

# **CMPT 733 – Big Data Programming II**

## **Deep Learning I**

Instructor

Steven Bergner

Course website

<https://coursys.sfu.ca/2024sp-cmpt-733-g1/pages>

# Overview

## Renaissance of artificial neural networks

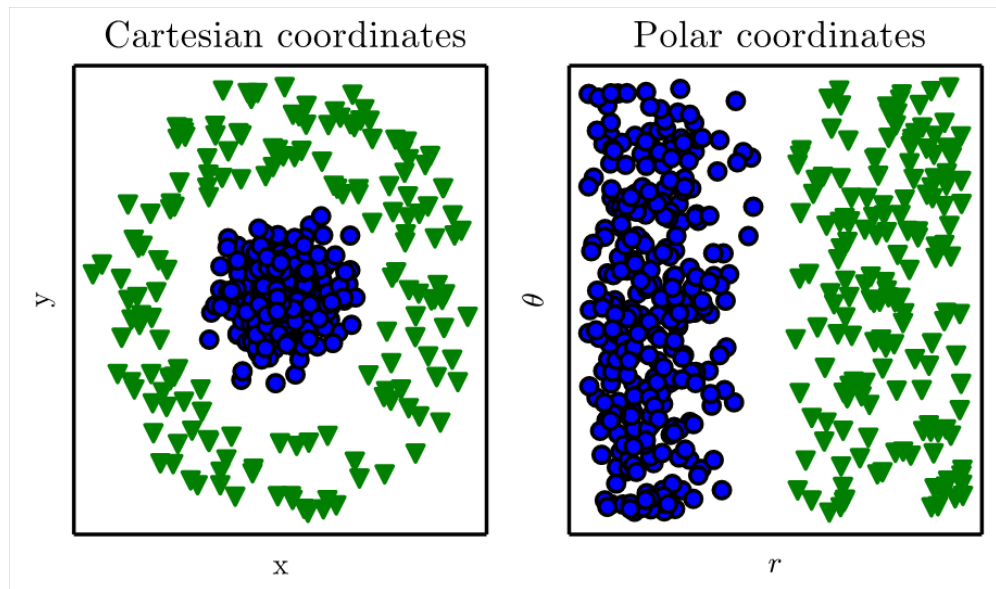
- Representation learning vs feature engineering

## Background

- Linear Algebra, Optimization
- Regularization
- Construction and training of layered models

