

STAT 453: Introduction to Deep Learning and Generative Models

Ben Lengerich

Lecture 07: Multi-layer Perceptrons, Backprop

September 24, 2026





Today: From Softmax to Multilayer Perceptrons

- **From Binary to Multi-Class Classification**
- Multilayer Perceptron Architecture
- Nonlinear Activation Functions
- Backpropagation
- Practical Considerations for Training MLPs

Last Time: Autograd, End to End

- Specified a distribution (Bernoulli), wrote the likelihood, took $-\log$ to get binary cross-entropy
- Derived $\partial L / \partial z = \hat{y} - y$ by hand
 - Confirmed autograd agrees
- Today: We generalize from 2 classes to K classes, then stack layers into an MLP

From Binary to Multi-Class Classification

- **Binary** (Lecture 6): $y \mid x \sim \text{Bernoulli}(\hat{y})$, $\hat{y} = \sigma(z)$
- **Multi-class**: $y \mid x \sim \text{Categorical}(\hat{y}_1, \dots, \hat{y}_K)$, one logit z_k per class
- **Softmax** turns K logits into a valid probability distribution:
$$\hat{y}_k = \text{softmax}(z)_k = \exp(z_k) / \sum_j \exp(z_j)$$
- **Sigmoid** is just the $K=2$ special case of softmax

Likelihood → Cross-Entropy Loss

- **Likelihood** for one example (one-hot y):

$$P(y \mid x) = \prod_k \hat{y}_k^{y_k} = \hat{y}_{\{\text{true class}\}}$$

- **Log-likelihood:** $\log P(y \mid x) = \sum_k y_k \log \hat{y}_k$
- **Loss** = negative log-likelihood = $-\sum_k y_k \log \hat{y}_k$ (categorical cross-entropy)
- Exact same recipe as Lecture 6, just more classes
- In **PyTorch**: `nn.CrossEntropyLoss` takes raw logits directly, just like `BCEWithLogitsLoss` did

The Gradient: Same Pattern, More Classes

- Just like the binary case $\frac{\partial L}{\partial z} = \hat{y} - y$, the softmax + cross-entropy gradient has a simple closed form:

$$\frac{\partial L}{\partial z_k} = \hat{y}_k - y_k \quad (\text{elementwise, } y \text{ one-hot})$$

i.e. $dL/dz = \text{softmax}(z) - \text{onehot}(y)$

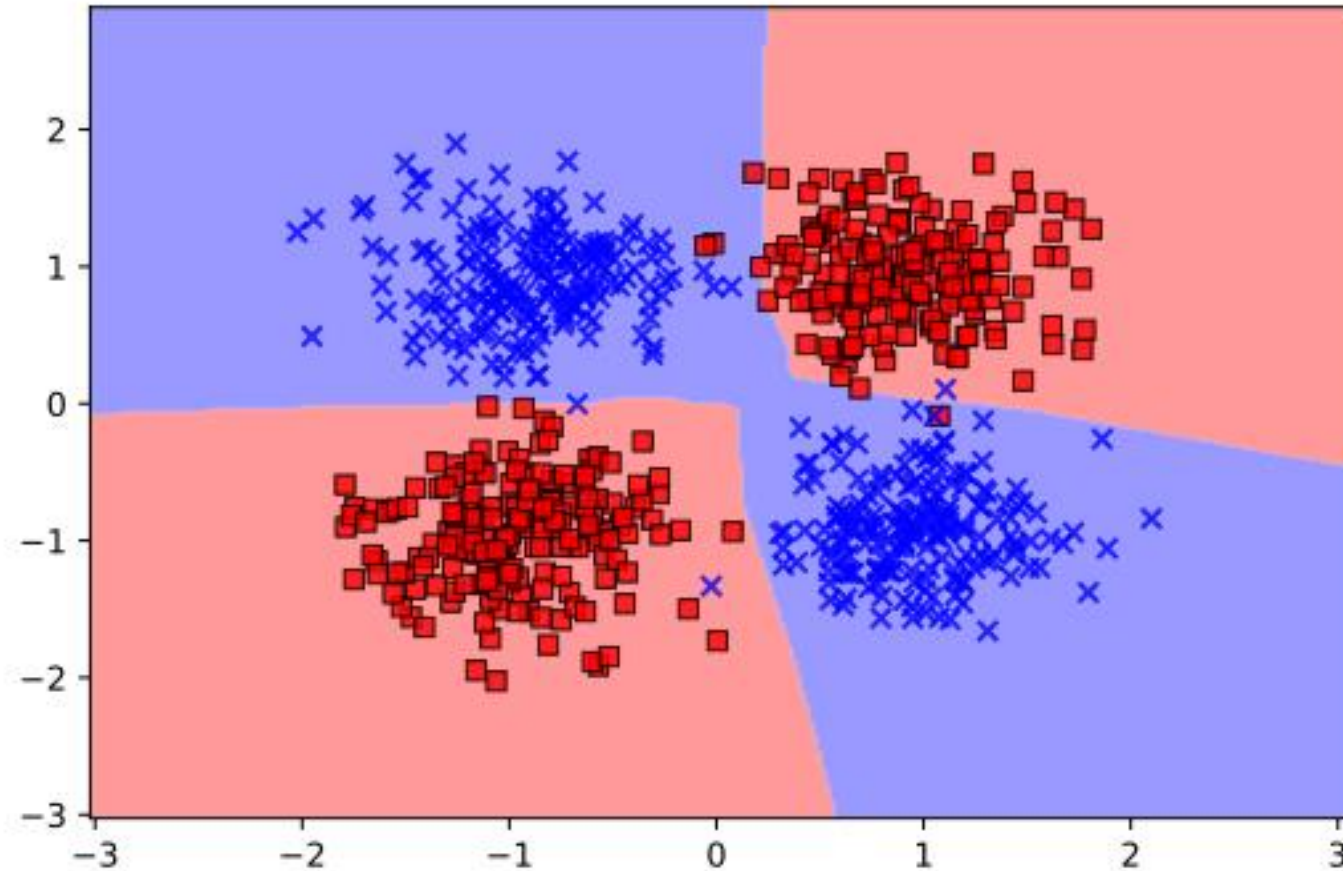
- Same derivation strategy as Lecture 6: chain rule through softmax's own derivative.
- Try it yourself with `torch` the way we checked the binary case last time.



