

Introduction to Neural Networks

November 17th 2020
KMI-2020
Nagoya University



Kazuhiro Terao
SLAC National Lab

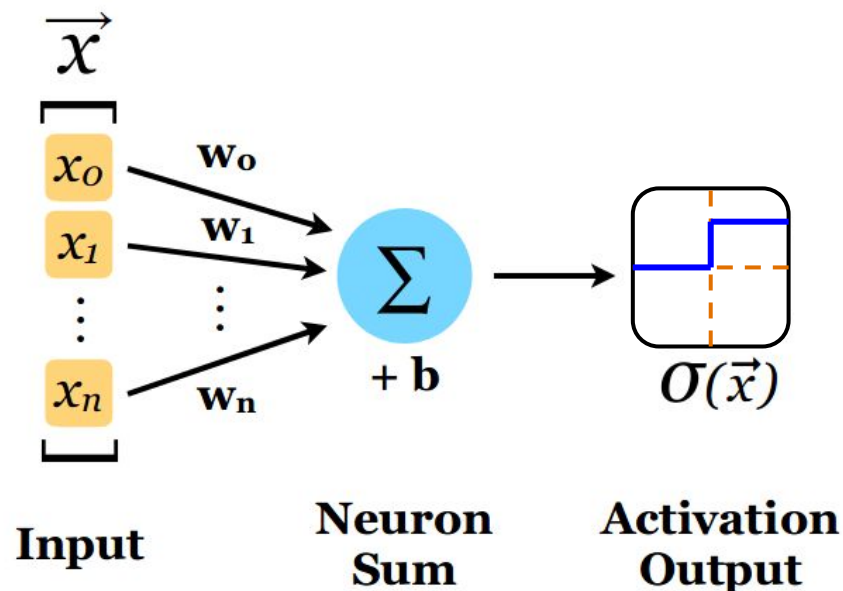


Introduction to Neural Networks

Perceptron

The basic unit of a neural net is the *perceptron* (loosely based on a real neuron)

Takes in a vector of inputs (x). Commonly inputs are summed with weights (w) and offset (b) then run through activation.

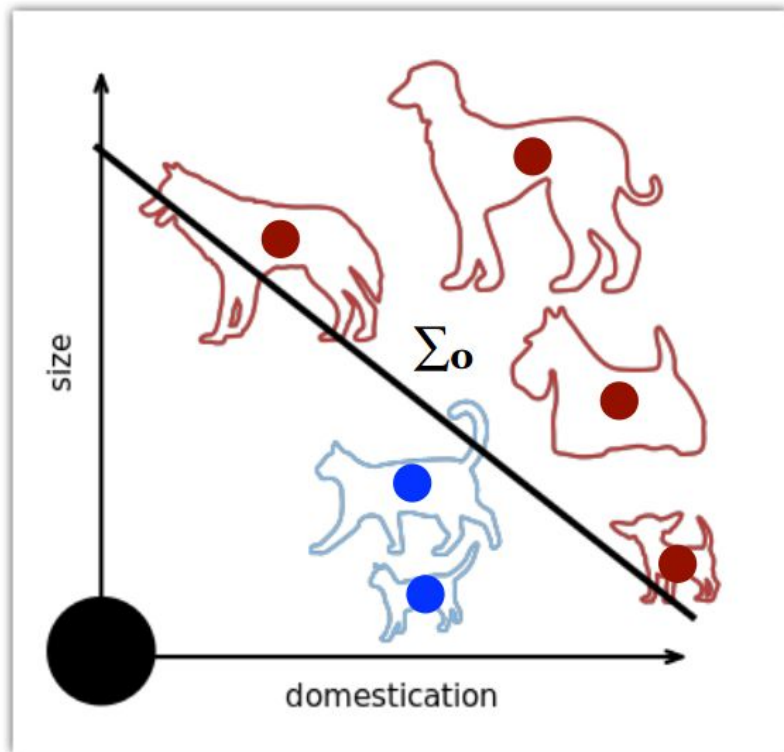


$$\sigma(\vec{x}) = \begin{cases} 1 & \vec{w}_i \cdot \vec{x} + b_i \geq 0 \\ 0 & \vec{w}_i \cdot \vec{x} + b_i < 0. \end{cases}$$

Introduction to Neural Networks

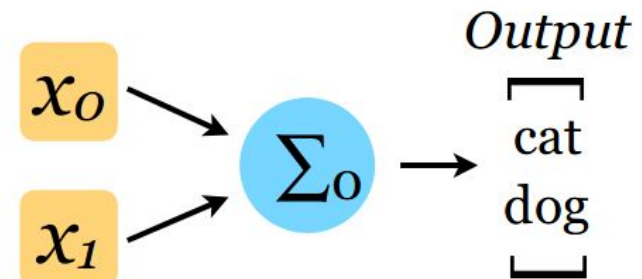
Perceptron

Imagine using two features to separate cats and dogs



from [wikipedia](https://en.wikipedia.org/wiki/Perceptron)

$$\sigma(\vec{x}) = \begin{cases} 1 & \vec{w}_i \cdot \vec{x} + b_i \geq 0 \\ 0 & \vec{w}_i \cdot \vec{x} + b_i < 0. \end{cases}$$

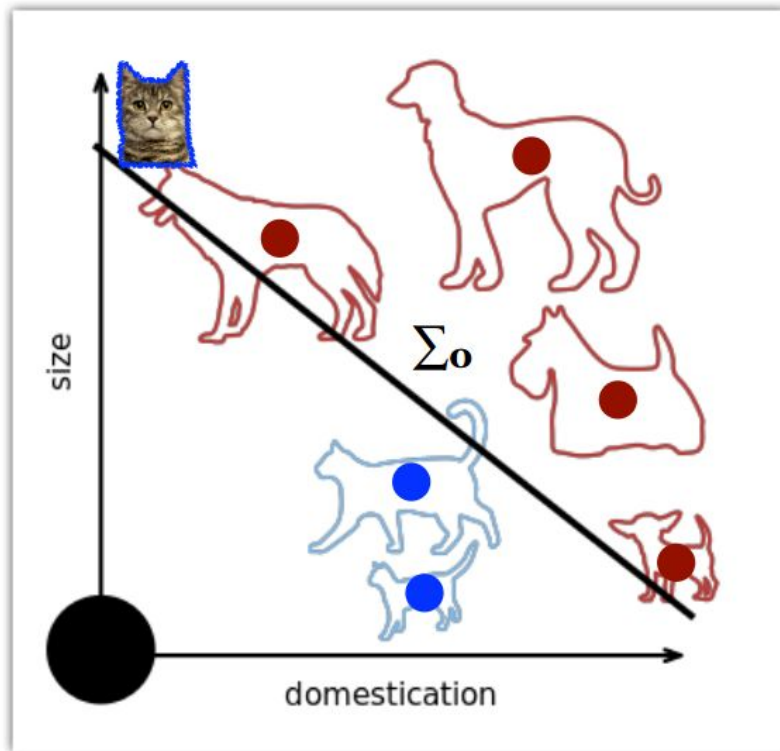


By picking a value for $\mathbf{w}$ and $\mathbf{b}$,
we define a boundary
between the two sets of data

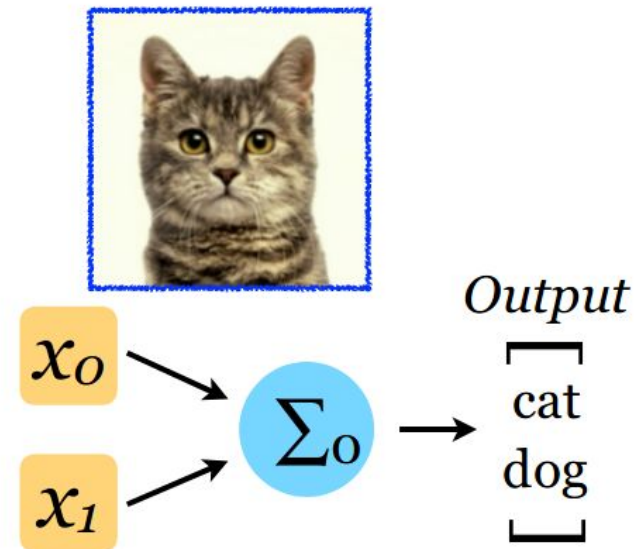
Introduction to Neural Networks

Perceptron

What if we have a new data point?