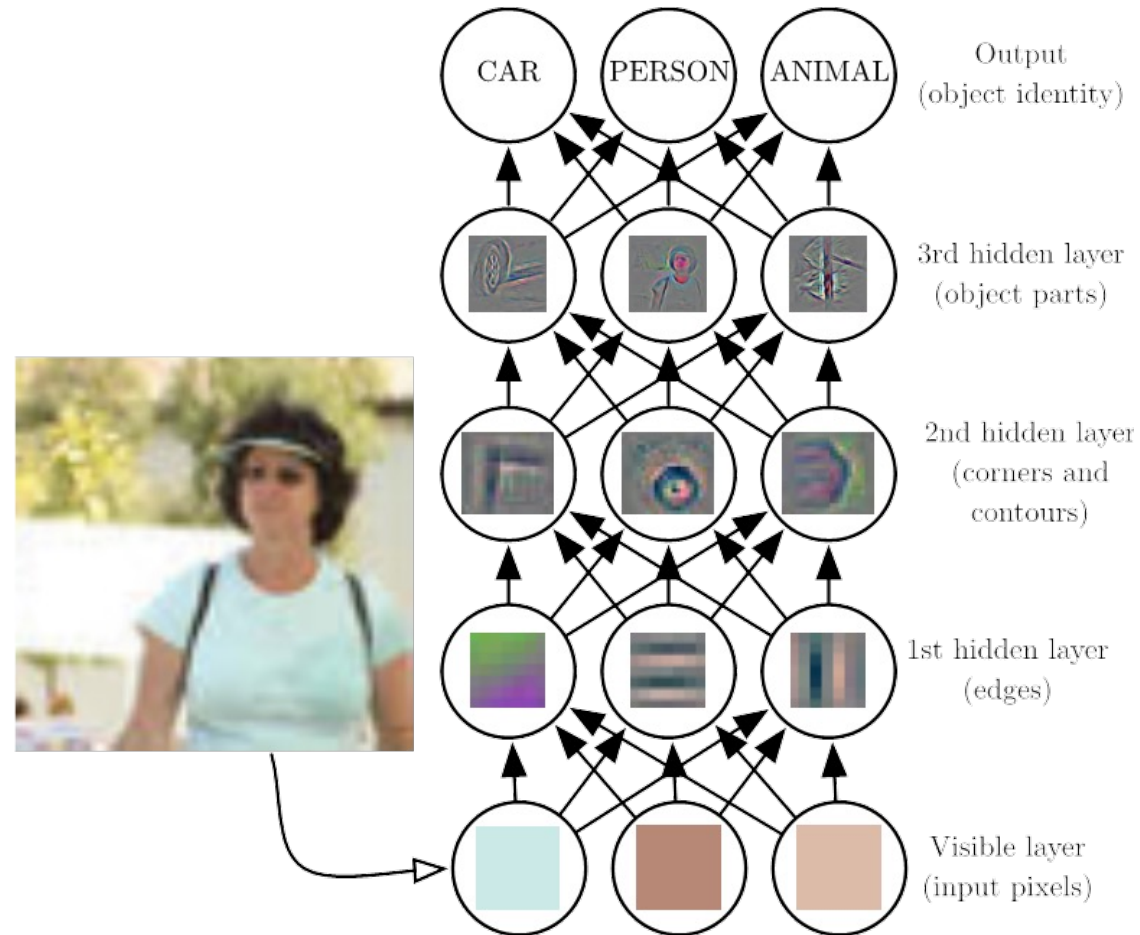
## Frameworks for deep learning

# Representations matter

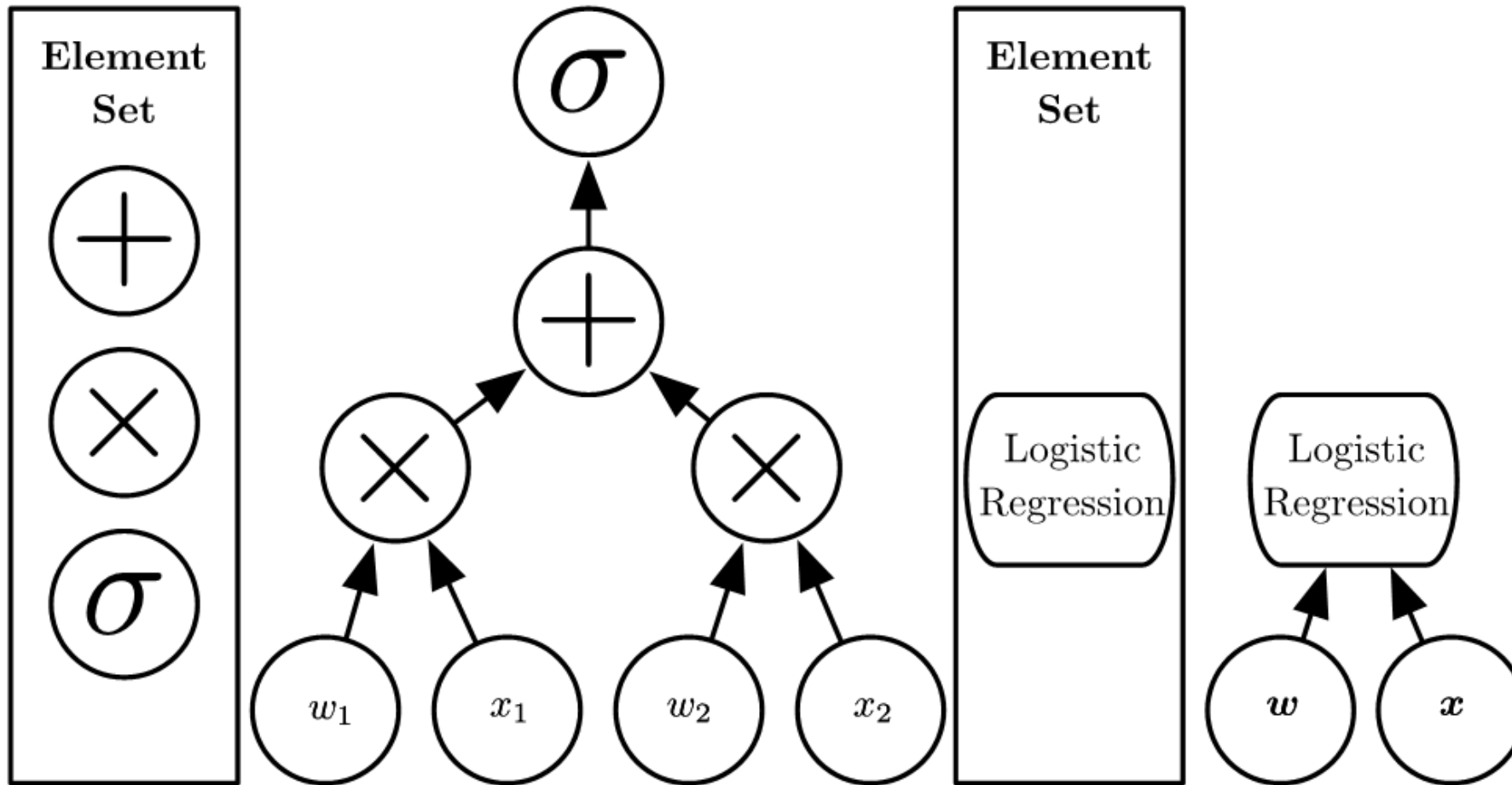


- Transform into the right representation
- Classify points simply by threshold on radius axis
- Single neuron with non-linearity can do this