Today: From Softmax to Multilayer Perceptrons

- From Binary to Multi-Class Classification
- **Multilayer Perceptron Architecture**
- Nonlinear Activation Functions
- Backpropagation
- Practical Considerations for Training MLPs

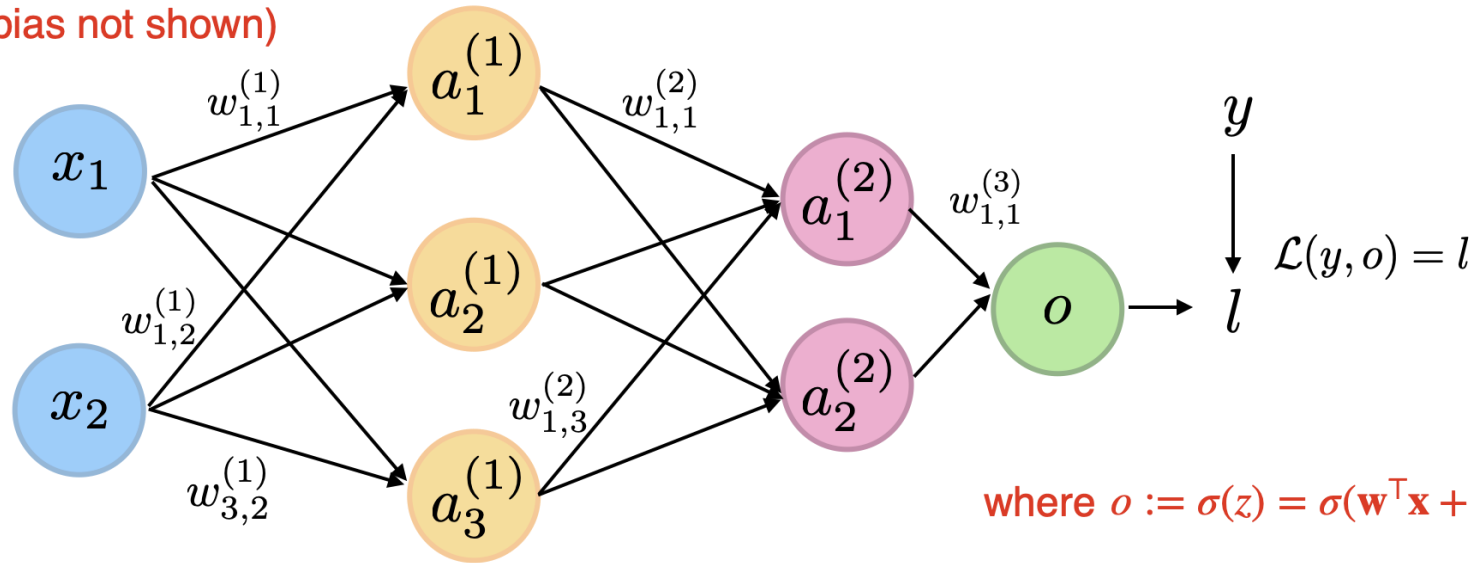
Today: We will finally be able to solve XOR



Multilayer Perceptron

- Computation Graph with Multiple Fully-Connected Layers

(bias not shown)



where $o := \sigma(z) = \sigma(\mathbf{w}^T \mathbf{x} + b)$

(Assume network for binary classification)

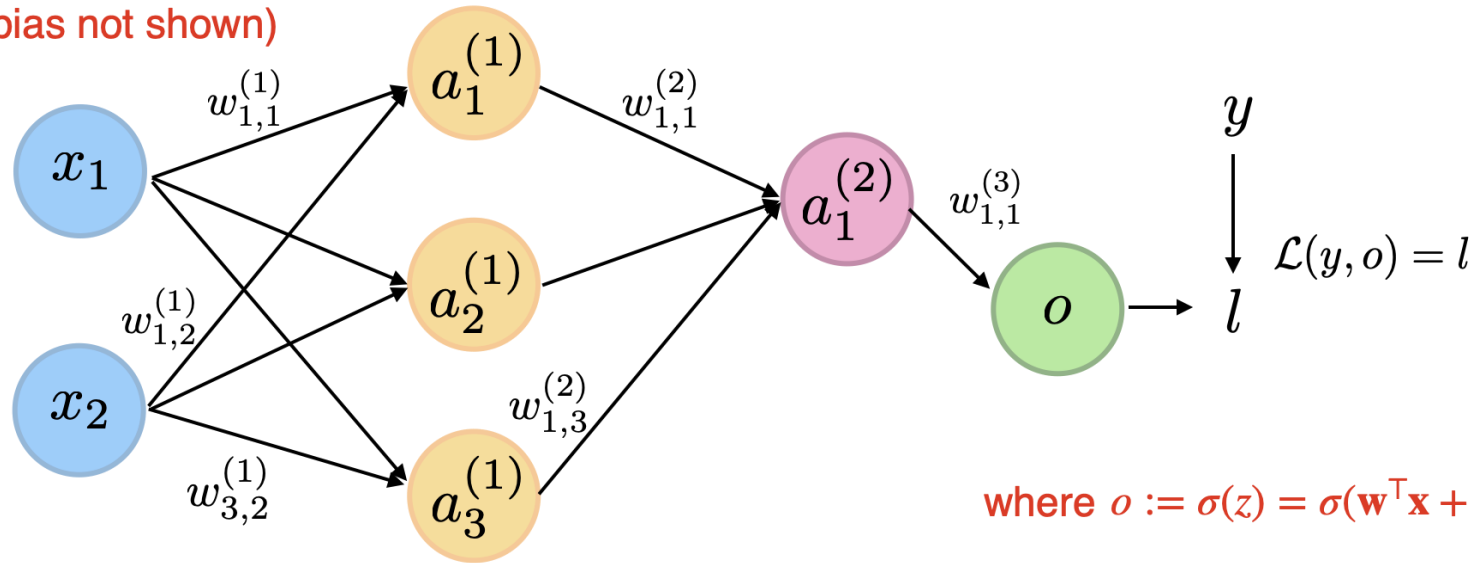
$$\begin{aligned} \frac{\partial l}{\partial w_{1,1}^{(1)}} &= \frac{\partial l}{\partial o} \cdot \frac{\partial o}{\partial a_1^{(2)}} \cdot \frac{\partial a_1^{(2)}}{\partial a_1^{(1)}} \cdot \frac{\partial a_1^{(1)}}{\partial w_{1,1}^{(1)}} \\ &+ \frac{\partial l}{\partial o} \cdot \frac{\partial o}{\partial a_2^{(2)}} \cdot \frac{\partial a_2^{(2)}}{\partial a_1^{(1)}} \cdot \frac{\partial a_1^{(1)}}{\partial w_{1,1}^{(1)}} \end{aligned}$$