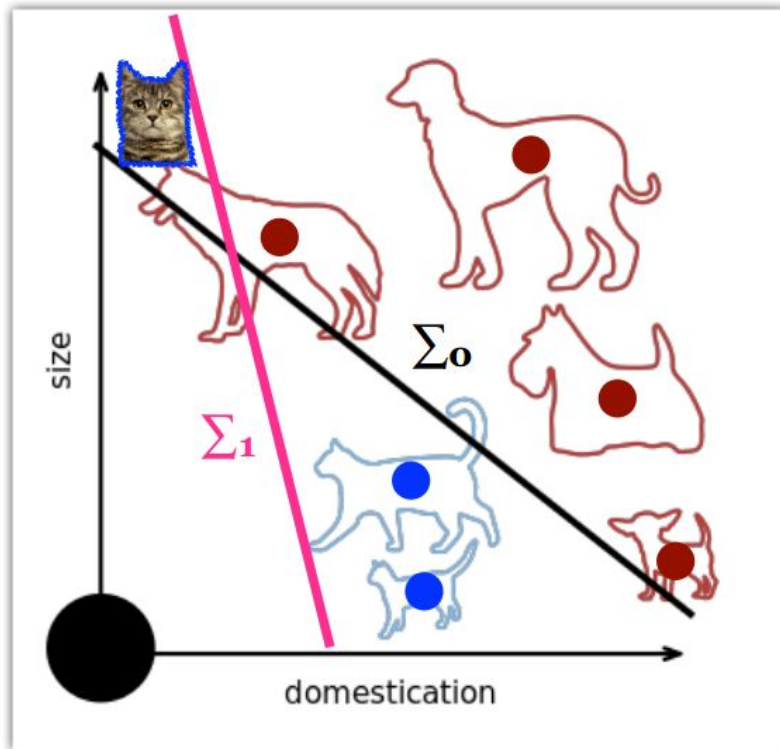
from [wikipedia](https://en.wikipedia.org/wiki/Perceptron)



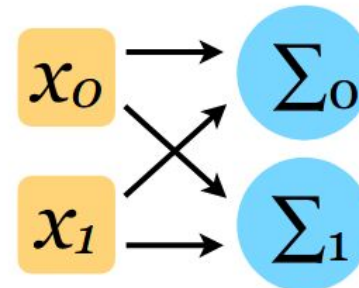
Introduction to Neural Networks

Perceptron

What if we have a new data point?



from [wikipedia](https://en.wikipedia.org/wiki/Perceptron)

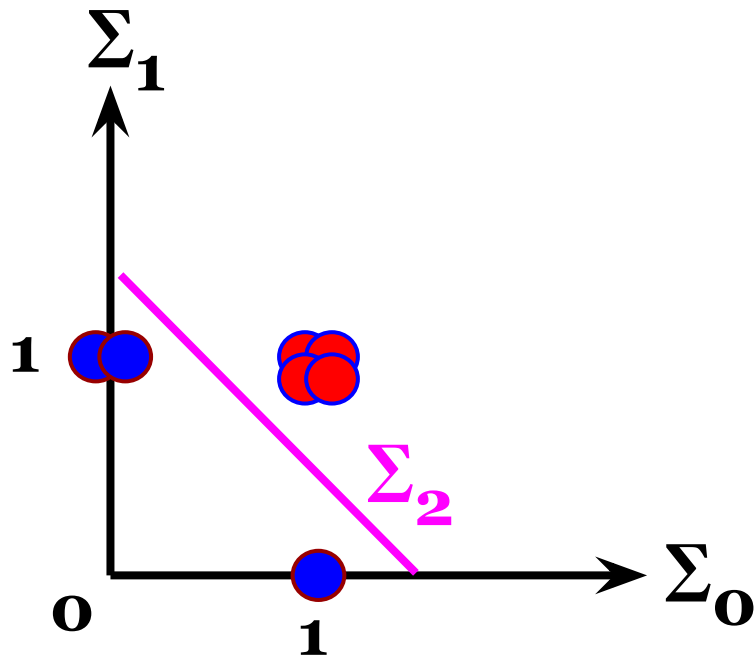


We can add **another neuron** to help (but does not yet solve the problem)

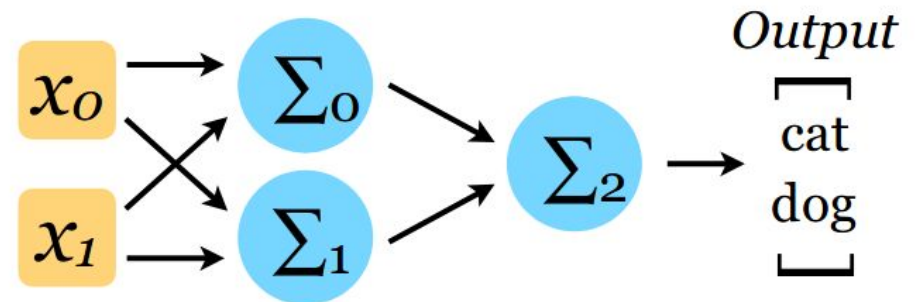
Introduction to Neural Networks

Multi-Layer Perceptron

What if we have a new data point?



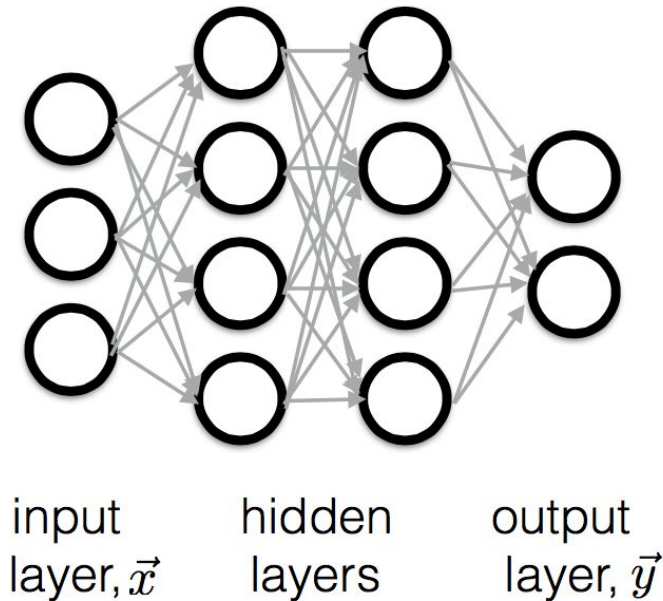
from [wikipedia](https://en.wikipedia.org/wiki/Linear_classifier)



Another layer can classify
based on preceding layer's output
(of **non-linear activation**)

“Traditional neural net”

Fully-Connected Multi-Layer Perceptrons



A traditional neural network consists of a stack of layers of such neurons where each neuron is **fully connected** to other neurons of the neighbor layers

Training Feed-forward Neural Networks Gradient-based Method

Introduction to Neural Networks

Training (Optimization)