# Depth: Layered composition

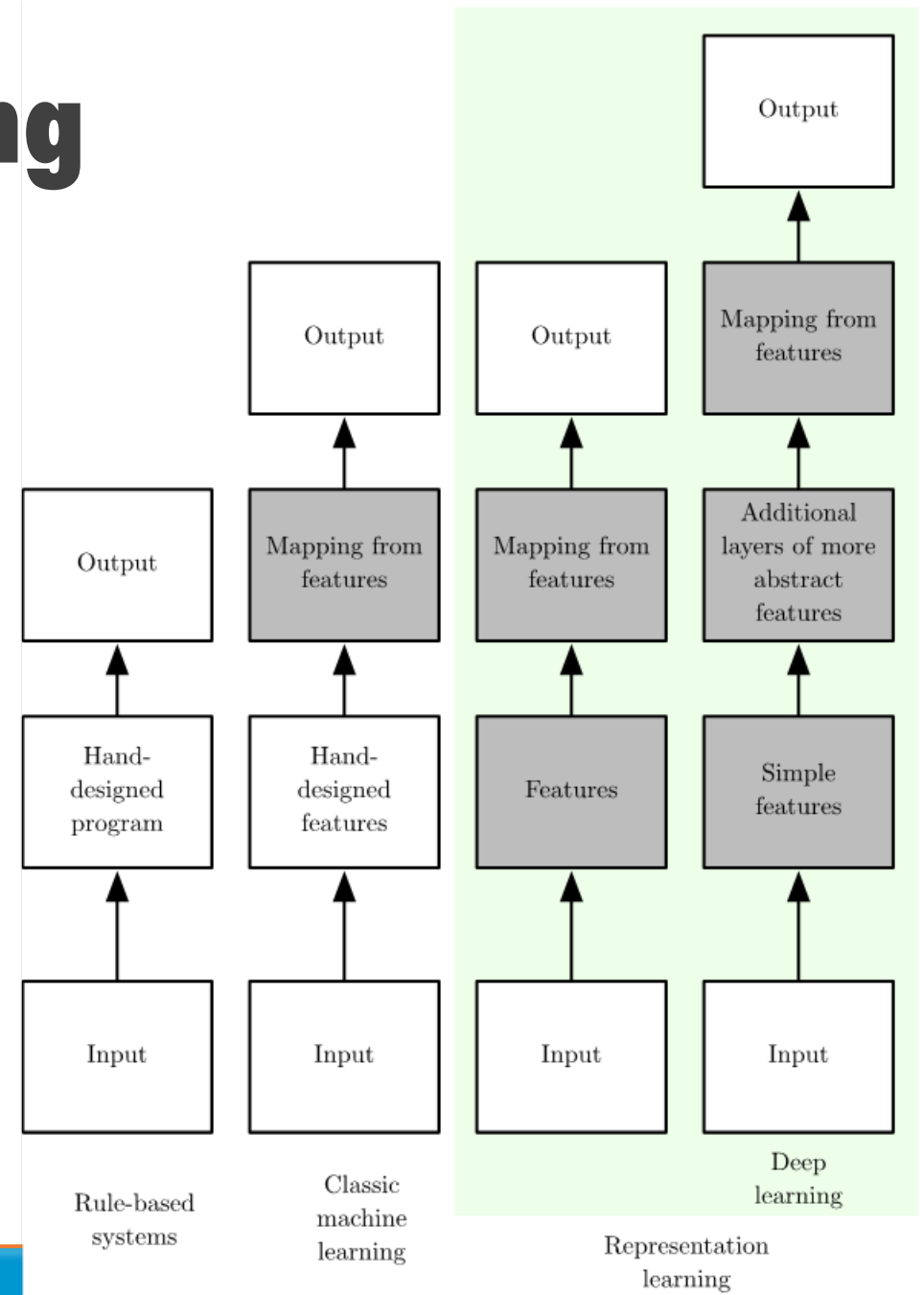


# Computational graph

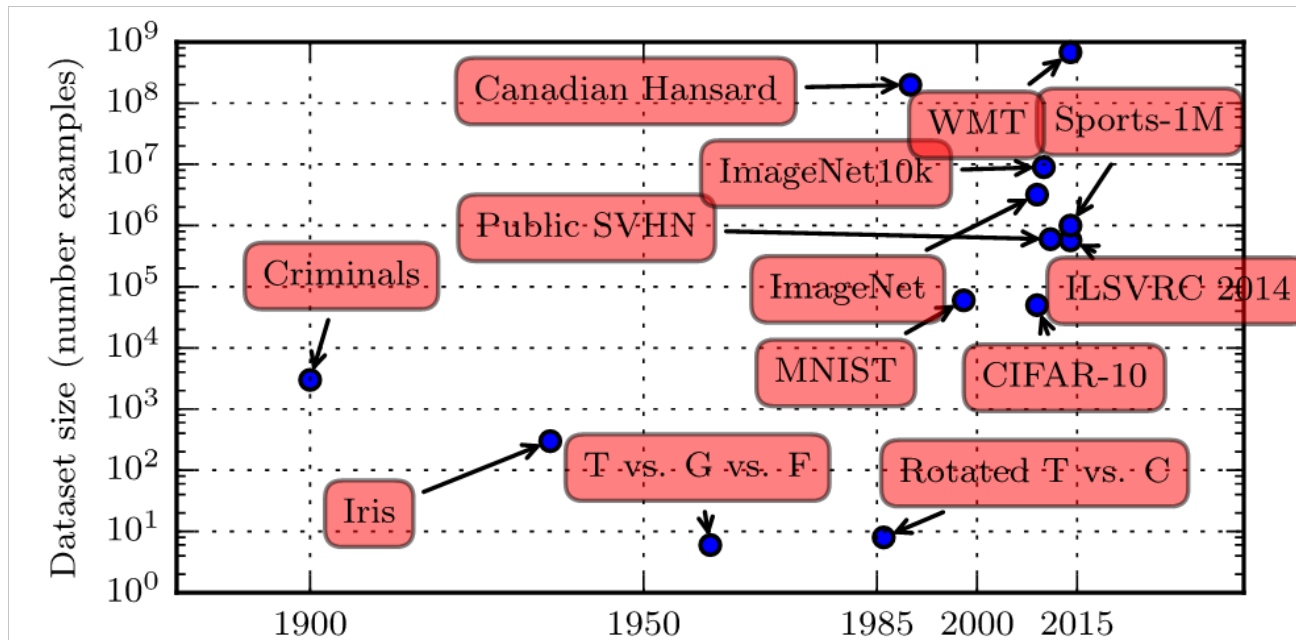


# Components of learning

- Hand designed program
  - Input → Output
- Increasingly automated
  - Simple features
  - Abstract features
  - Mapping from features



# Growing dataset size



MNIST dataset

|   |   |   |   |   |   |   |   |   |   |   |   |   |   |   |   |   |   |   |   |
|---|---|---|---|---|---|---|---|---|---|---|---|---|---|---|---|---|---|---|---|
| 8 | 9 | 0 | 1 | 2 | 3 | 4 | 7 | 8 | 9 | 0 | 1 | 2 | 3 | 4 | 5 | 6 | 7 | 8 | 6 |
| 4 | 2 | 6 | 4 | 7 | 5 | 5 | 4 | 7 | 8 | 9 | 2 | 9 | 3 | 9 | 3 | 8 | 2 | 0 | 5 |
| 0 | 7 | 0 | 4 | 2 | 6 | 5 | 3 | 5 | 3 | 8 | 0 | 0 | 3 | 4 | 1 | 5 | 3 | 0 | 8 |
| 3 | 0 | 6 | 2 | 7 | 1 | 1 | 8 | 1 | 7 | 1 | 3 | 8 | 9 | 7 | 6 | 7 | 4 | 1 | 6 |
| 7 | 5 | 1 | 7 | 1 | 9 | 8 | 0 | 6 | 9 | 4 | 9 | 9 | 3 | 7 | 1 | 9 | 2 | 2 | 5 |
| 3 | 7 | 8 | 2 | 3 | 4 | 5 | 6 | 7 | 8 | 9 | 0 | 1 | 2 | 3 | 4 | 5 | 6 | 7 | 0 |
| 1 | 2 | 3 | 4 | 5 | 6 | 7 | 8 | 9 | 8 | 1 | 0 | 5 | 5 | 1 | 9 | 0 | 4 | 1 | 9 |
| 3 | 8 | 4 | 7 | 7 | 8 | 5 | 0 | 6 | 5 | 5 | 3 | 3 | 3 | 9 | 8 | 1 | 4 | 0 | 6 |
| 1 | 0 | 0 | 6 | 2 | 1 | 1 | 3 | 2 | 8 | 8 | 7 | 8 | 4 | 6 | 0 | 2 | 0 | 3 | 6 |
| 8 | 7 | 1 | 5 | 9 | 9 | 3 | 2 | 4 | 9 | 4 | 6 | 5 | 3 | 2 | 5 | 5 | 9 | 4 | 1 |
| 6 | 5 | 0 | 1 | 2 | 3 | 4 | 5 | 6 | 7 | 8 | 9 | 0 | 1 | 2 | 3 | 4 | 5 | 6 | 7 |
| 8 | 9 | 0 | 1 | 2 | 3 | 4 | 5 | 6 | 7 | 8 | 9 | 6 | 4 | 2 | 6 | 4 | 7 | 5 | 5 |
| 4 | 7 | 8 | 9 | 2 | 9 | 3 | 9 | 3 | 8 | 2 | 0 | 9 | 8 | 0 | 5 | 6 | 0 | 1 | 0 |
| 4 | 2 | 6 | 5 | 5 | 5 | 4 | 3 | 4 | 1 | 5 | 3 | 0 | 8 | 3 | 0 | 6 | 2 | 7 | 1 |
| 1 | 8 | 1 | 7 | 1 | 3 | 8 | 5 | 4 | 2 | 0 | 9 | 7 | 6 | 7 | 4 | 1 | 6 | 8 | 4 |
| 7 | 5 | 1 | 2 | 6 | 7 | 1 | 9 | 8 | 0 | 6 | 9 | 4 | 9 | 9 | 6 | 2 | 3 | 7 | 1 |
| 9 | 2 | 2 | 5 | 3 | 7 | 8 | 0 | 1 | 2 | 3 | 4 | 5 | 6 | 7 | 8 | 0 | 1 | 2 | 3 |
| 4 | 5 | 6 | 7 | 8 | 0 | 1 | 2 | 3 | 4 | 5 | 6 | 7 | 8 | 9 | 2 | 1 | 2 | 1 | 3 |
| 9 | 9 | 8 | 5 | 3 | 7 | 0 | 7 | 7 | 5 | 7 | 9 | 9 | 4 | 7 | 0 | 3 | 4 | 1 | 4 |
| 4 | 7 | 5 | 8 | 1 | 4 | 8 | 4 | 1 | 8 | 6 | 6 | 4 | 6 | 3 | 5 | 7 | 2 | 5 | 9 |