Nothing new, really

Multilayer Perceptron

- Computation Graph with Multiple Fully-Connected Layers

(bias not shown)



where $o := \sigma(z) = \sigma(\mathbf{w}^T \mathbf{x} + b)$

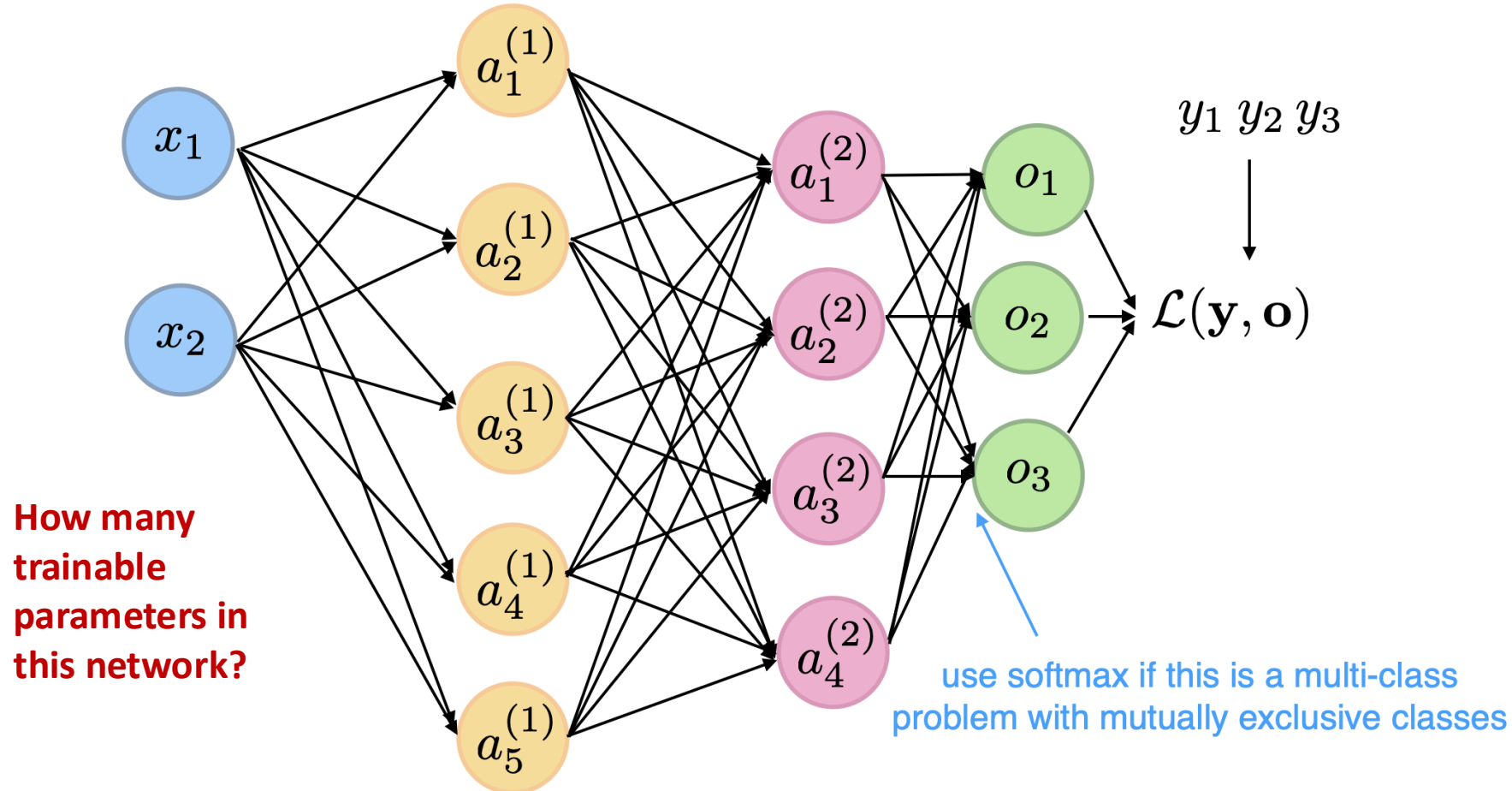
(Assume network for binary classification)

$$\begin{aligned} \frac{\partial l}{\partial w_{1,1}^{(1)}} &= \frac{\partial l}{\partial o} \cdot \frac{\partial o}{\partial a_1^{(2)}} \cdot \frac{\partial a_1^{(2)}}{\partial a_1^{(1)}} \cdot \frac{\partial a_1^{(1)}}{\partial w_{1,1}^{(1)}} \\ &\quad + \frac{\partial l}{\partial o} \cdot \frac{\partial o}{\partial a_2^{(2)}} \cdot \frac{\partial a_2^{(2)}}{\partial a_1^{(1)}} \cdot \frac{\partial a_1^{(1)}}{\partial w_{1,1}^{(1)}} \end{aligned}$$