N-layer MLP

x_0 ... input data

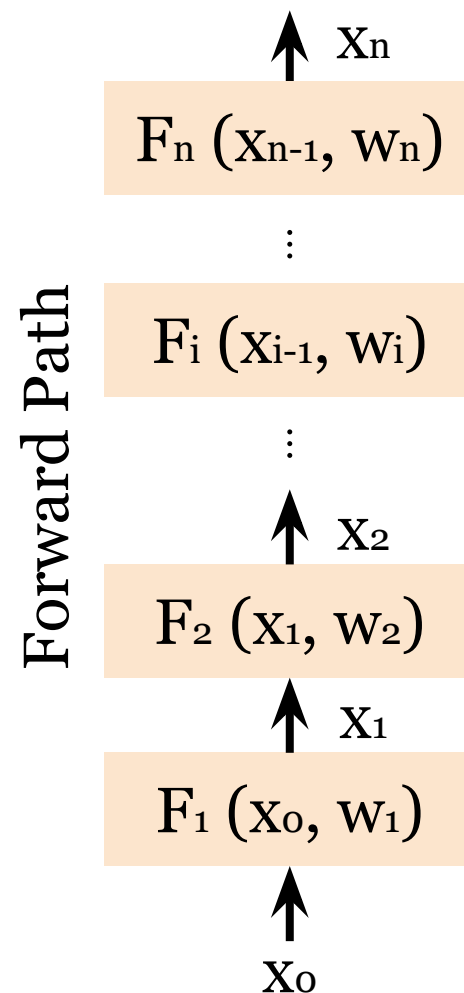
F_i ... i -th hidden layer

x_i ... i -th layer output

w_i ... i -th layer weights

y ... correct answer

$\mathcal{L}$... loss (cost) function



Introduction to Neural Networks

Training (Optimization)

Loss: a cost to minimize (e.g. error).

N-layer MLP

x_0 ... input data

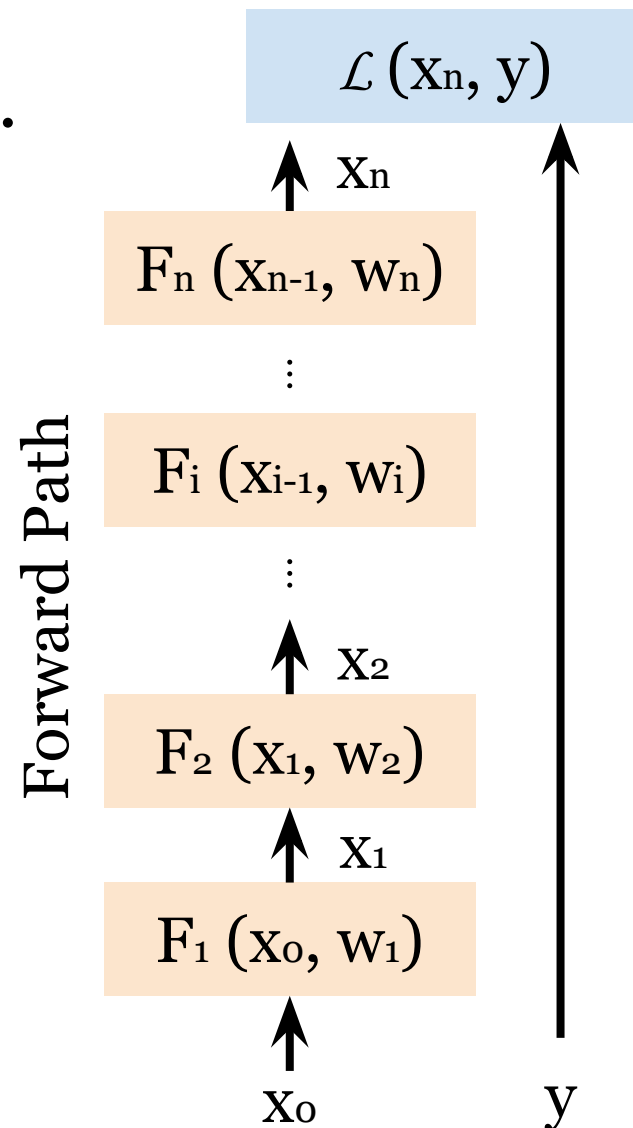
F_i ... i -th hidden layer

x_i ... i -th layer output

w_i ... i -th layer weights

y ... correct answer

$\mathcal{L}$... loss (cost) function

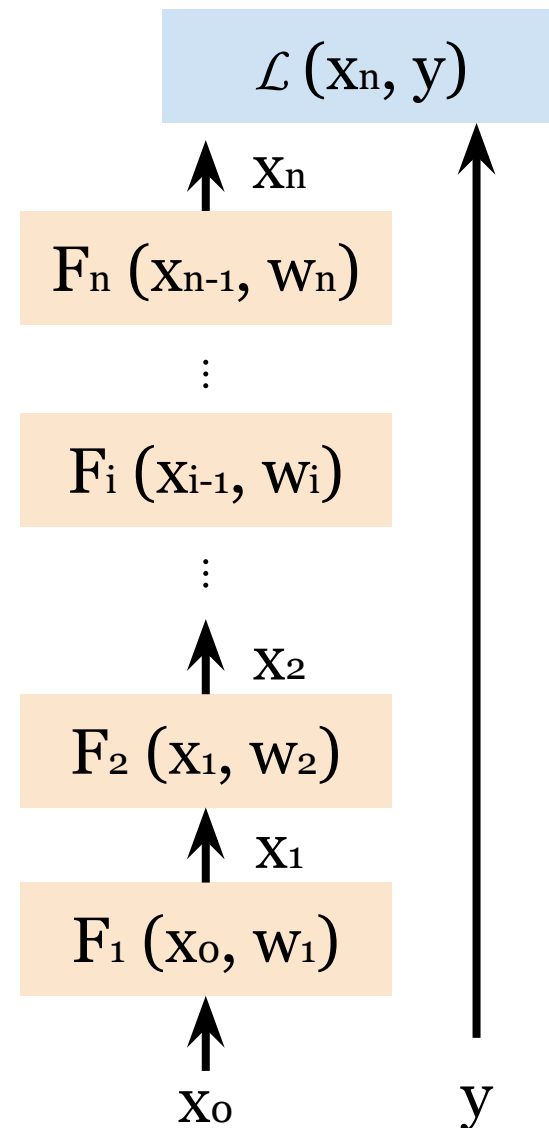


Introduction to Neural Networks

Training (Optimization)

Loss: a cost to minimize (e.g. error).

Training: an optimization of parameters by minimizing loss.



Introduction to Neural Networks

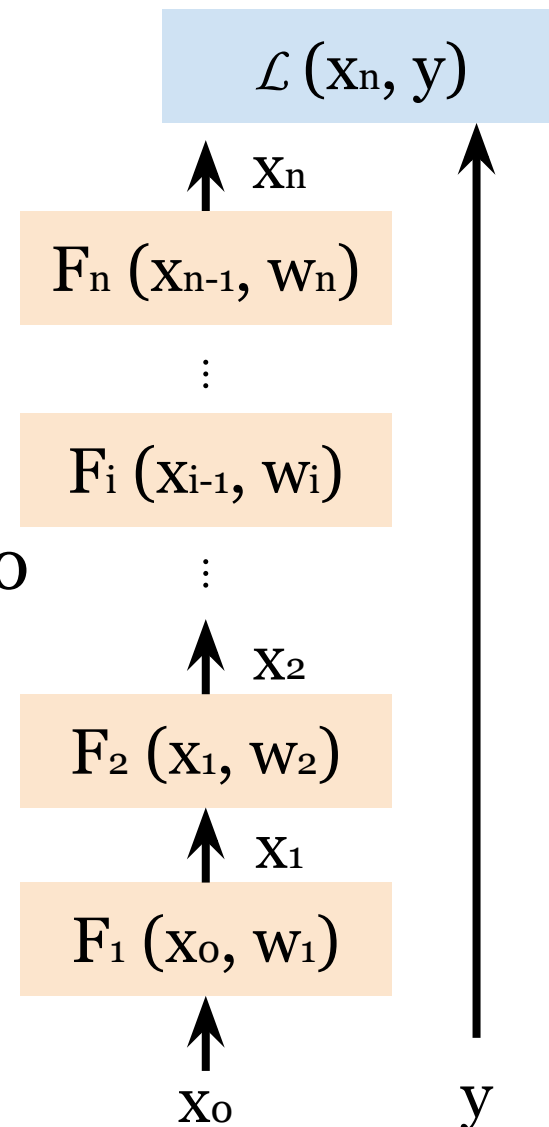
Training (Optimization)

Loss: a cost to minimize (e.g. error).