# **Basics**

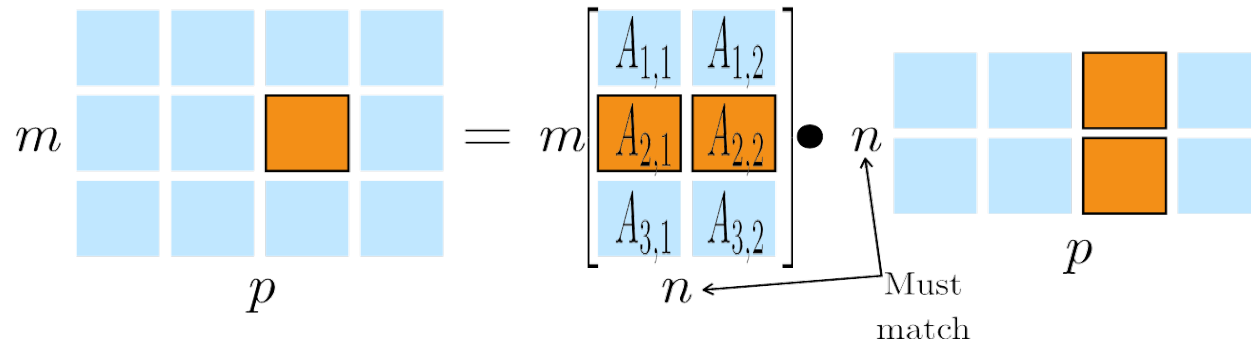
## **Linear Algebra and Optimization**



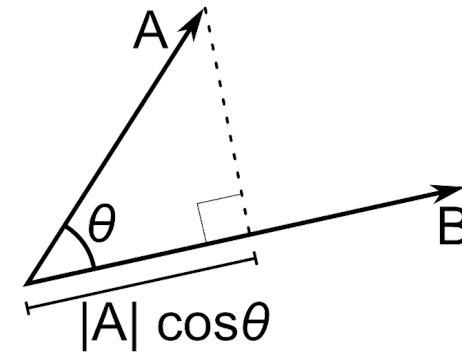
# Linear algebra

- Tensor is an array of numbers
  - Multi-dim: 0d scalar, 1d vector, 2d matrix/image, 3d RGB image

- Matrix (dot) product  $C = AB$   $C_{i,j} = \sum_k A_{i,k} B_{k,j}$



- Dot product of vectors A and B
  - (m = p = 1 in above notation)



# Linear algebra: Norms

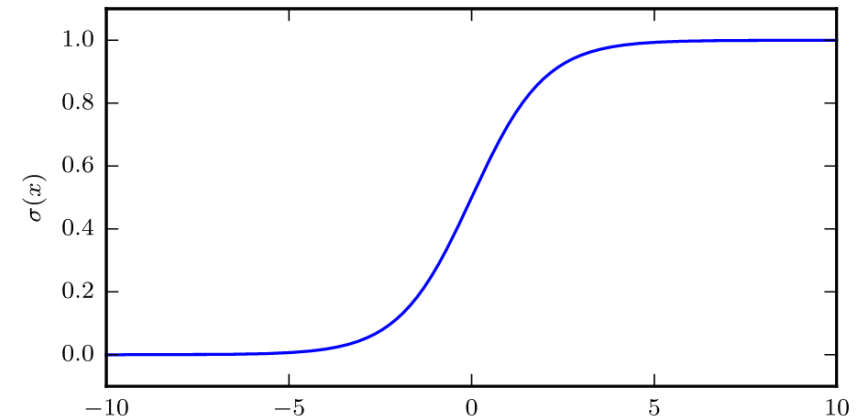
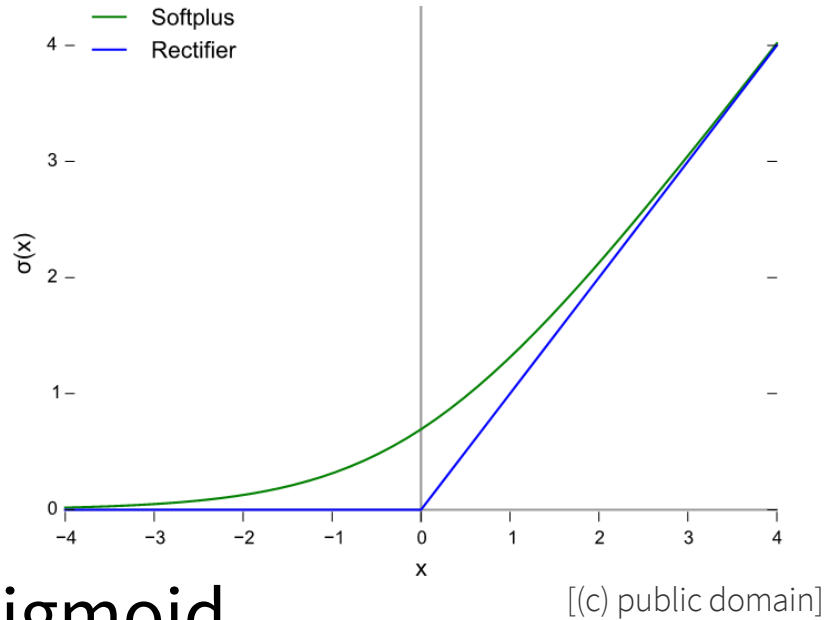
- $L^p$  norm

$$\|\mathbf{x}\|_p = \left( \sum_i |x_i|^p \right)^{\frac{1}{p}}$$

- Most popular norm: L2 norm,  $p=2$
- L1 norm,  $p=1$ :  $\|\mathbf{x}\|_1 = \sum_i |x_i|$ .
- Max norm, infinite  $p$ :  $\|\mathbf{x}\|_\infty = \max_i |x_i|$ .