Nothing new, really

Multilayer Perceptron

- Computation Graph with Multiple Fully-Connected Layers

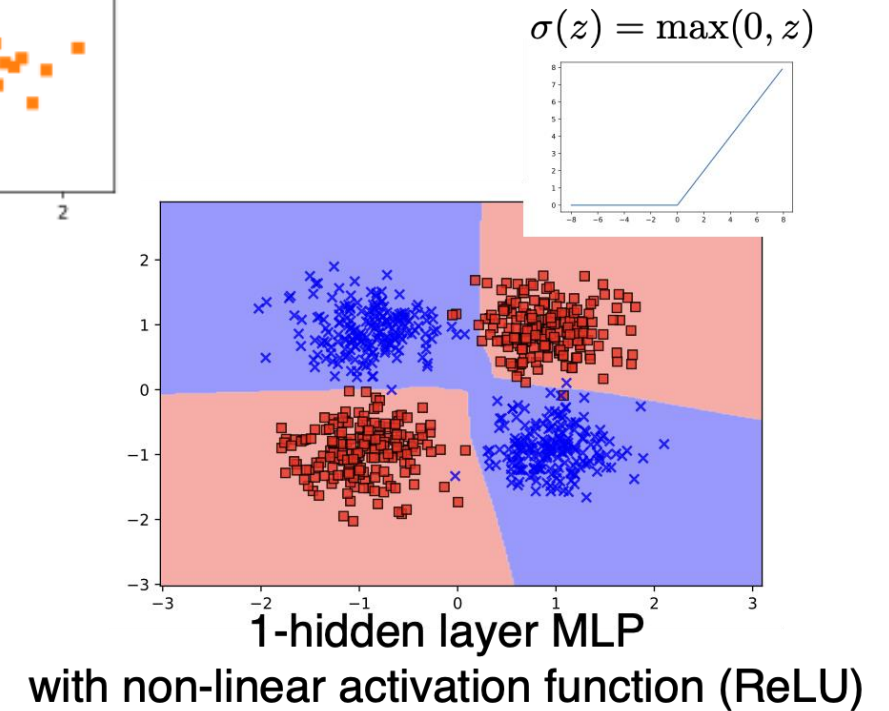
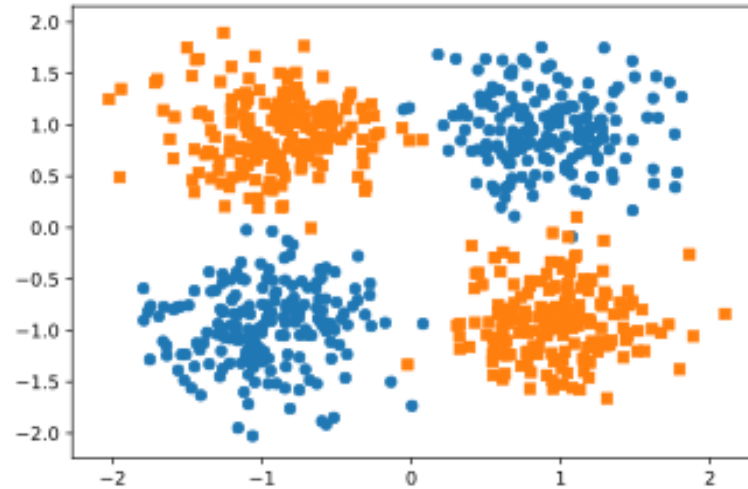
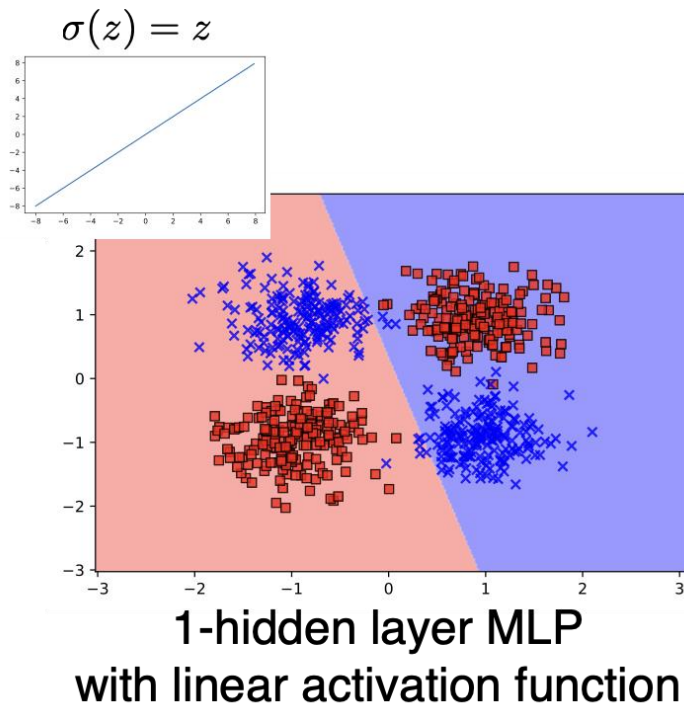