Training: an optimization of parameters by minimizing loss.

Gradient Descend: an iterative approach to optimize weights in order to minimize the loss.

- Define & measure the error
 - Compute $\partial \mathcal{L} / \partial \mathbf{w}_n = \nabla \mathbf{W}$
 - Update weights: $\mathbf{W}_{\text{new}} = \mathbf{W} - \lambda \nabla \mathbf{W}$
- ... where λ is called a “learning rate”



Introduction to Neural Networks

Backpropagation: how to compute gradients

How can we compute the gradient for the i -th layer?

$$\mathcal{L}(\mathbf{x}_n, y) = \mathcal{L}(\mathbf{F}_n(\mathbf{x}_{n-1}, \mathbf{w}_n), y) = \dots = \mathcal{L}(\mathbf{F}_n(\dots \mathbf{F}_i(\mathbf{x}_{i-1}, \mathbf{w}_i) \dots), y)$$

Introduction to Neural Networks

Backpropagation: how to compute gradients

How can we compute the gradient for the i -th layer?

$$\mathcal{L}(\mathbf{x}_n, y) = \mathcal{L}(\mathbf{F}_n(\mathbf{x}_{n-1}, \mathbf{w}_n), y) = \dots = \mathcal{L}(\mathbf{F}_n(\dots \mathbf{F}_i(\mathbf{x}_{i-1}, \mathbf{w}_i) \dots), y)$$

$$\frac{\partial \mathcal{L}}{\partial \mathbf{w}_i} = \frac{\partial \mathcal{L}}{\partial \mathbf{x}_n} \cdot \frac{\partial \mathbf{x}_n}{\partial \mathbf{x}_{n-1}} \cdot \frac{\partial \mathbf{x}_{n-1}}{\partial \mathbf{x}_{n-2}} \dots \frac{\partial \mathbf{x}_{i+1}}{\partial \mathbf{x}_i} \cdot \frac{\partial \mathbf{x}_i}{\partial \mathbf{w}_i} \quad (\text{chain rule})$$

Introduction to Neural Networks

Backpropagation: how to compute gradients

How can we compute the gradient for the i -th layer?

$$\mathcal{L}(\mathbf{x}_n, y) = \mathcal{L}(\mathbf{F}_n(\mathbf{x}_{n-1}, \mathbf{w}_n), y) = \dots = \mathcal{L}(\mathbf{F}_n(\dots \mathbf{F}_i(\mathbf{x}_{i-1}, \mathbf{w}_i) \dots), y)$$

$$\frac{\partial \mathcal{L}}{\partial \mathbf{w}_i} = \boxed{\frac{\partial \mathcal{L}}{\partial \mathbf{x}_n}} \cdot \frac{\partial \mathbf{x}_n}{\partial \mathbf{x}_{n-1}} \cdot \frac{\partial \mathbf{x}_{n-1}}{\partial \mathbf{x}_{n-2}} \dots \frac{\partial \mathbf{x}_{i+1}}{\partial \mathbf{x}_i} \cdot \frac{\partial \mathbf{x}_i}{\partial \mathbf{w}_i} \quad (\text{chain rule})$$

Example: quadratic loss $= \sum_{m=1}^m (\mathbf{x}_n - y)^2$

$$\frac{\partial \mathcal{L}}{\partial \mathbf{x}_n} = \sum_{m=1}^m (2\mathbf{x}_n - 2y)$$

... or $\frac{\partial \mathcal{L}}{\partial \mathbf{x}_n} = \sum_{m=1}^m (2\mathbf{x}_n - 2y)^T$ if $\mathbf{x}_n$ and y are column vector

Introduction to Neural Networks

Backpropagation: how to compute gradients

How can we compute the gradient for the i -th layer?

$$\mathcal{L}(\mathbf{x}_n, y) = \mathcal{L}(\mathbf{F}_n(\mathbf{x}_{n-1}, \mathbf{w}_n), y) = \dots = \mathcal{L}(\mathbf{F}_n(\dots \mathbf{F}_i(\mathbf{x}_{i-1}, \mathbf{w}_i) \dots), y)$$

$$\frac{\partial \mathcal{L}}{\partial \mathbf{w}_i} = \boxed{\frac{\partial \mathcal{L}}{\partial \mathbf{x}_n} \cdot \frac{\partial \mathbf{x}_n}{\partial \mathbf{x}_{n-1}} \cdot \frac{\partial \mathbf{x}_{n-1}}{\partial \mathbf{x}_{n-2}} \dots \frac{\partial \mathbf{x}_{i+1}}{\partial \mathbf{x}_i}} \cdot \frac{\partial \mathbf{x}_i}{\partial \mathbf{w}_i} \quad (\text{chain rule})$$

$$= \frac{\partial \mathcal{L}}{\partial \mathbf{x}_i} \dots \text{can compute iteratively} \\ (\text{"Back" propagation})$$

If we know $\frac{\partial \mathcal{L}}{\partial \mathbf{x}_i}$ (gradient at i -th layer), then we can calculate:

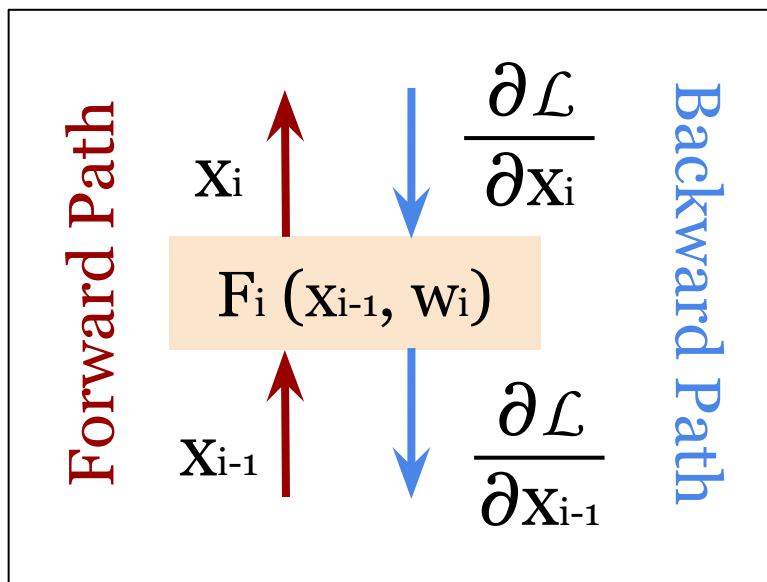
$$\frac{\partial \mathcal{L}}{\partial \mathbf{x}_{i-1}} = \frac{\partial \mathcal{L}}{\partial \mathbf{x}_i} \cdot \frac{\partial \mathbf{x}_i}{\partial \mathbf{x}_{i-1}} = \frac{\partial \mathcal{L}}{\partial \mathbf{x}_i} \cdot \frac{\partial \mathbf{F}_i(\mathbf{x}_{i-1}, \mathbf{w}_i)}{\partial \mathbf{x}_{i-1}}$$

Introduction to Neural Networks

Backpropagation: how to compute gradients

How can we compute the gradient for the i -th layer?