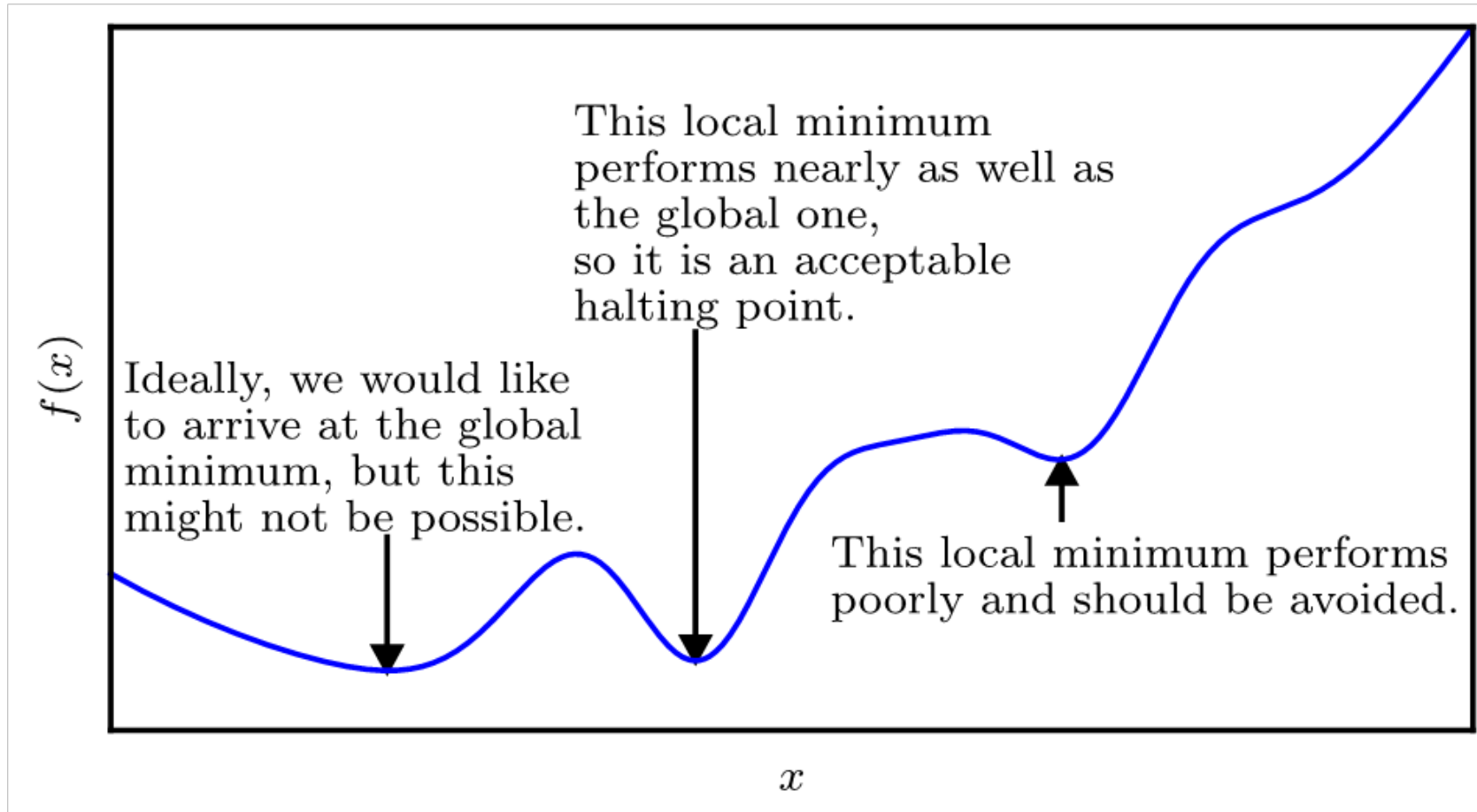
# Nonlinearities

- ReLU
- Softplus
- Logistic Sigmoid

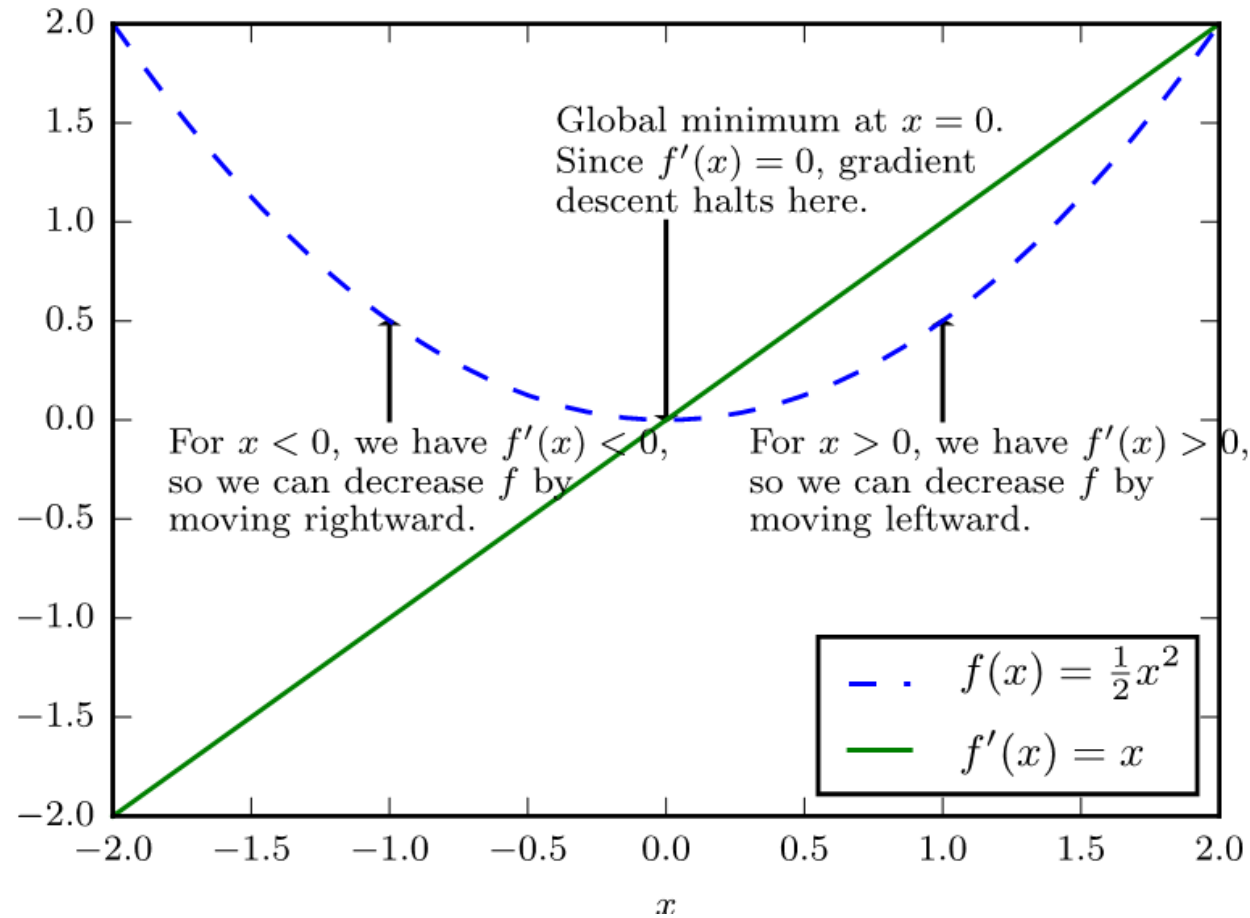


[Goodfellow, Bengio, Courville 2016]