Today: From Softmax to Multilayer Perceptrons

- From Binary to Multi-Class Classification
- Multilayer Perceptron Architecture
- **Nonlinear Activation Functions**
- Backpropagation
- Practical Considerations for Training MLPs

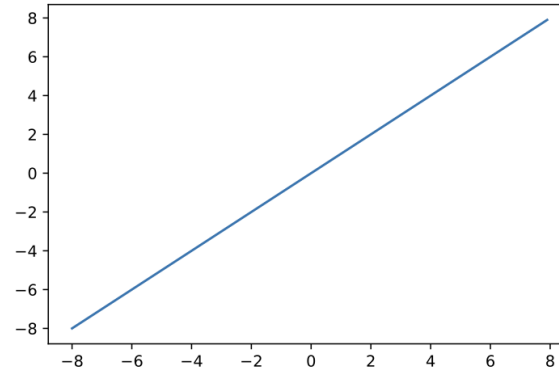
Solving the XOR Problem with Non-Linear Activations



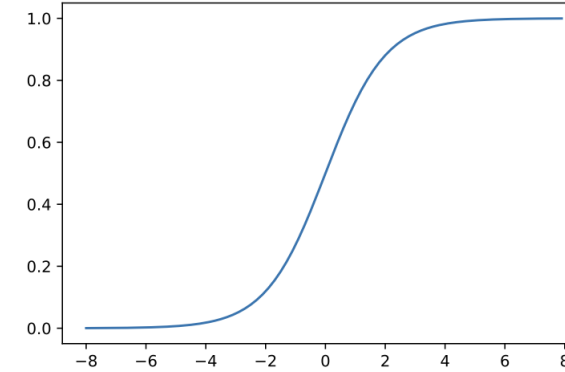
<https://github.com/rasbt/stat453-deep-learning-ss21/blob/master/L09/code/xor-problem.ipynb>

A Selection of Common Activation Functions

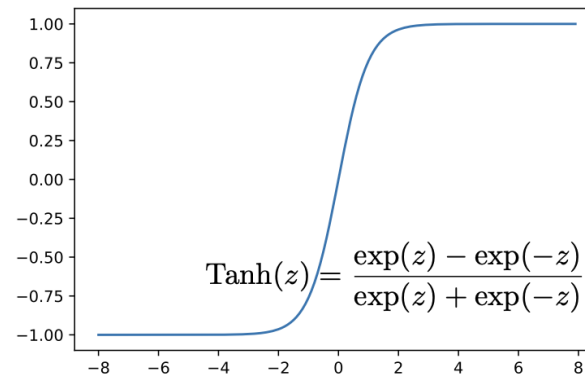
Identity



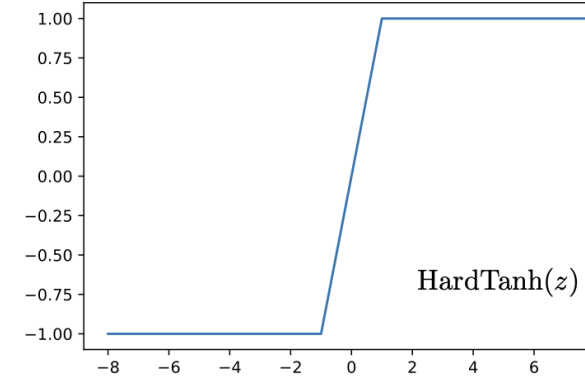
(Logistic) Sigmoid



Tanh ("tanH")



Hard Tanh

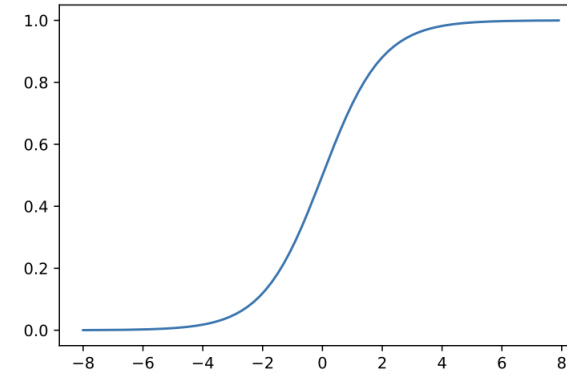


A Selection of Common Activation Functions

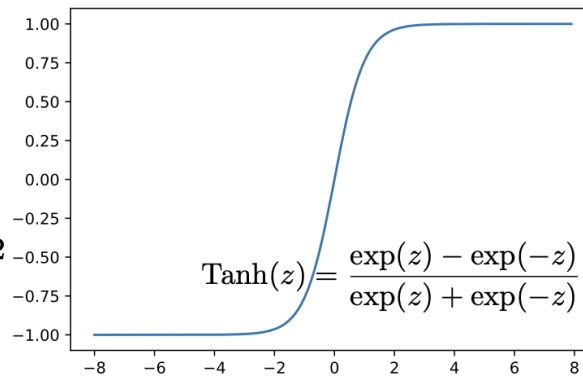
Advantages of Tanh

- Mean centering
- Positive and negative values
- Larger gradients

(Logistic) Sigmoid



Tanh ("tanH")