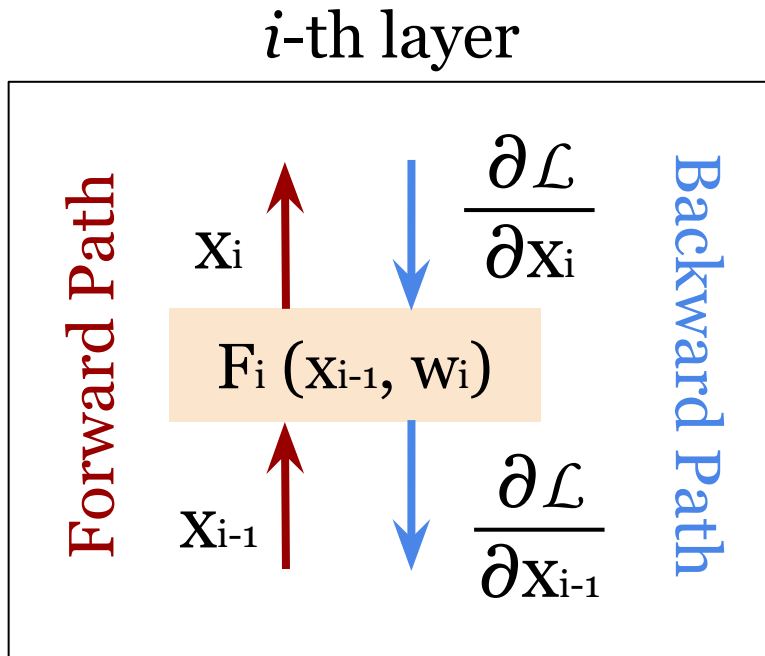
i -th layer



Introduction to Neural Networks

Backpropagation: how to compute gradients

How can we compute the gradient for the i -th layer?



Input: $x_{i-1} \quad \frac{\partial \mathcal{L}}{\partial x_i}$ (given)

Output (we compute):

$$x_i = F_i(x_{i-1}, w_i)$$

$$\frac{\partial \mathcal{L}}{\partial x_{i-1}} = \frac{\partial \mathcal{L}}{\partial x_i} \cdot \frac{\partial F_i(x_{i-1}, w_i)}{\partial x_{i-1}}$$

Weight Update: $w_i - \lambda \nabla w_i$

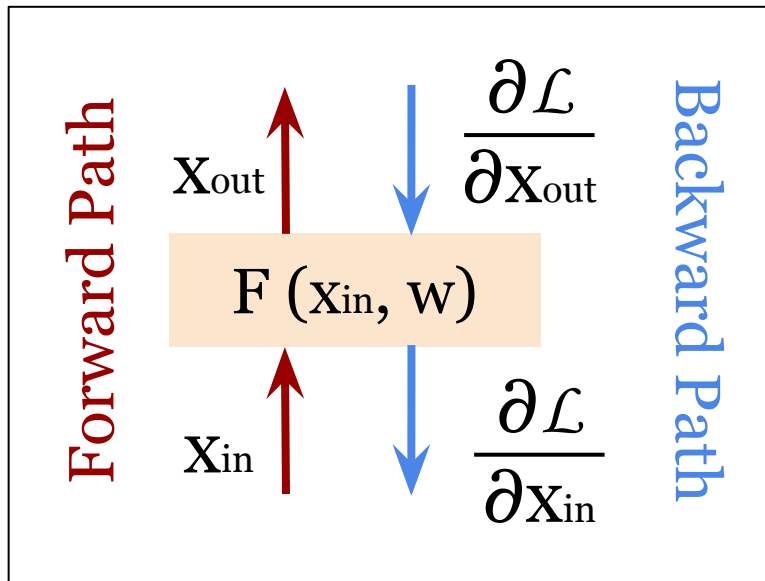
where: $\nabla w_i = \frac{\partial \mathcal{L}}{\partial w_i} = \frac{\partial \mathcal{L}}{\partial x_i} \cdot \frac{\partial F_i(x_{i-1}, w_i)}{\partial w_i}$

Introduction to Neural Networks

Backpropagation: how to compute gradients

How can we compute the gradient for the i -th layer?

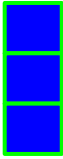
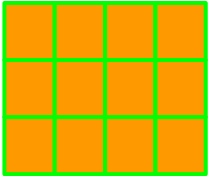
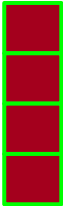
i -th layer



Example: linear transformation

$$F(X_{in}, w) = w \cdot X_{in}$$

Let's generalize for a tensor input

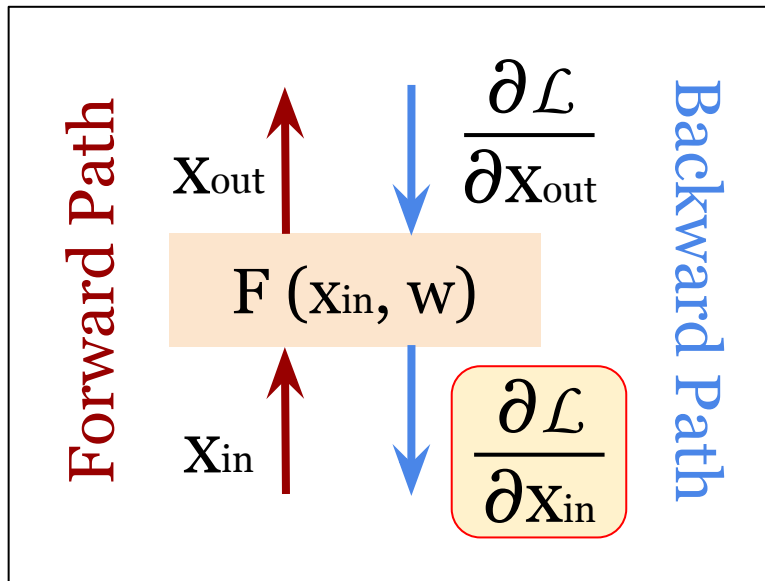
e.g.)  =  · 

Introduction to Neural Networks

Backpropagation: how to compute gradients

How can we compute the gradient for the i -th layer?

i -th layer

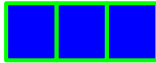
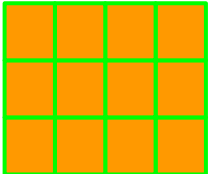


Example: linear transformation

$$F(X_{in}, w) = w \cdot X_{in}$$

Let's generalize for a tensor input

$$\frac{\partial \mathcal{L}}{\partial X_{in}} = \frac{\partial \mathcal{L}}{\partial X_{out}} \cdot w$$

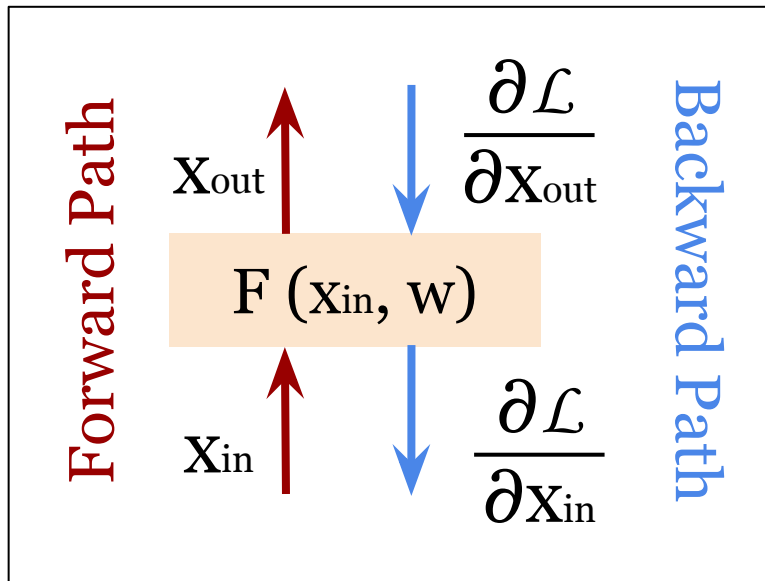
e.g.)  =  · 

Introduction to Neural Networks

Backpropagation: how to compute gradients

How can we compute the gradient for the i -th layer?