# Approximate optimization

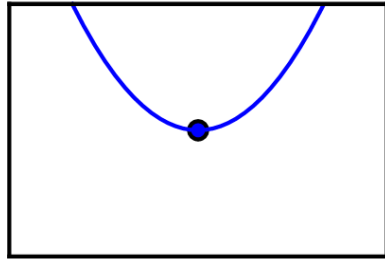


# Gradient descent

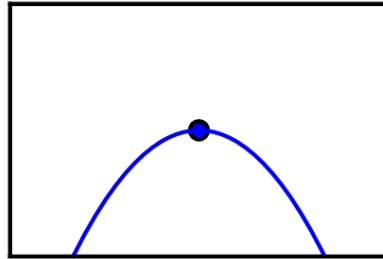


# Critical points

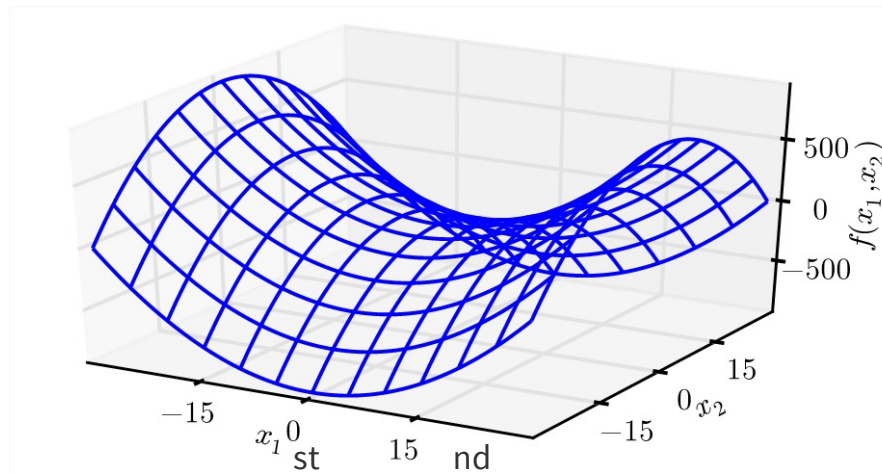
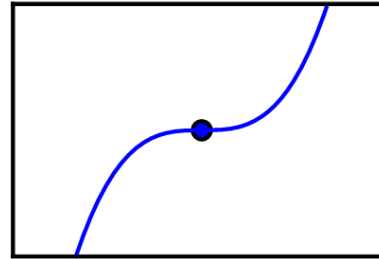
Minimum



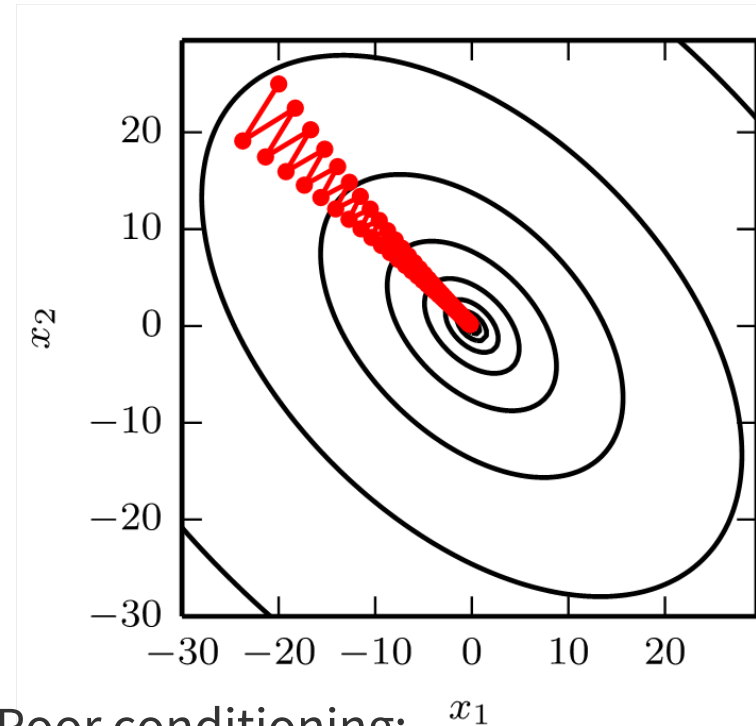
Maximum



Saddle point



Saddle point – 1 and 2 derivative vanish



Poor conditioning:  
1st deriv large in one and small in another direction

# Optimization algorithm

- Lots of variants address choice of learning rate
- See [Visualization of Algorithms](#)
- AdaDelta and RMSprop often work well

# Neural network playgrounds

- <http://playground.tensorflow.org/>
  - Try out simple network configurations on TF Playground
- <https://cs.stanford.edu/people/karpathy/convnetjs/demo/classify2d.html>
  - Visualize linear and non-linear mappings