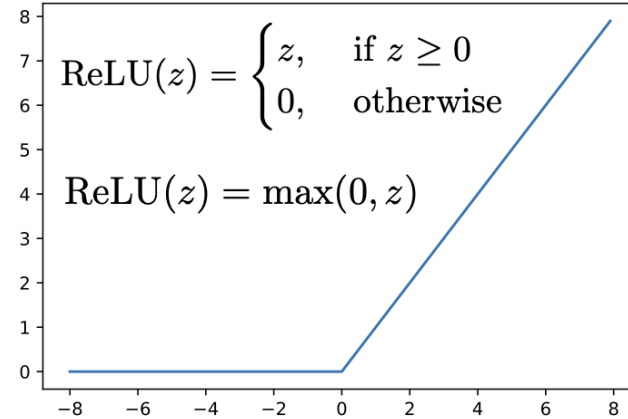
Also simple
derivative:

$$\frac{d}{dz} \text{Tanh}(z) = 1 - \text{Tanh}(z)^2$$

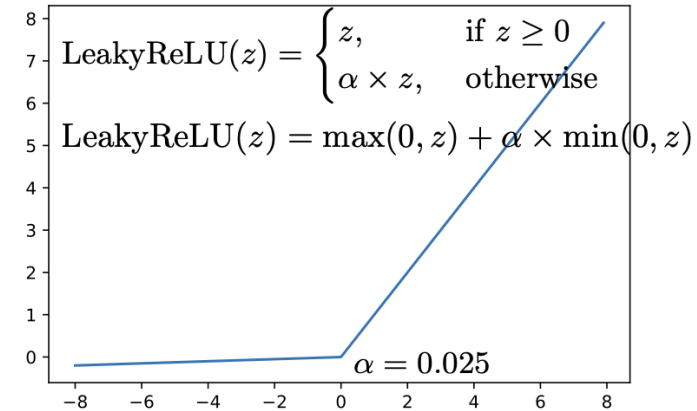
Important to normalize inputs to mean zero and use random weight initialization with **avg. weight centered at zero**

A Selection of Common Activation Functions (cont.)

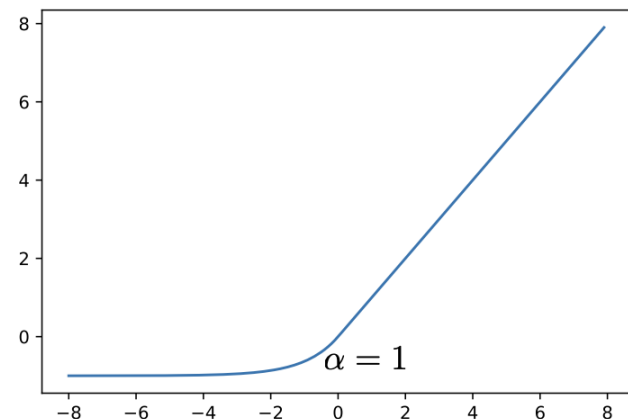
ReLU (Rectified Linear Unit)



Leaky ReLU



ELU (Exponential Linear Unit)



$$\text{ELU}(z) = \max(0, z) + \min(0, \alpha \times (\exp(z) - 1))$$

PReLU (Parameterized Rectified Linear Unit)

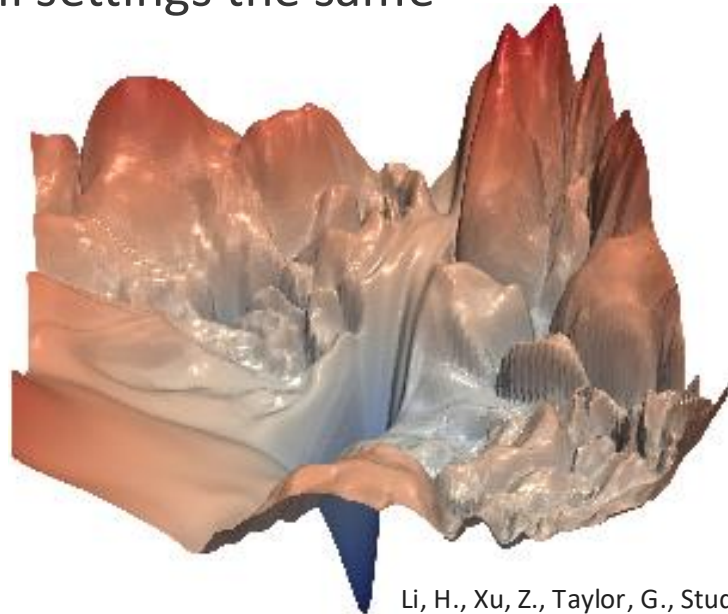
here, alpha is a trainable parameter

$$\text{PReLU}(z) = \begin{cases} z, & \text{if } z \geq 0 \\ \alpha z, & \text{otherwise} \end{cases}$$

$$\text{PReLU}(z) = \max(0, z) + \alpha \times \min(0, z)$$

Note that our Loss is Not Convex Anymore

- Linear regression, Adaline, Logistic Regression, and Softmax Regression have convex loss functions
- But our deep loss is no longer convex (most of the time)
 - In practice, we usually end up at different local minima if we repeat the training (e.g. by changing the random seed for weight initialization or shuffling the dataset while leaving all settings the same)



Li, H., Xu, Z., Taylor, G., Studer, C. and Goldstein, T., 2018. Visualizing the loss landscape of neural nets. In Advances in Neural Information Processing Systems (pp. 6391-6401).



But...

Convexity is Overrated

- **Using a suitable architecture (even if it leads to non-convex loss functions) is more important than insisting on convexity (particularly if it restricts us to unsuitable architectures)**
 - ▶ e.g.: Shallow (convex) classifiers versus Deep (non-convex) classifiers
- **Even for shallow/convex architecture, such as SVM, using non-convex loss functions actually improves the accuracy and speed**
 - ▶ See “trading convexity for efficiency” by Collobert, Bottou, and Weston, ICML 2006 (best paper award)

- Yann LeCun, [“Who’s afraid of Non-convex loss functions?”](#) – **2007**



Hmm...

What happens if we initialize the multilayer perceptron to all-zero weights?

Answer: The units stay symmetric forever.

This is exactly why we initialize randomly: it breaks the symmetry so different units can learn different features.