i -th layer



Example: linear transformation

$$F(X_{in}, w) = w \cdot X_{in}$$

Let's generalize for a tensor input

$$\frac{\partial \mathcal{L}}{\partial X_{in}} = \frac{\partial \mathcal{L}}{\partial X_{out}} \cdot w$$

$$\frac{\partial \mathcal{L}}{\partial w} = X_{in} \cdot \frac{\partial \mathcal{L}}{\partial X_{out}}$$

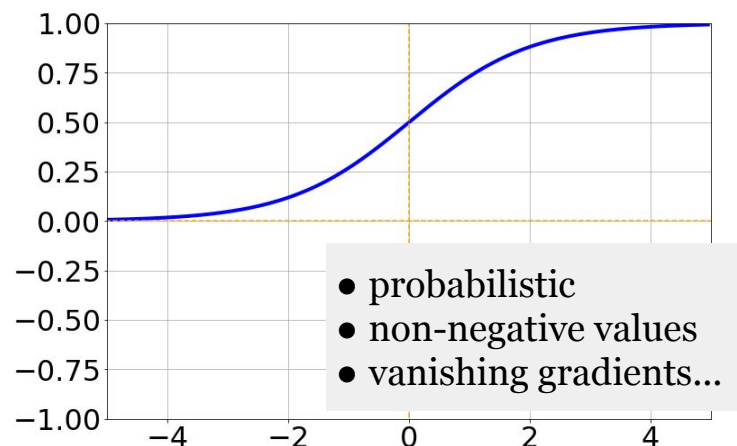
$$w_{new} = w - \lambda \left(\frac{\partial \mathcal{L}}{\partial w} \right)^T$$

A diagram illustrating the matrix multiplication $\frac{\partial \mathcal{L}}{\partial w} = X_{in} \cdot \frac{\partial \mathcal{L}}{\partial X_{out}}$. On the left is a 4x4 orange matrix. To its right is an equals sign, followed by a 4x1 blue column vector, a dot, and a 1x4 red row vector. All three matrices have green borders.

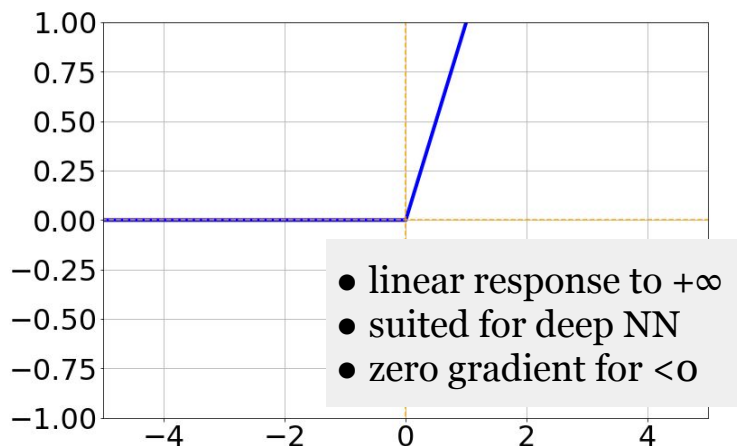
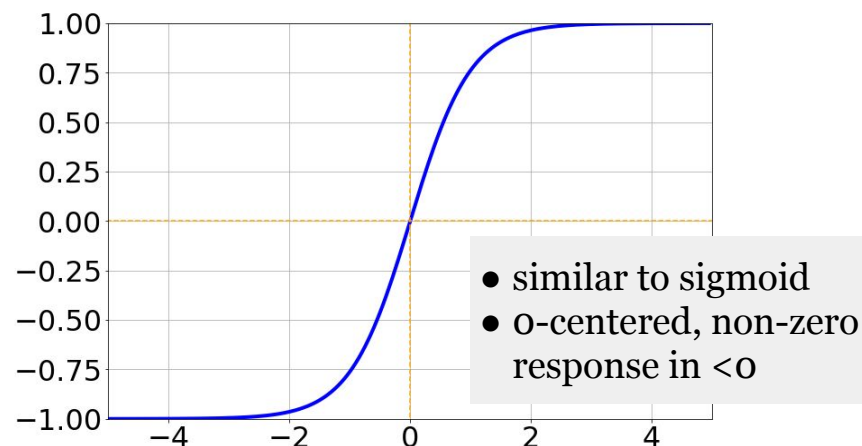
Introduction to Neural Networks

Activation Functions

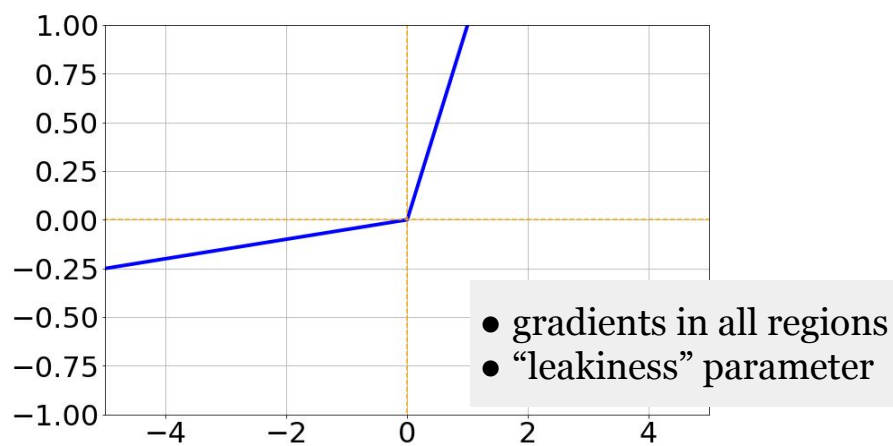
Sigmoid



Tanh



ReLU



Leaky ReLU

Weights Initialization

- **Random values** to set different neuron states
- **Not too large or too small:** gradients become negligible for *tanh* and *sigmoid* activation layers! (positive large value not an issue for ReLU)
- **Gaussian @ $(\mu, \sigma) = (0, 1)$**
 - Slight modification: normalize such that the variance of a signal strength stays similar (0~1) across layers (depends on # of filters)
 - See more details [here \(CS231\)](#)
 - [Xavier initialization](#), [He \(MSRA\) initialization](#)

