P^{-1}$$

Now adding the gradients paths through A and A^T we get:

$$\begin{aligned} & \frac{\partial L_2}{\partial A} + \left[\frac{\partial L_2}{\partial A^T} \right]^T \\ &= \frac{\alpha}{2} P^{-1} B B^T P^{-1} A + \left[\frac{\alpha}{2} A^T P^{-1} B B^T P^{-1} \right]^T \\ &= \frac{\alpha}{2} P^{-1} B B^T P^{-1} A + \frac{\alpha}{2} (P^T)^{-1} B B^T (P^T)^{-1} A \end{aligned}$$

Since P is symmetric so

$$= \frac{\alpha}{2} P^{-1} B B^T P^{-1} A + \frac{\alpha}{2} P^{-1} B B^T P^{-1} A$$

$$= \alpha P^{-1} B B^T P^{-1} A$$

3

We will use backpropagation to compute the required gradients.

Starting from the loss function

L

$$L = \|q\|_2^2$$

$$\frac{\partial L}{\partial q} = 2q \in \mathbb{R}^k$$

Now, $q = P - y$

Computing the local gradient

at $\ominus$

$$\frac{\partial q}{\partial P} = I \in \mathbb{R}^{k \times k}$$

$$\frac{\partial q}{\partial y} = -I \in \mathbb{R}^{k \times k}$$

Using chain rule, backpropagating to P , we have

$$\frac{\partial L}{\partial P} = \frac{\partial a}{\partial P} \frac{\partial L}{\partial a}$$

$$\frac{\partial L}{\partial P} = 2a \in \mathbb{R}^K$$

Now,

$$P = 2^{z_2}$$

where we raise each entry of z_2 to the power of 2

$$2^{z_2} = \begin{bmatrix} 2^{z_{2,1}} \\ 2^{z_{2,2}} \\ \vdots \\ 2^{z_{2,K}} \end{bmatrix} \in \mathbb{R}^K$$

Computing the local gradient

@ $z^{(1)}$:

$$\frac{\partial P}{\partial z_2} = \text{diag}(z^{z_2} \ln z) \in \mathbb{R}^{k \times k}$$

Using chain rule backpropagating
to z_2 , we have

$$\frac{\partial L}{\partial z_2} = \frac{\partial P}{\partial z_2} \frac{\partial L}{\partial P}$$

$$\frac{\partial L}{\partial z_2} = z^{z_2} \ln z \odot z a \in \mathbb{R}^k$$

Since $\oplus$ gate passes the gradient, so

$$\frac{\partial L}{\partial u} = 2^{z_2} \ln 2 \odot z_2 \in \mathbb{R}^k$$

$$\frac{\partial L}{\partial b_2} = 2^{z_2} \ln 2 \odot z_2 \in \mathbb{R}^k$$

Now,

$$u = W_2 h_1$$

Computing the local gradient @

$$\odot \frac{\partial u}{\partial h_1} = W_2^T \in \mathbb{R}^{m \times k}$$

Using chain rule, backpropagating to h_1 we have

$$\frac{dL}{dh_1} = \frac{du}{dh_1} \frac{dL}{du}$$

$$\frac{dL}{dh_1} = w_2^T \frac{dL}{du} \in \mathbb{R}^m$$

By the 3-D tensor derivative trick learned in class,

$$\frac{dL}{dW_2} = \frac{dL}{du} h_1^T \in \mathbb{R}^{K \times m}$$

now, $h_1 = \text{relu}(z_1)$

Computing the local gradient @ relu:

$$\frac{dh_1}{dz_1} = \text{diag}(z_1 > 0) \in \mathbb{R}^{m \times m}$$

using chain rule, backpropagating to z_1 , we have

$$\frac{\partial L}{\partial z_1} = \frac{\partial h_1}{\partial z_1} \frac{\partial L}{\partial h_1}$$

$$= \text{diag}(z_i > 0) \frac{\partial L}{\partial h_1}$$

$$\frac{\partial L}{\partial z_1} = \mathbb{I}(z_i > 0) \odot \frac{\partial L}{\partial h_1} \in \mathbb{R}^m$$

since $\oplus$ passes the gradient, so

$$\frac{\partial L}{\partial b_1} = \frac{\partial L}{\partial z_1} \in \mathbb{R}^m$$

$$\frac{\partial L}{\partial t} = \frac{\partial L}{\partial z_1} \in \mathbb{R}^m$$

Now,

$$t = w_1 X$$

Computing the local gradient @ $\odot X$

$$\frac{\partial t}{\partial X} = w_1^T \in \mathbb{R}^{n \times m}$$

Using chain rule, backpropagating to X we have

$$\frac{\partial L}{\partial X} = \frac{\partial t}{\partial X} \frac{\partial L}{\partial t}$$

$$\frac{\partial L}{\partial X} = w_1^T \frac{\partial L}{\partial z_1} \in \mathbb{R}^n$$

By the 3-D tensor derivative
trick learned in class,

$$\frac{\partial L}{\partial W_1} = \frac{\partial L}{\partial t} X^T$$

$$\frac{\partial L}{\partial W_1} = \frac{\partial L}{\partial z_1} X^T \in \mathbb{R}^{m \times n}$$