Today: From Softmax to Multilayer Perceptrons

- From Binary to Multi-Class Classification
- Multilayer Perceptron Architecture
- Nonlinear Activation Functions
- **Backpropagation**
- Practical Considerations for Training MLPs



Backprop: From L5's Sums to Matrix Form

Recall in Lecture 5, we discussed a fully-connected-layer gradient:
for every weight, you sum the chain rule over every path it reaches the loss through.

That sounds like matrix multiplication. Turns out it is!

Stack all those per-weight sums into vectors, and "sum over paths j " becomes "multiply by the transpose of the weight matrix":

$\text{delta_L} = \text{yhat} - y$	(output layer, Lecture 6)
$\text{delta_l} = (W_{(l+1)})^T \text{delta}_{(l+1)}$	(elementwise*) $\sigma'(z_l)$ (hidden layers)
$dL/dW_l = \text{delta_l} (a_{(l-1)})^T$	$dL/db_l = \text{delta_l}$

This is why code uses W , delta , and matmul instead of writing out each path by hand.

Backprop by the Numbers: Forward Pass on XOR

Let's define a network:

$x \rightarrow [\text{Linear}(2,2), \text{ReLU}] \rightarrow h \rightarrow [\text{Linear}(2,1), \text{sigmoid}]$
 $\rightarrow \hat{y} \rightarrow \text{BCE loss}$

$W1 = \begin{bmatrix} 1 & 1 \\ 1 & 1 \end{bmatrix}, \quad b1 = [0, -1]$
 $W2 = [1, -2], \quad b2 = -0.5$

$x = (1, 0), \quad y = 1 \quad (\text{XOR}(1,0) = 1)$

Forward pass:

$z1 = W_1 x + b1 = (1, 0)$

$h = \text{ReLU}(z1) = (1, 0) \quad \leftarrow \text{hidden unit 2 is "dead" for this input}$

$z2 = W_2 h + b2 = 1(1) + (-2)(0) - 0.5 = 0.5$

$\hat{y} = \text{sigmoid}(0.5) = 0.62, \quad L = -\log(\hat{y}) = 0.47$

Backprop by the Numbers: Backward Pass

$$\begin{aligned} dL/dz_2 &= \hat{y} - y = 0.62 - 1 = -0.38 && \text{(Lecture 6's formula)} \\ dL/dW_2 &= (dL/dz_2) h^T = (-0.38, 0) && dL/db_2 = -0.38 \end{aligned}$$

$$\begin{aligned} dL/dh &= (W_2)^T (dL/dz_2) = (-0.38, 0.75) && \text{(from prev. slide)} \\ dL/dz_1 &= (dL/dh) \text{ (elementwise*) } \text{ReLU}'(z_1) = (-0.38, 0) \\ dL/dW_1 &= (dL/dz_1) x^T = \begin{bmatrix} -0.38 & 0 \\ 0 & 0 \end{bmatrix} \\ dL/db_1 &= (-0.38, 0) \end{aligned}$$

Hidden unit 2's gradient is exactly zero — it was "dead" on this example.

Depth Makes Gradients Fragile

Notice, the gradient reaching layer 1 is a PRODUCT of every layer's local derivatives:

$$dL/dz_1 = (dL/dz_L) (dz_L/dz_{L-1}) \dots (dz_2/dz_1)$$

Each factor bakes in an activation derivative:

$$\sigma'(\text{sigmoid}) \leq 0.25, \quad \sigma'(\tanh) \leq 1, \quad \text{ReLU}' \in \{0, 1\}$$

Mostly-small factors $\rightarrow$ product shrinks exponentially with depth (vanishing).
Mostly-large factors (e.g. large weights) $\rightarrow$ product grows (exploding).

This is why Tanh's larger gradients matter. We will see more when we discuss Normalization & Initialization (10/6).



Today: From Softmax to Multilayer Perceptrons

- From Binary to Multi-Class Classification
- Multilayer Perceptron Architecture
- Nonlinear Activation Functions
- Backpropagation: Matrices, a Worked Example, and Gradient Flow
- **Practical Considerations for Training MLPs**

Diagnosing loss curves

- Checking the train/test loss/error is one simple diagnosis.
 - Tools such as Weights & Biases (<https://wandb.ai>) are helpful
- Test/val. error plateau or increases may suggest overfitting
 - Regularization / early stopping is needed (we will discuss next time)
- Train loss may be smoothed over mini-batches
- Some models such as Transformers may be more difficult to train, sharp transition followed by long plateau, e.g. “grokking” (<https://arxiv.org/abs/2201.02177>)
- Learning rate (LR), LR scheduler, initialization all effect training curves. Lots of trial and error.
 - “Grad student descent”



Parameters vs Hyperparameters

weights (weight parameters)
biases (bias units)

minibatch size
data normalization schemes
number of epochs
number of hidden layers
number of hidden units
learning rates
(random seed, why?)
loss function
various weights (weighting terms)
activation function types
regularization schemes (more later)
weight initialization schemes (more later)
optimization algorithm type (more later)
...

Wide vs Deep Architectures

- MLPs with one (large) hidden layer are universal functions approximators [1-4]
- So why do we want to use deeper architectures?

[1] Balázs Csanád Csáji (2001) Approximation with Artificial Neural Networks; Faculty of Sciences; Eötvös Loránd University, Hungary

[2] Barron, Andrew R. "Universal approximation bounds for superpositions of a sigmoidal function." IEEE Transactions on Information theory 39.3 (1993): 930-945.