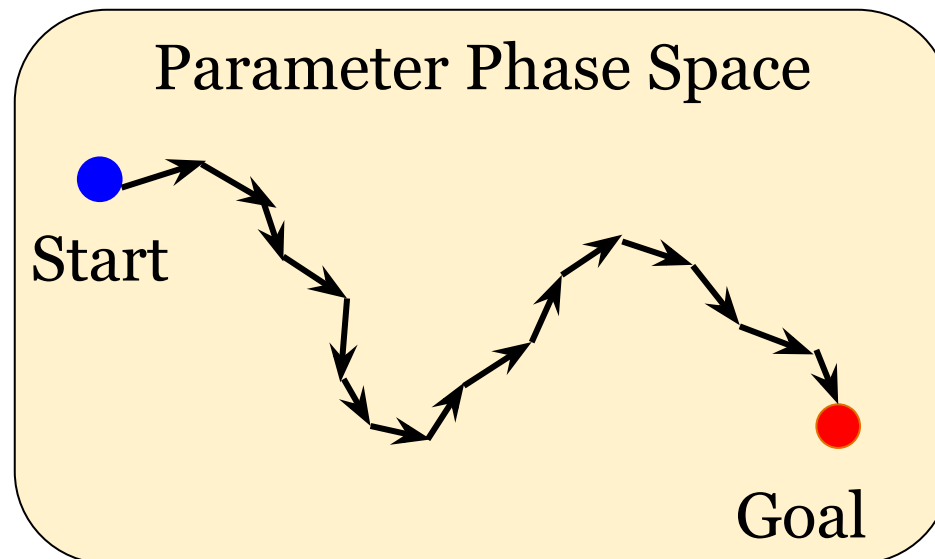
Introduction to Neural Networks

Gradient-based Optimization

Start: some random initial set of parameters

Goal: minimize the overall loss = sum over all samples

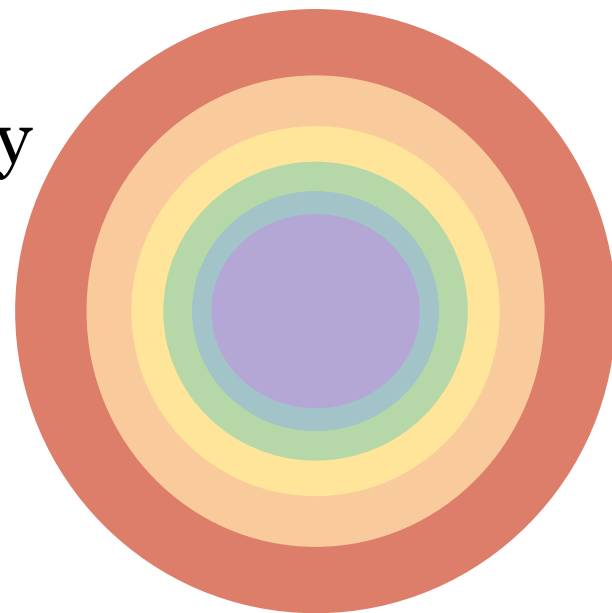
How: SGD from yesterday :)



Adaptive learning rate (LR)

Ideally start with a larger LR to quickly converge, then decrease the LR to fine-tune near the minimum.

Warning! too small LR rate is not only slow but can trap in false minimum.



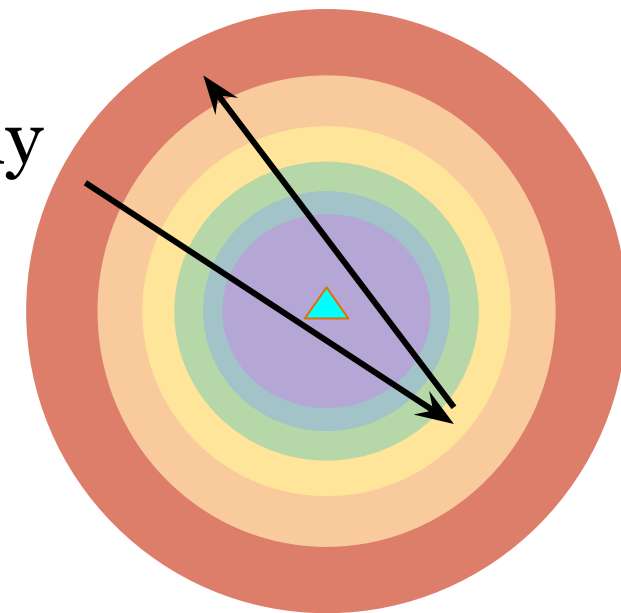
Introduction to Neural Networks

Gradient-based Optimization

Adaptive learning rate (LR)

Ideally start with a larger LR to quickly converge, then decrease the LR to fine-tune near the minimum.

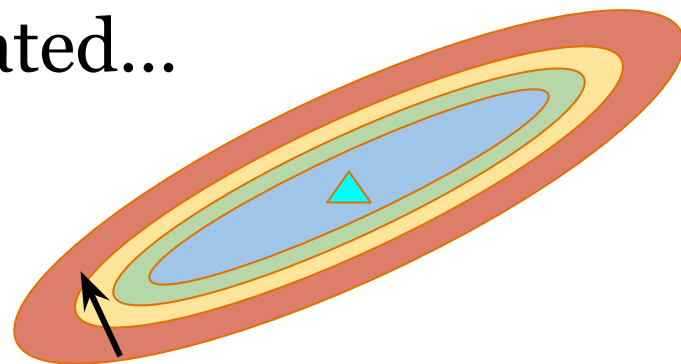
Warning! too small LR rate is not only slow but can trap in false minimum.



Momentum

The direction of the steepest descent can be almost perpendicular if the ellipse is elongated...

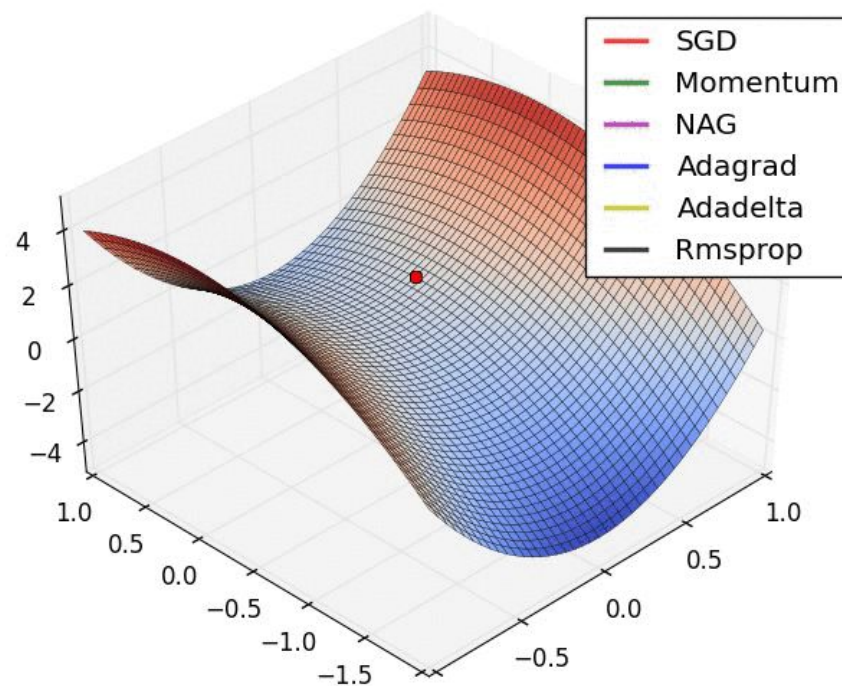
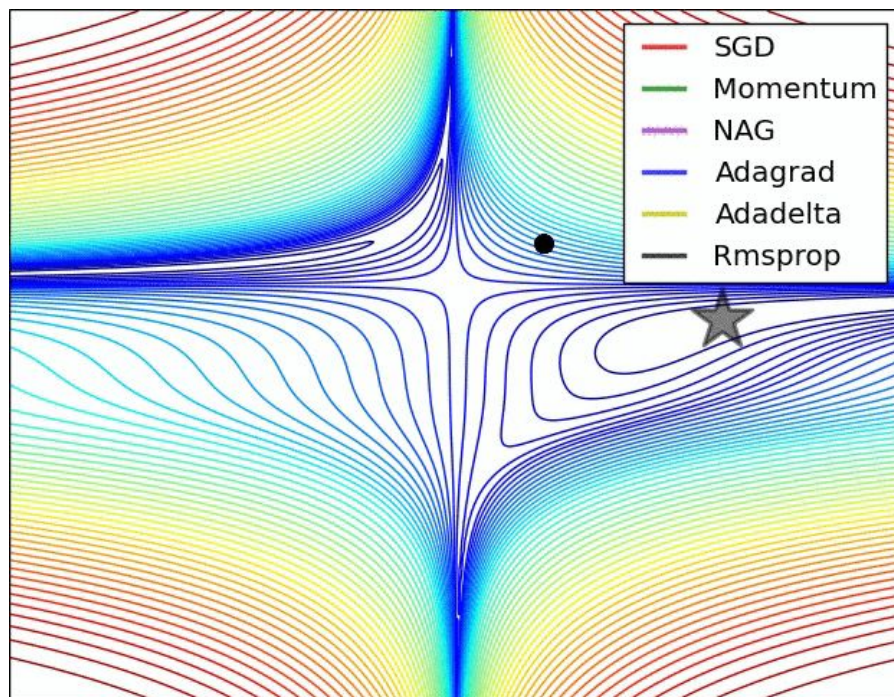
A solution is to average the history of past updates. This accumulates small updates in the right direction