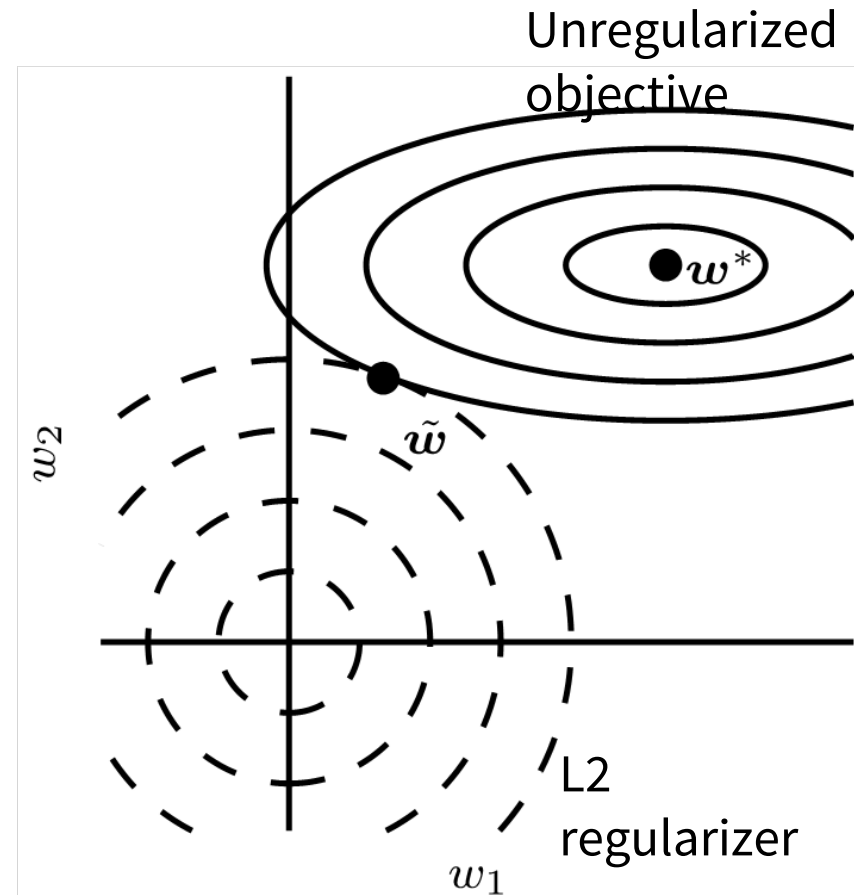


# Regularization

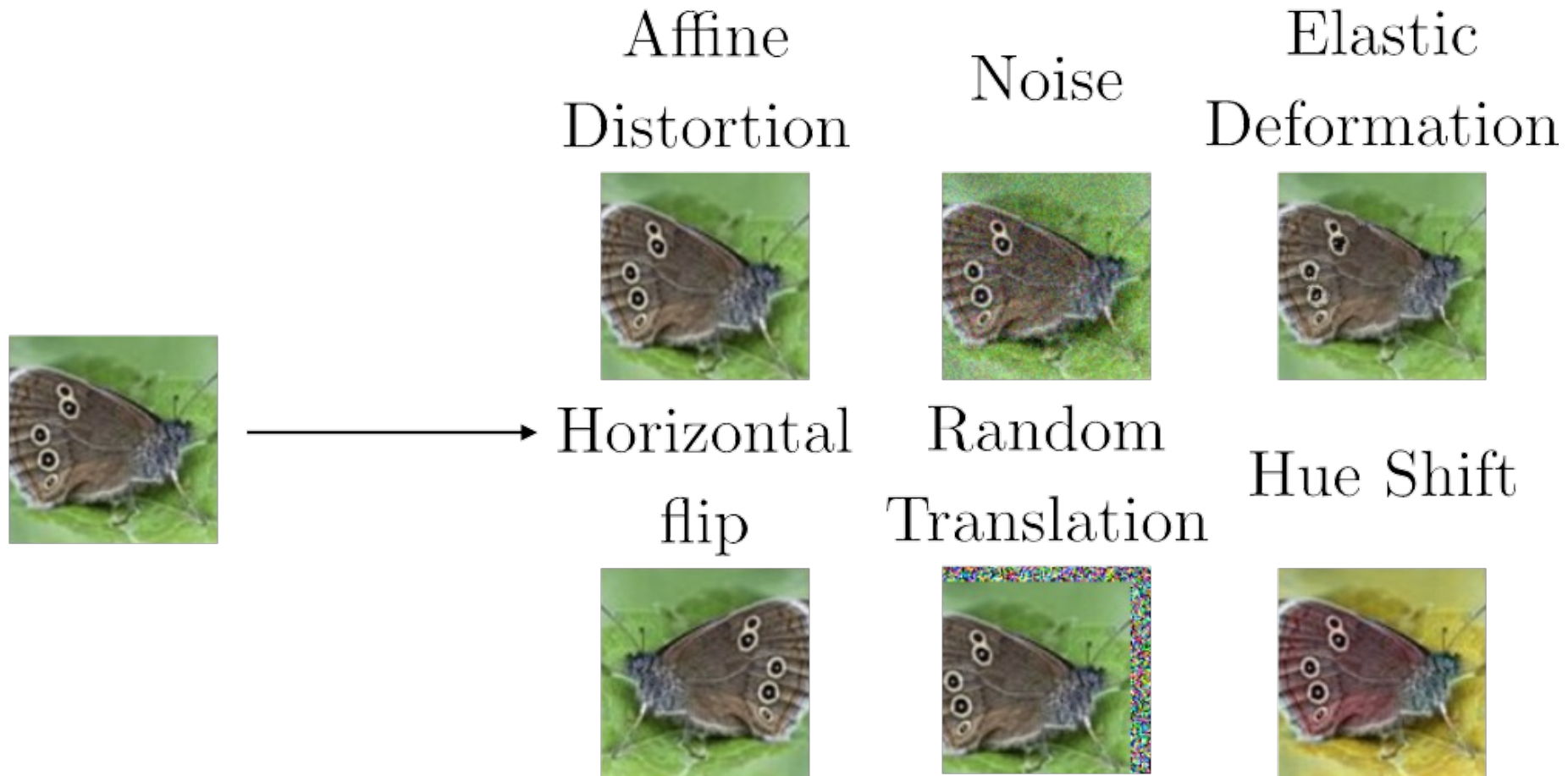
Reduced generalization error without impacting training error

# Constrained optimization

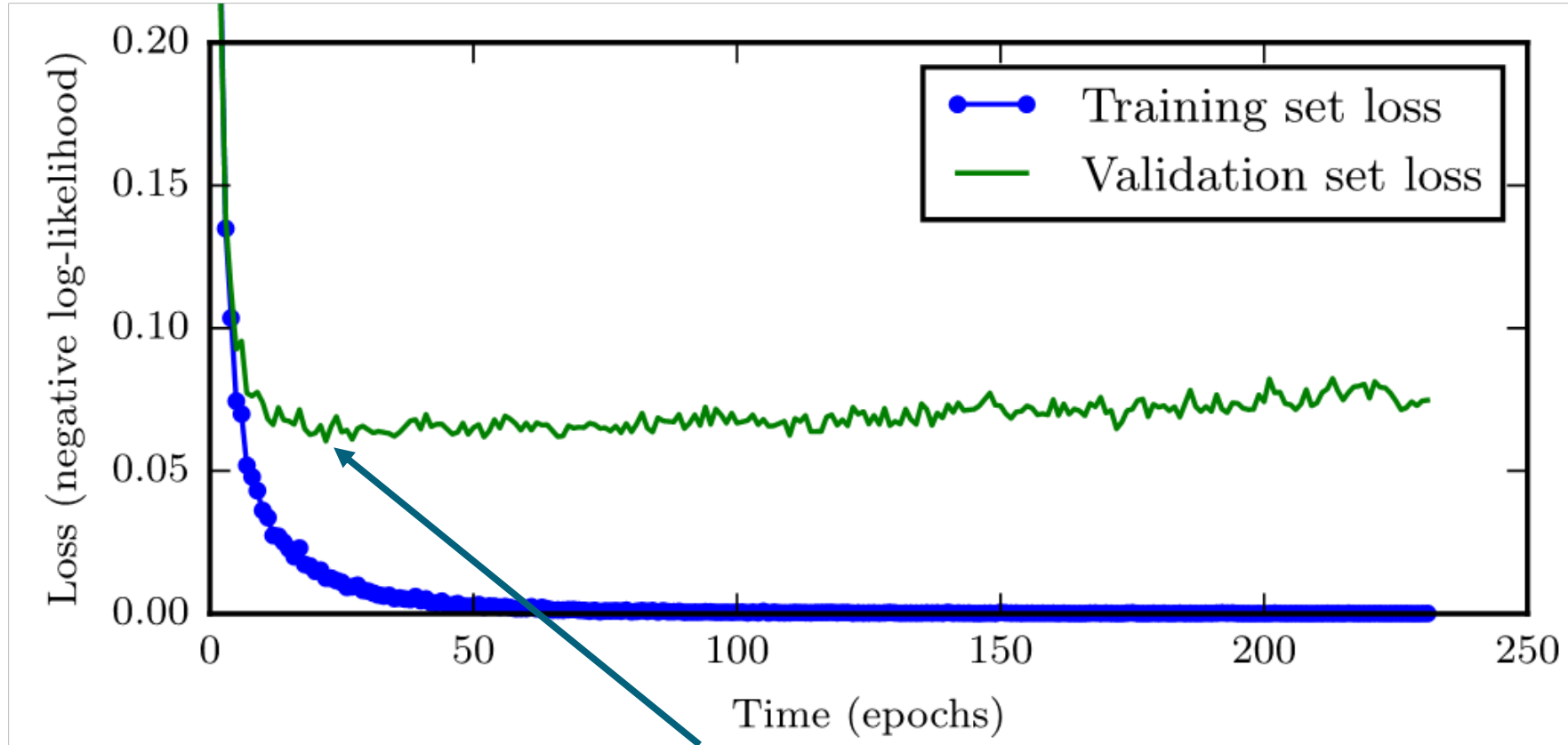
- Squared L2 encourages small weights
- L1 encourages sparsity of model parameters (weights)