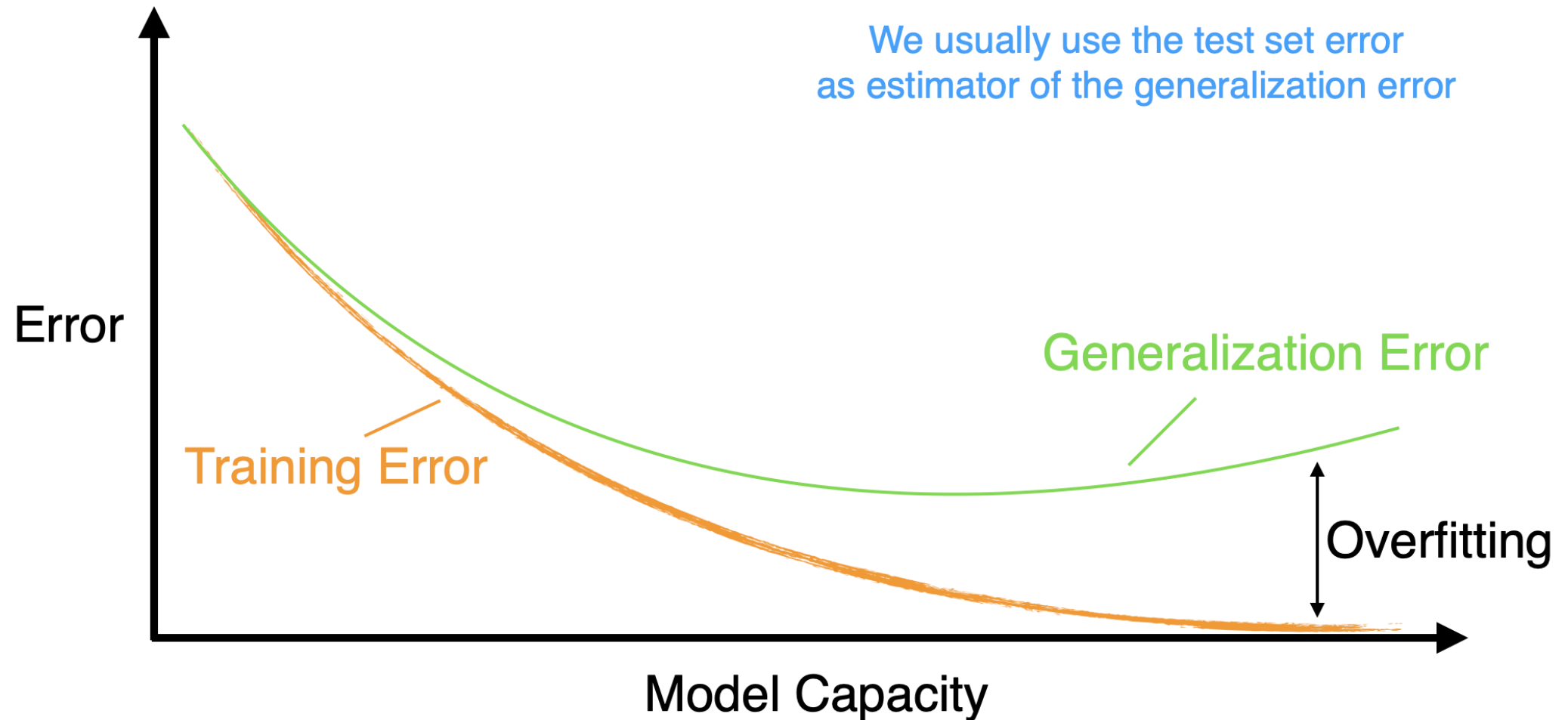
[3] Cybenko, G. (1989) "Approximations by superpositions of sigmoidal functions", Mathematics of Control, Signals, and Systems, 2(4), 303–314. doi:10.1007/BF02551274

[4] Hornik, K., Stinchcombe, M., & White, H. (1989). Multilayer feedforward networks are universal approximators. Neural networks, 2(5), 359-366.

Wide vs Deep Architectures

- MLPs with one (large) hidden layer are universal functions approximators [1-4]
- So why do we want to use deeper architectures?
 - Can achieve the same expressiveness with more layers but fewer parameters (combinatorics); fewer parameters => less overfitting?, faster computation
 - Also, having more layers provides some form of regularization: later layers are constrained on the behavior of earlier layers
 - However, more layers => vanishing/exploding gradients
 - Later: different layers for different levels of feature abstraction (DL is really more about feature learning than just stacking multiple layers)
 - Inductive bias? [<https://arxiv.org/abs/1608.08225>]

Overfitting and Underfitting



Good Practices for Training MLPs

Multilayer Perceptron with Sigmoid
Activation and MSE Loss

https://github.com/rasbt/stat453-deep-learning-ss20/blob/master/L08-mlp/code/mlp-pytorch_sigmoid-mse.ipynb

Multilayer Perceptron with Softmax
Activation and Cross-Entropy Loss

https://github.com/rasbt/stat453-deep-learning-ss20/blob/master/L08-mlp/code/mlp-pytorch_sigmoid-crossentr.ipynb

- Set a seed for reproducibility
- Check that your data/shapes look as expected before training
- Log intermediate values (loss, accuracy) and time your epochs
- If loss behaves strangely (e.g. increases early on), stop and debug before running longer
- Change one hyperparameter at a time and repeat experiments
- See Lecture 6's notebook for a worked PyTorch example of the `zero_grad` → `backward` → `step` loop



From Dense to Structured: Setting Up CNNs

A conv layer is the SAME equation as a fully-connected one: $z = Wx + b$.
The only difference is W is constrained – sparse (local receptive field) and tied (the same filter reused at every position).

So the backward pass doesn't change either: $\text{delta} = W^T \text{delta_next} (\text{elementwise}^*) \sigma'(z)$

What changes: a shared weight now gets a gradient contribution from every position it was used at, so you SUM them (Lecture 5's weight-sharing slide).

Teaser for next time: backprop through a convolution turns out to be a convolution too, with the kernel flipped. We'll prove this in Lecture 8.

Next time (9/29): we prove the flipped-kernel claim above, concretely — Structure and weight sharing, CNNs.

Questions?