Introduction to Neural Networks

Gradient-based Optimization

Popular Visualization of Optimizers (not mine)



- Excellent lecture [here \(CS232\)](#)
- Want to try by yourself ? You can get started [here](#).
- Popular choices: vanilla SGD and Adam

Image Credit: [Alec Radford](#)

Neural Network

- Each neuron produces a “feature”, activation function can add non-linearity, and additional layer can represent non-linear functional approximation.

Gradient-based Optimization

- Essentially a composite function = chain rule!
- Initialization should be non-zero random values
- Flavors of gradient-based optimizers to take into account for the error surface

Backup Slides

Cost Functions & Activation Functions

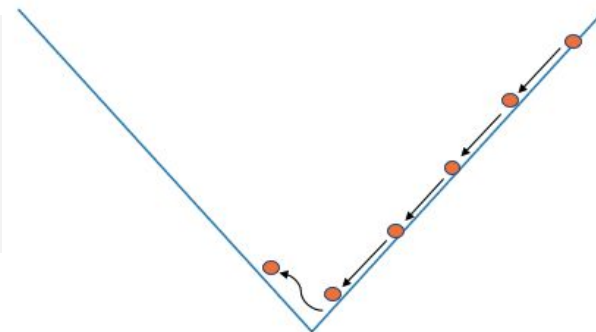
Introduction to Neural Networks

Loss/Cost Functions (Regressions)

1. Mean Absolute Error (MAE, L1 loss)

$$\text{MAE} = \frac{\sum^n |x_{\text{pred}} - y_{\text{true}}|}{n}$$

- Constant magnitude gradient everywhere
- Adaptive LR critical near the minimum



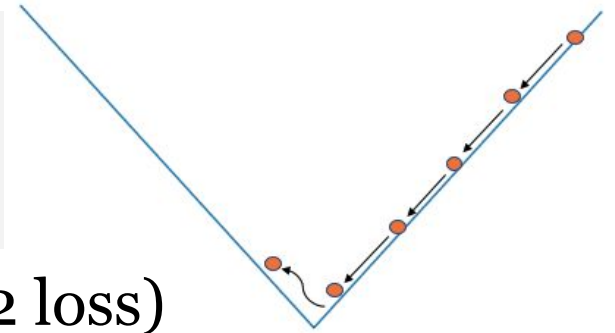
Introduction to Neural Networks

Loss/Cost Functions (Regressions)

1. Mean Absolute Error (MAE, L1 loss)

$$\text{MAE} = \frac{\sum_{n} |x_{\text{pred}} - y_{\text{true}}|}{n}$$

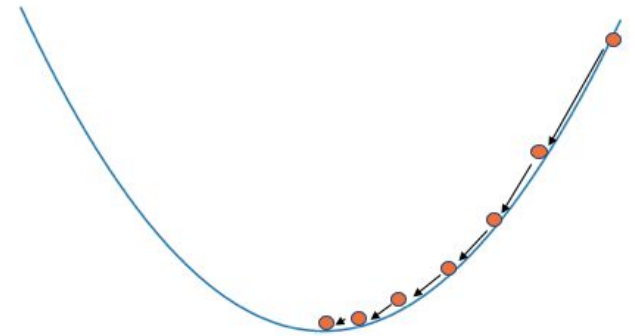
- Constant magnitude gradient everywhere
- Adaptive LR critical near the minimum



2. Mean Square Error (MSE, L2 loss)

$$\text{MSE} = \frac{\sum_{n} (x_{\text{pred}} - y_{\text{true}})^2}{n}$$

- Adaptive LR not needed
- Huge gradient far from the minimum



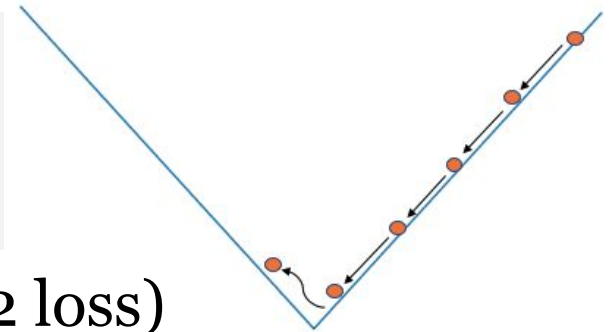
Introduction to Neural Networks

Loss/Cost Functions (Regressions)

1. Mean Absolute Error (MAE, L1 loss)

$$\text{MAE} = \frac{\sum_{n=1}^n |x_{\text{pred}} - y_{\text{true}}|}{n}$$

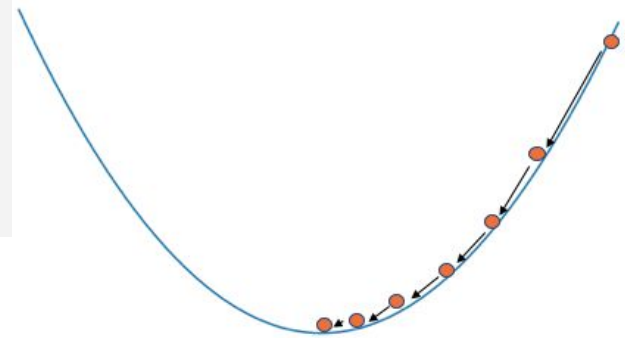
- Constant magnitude gradient everywhere
- Adaptive LR critical near the minimum



2. Mean Square Error (MSE, L2 loss)

$$\text{MSE} = \frac{\sum_{n=1}^n (x_{\text{pred}} - y_{\text{true}})^2}{n}$$

- Adaptive LR not needed
- Huge gradient far from the minimum



3. Hubar Loss

$$\text{Hubar} = \begin{cases} 0.5 (x_{\text{pred}} - y_{\text{true}})^2 \cdots & \text{if } |x_{\text{pred}} - y_{\text{true}}| < \delta \\ \delta |x_{\text{pred}} - y_{\text{true}}| - 0.5 \delta^2 \cdots & \text{otherwise} \end{cases}$$

Introduction to Neural Networks

Loss/Cost Functions (Classification)

Cross-Entropy Loss (entropy, softmax)

- Also called multinomial logistic loss
- Probability of being a class i in N categories

$$p_i(y_i) = \frac{\exp(x_i)}{\sum_{j=1}^N \exp(x_j)}$$

- Binomial is a special case where probability becomes **sigmoid**
- Entropy loss minimization is **maximum likelihood** method.
- Others: hinge (SVM) loss

- **Cross-entropy loss**

$$\mathcal{L}(x_i, y_i) = -\log \left[\frac{\exp(x_i)}{\sum_{j=1}^N \exp(x_j)} \right]$$