# Dataset augmentation



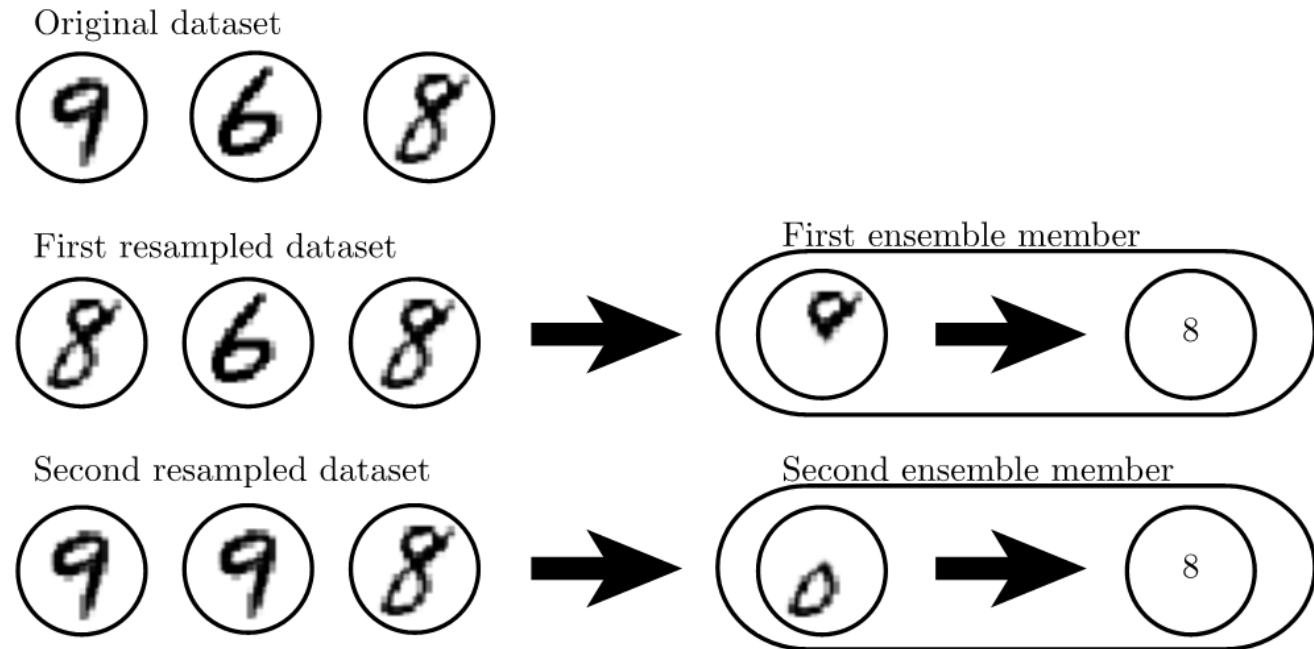
# Learning curves



- Early stopping before validation error starts to increase

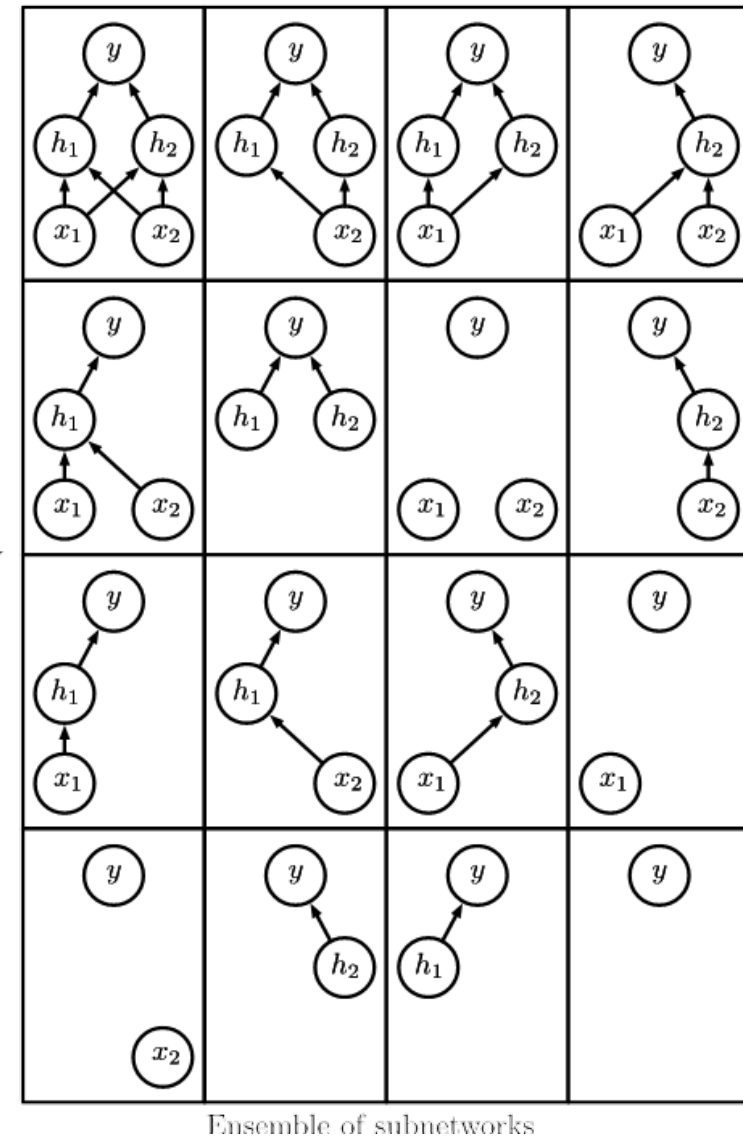
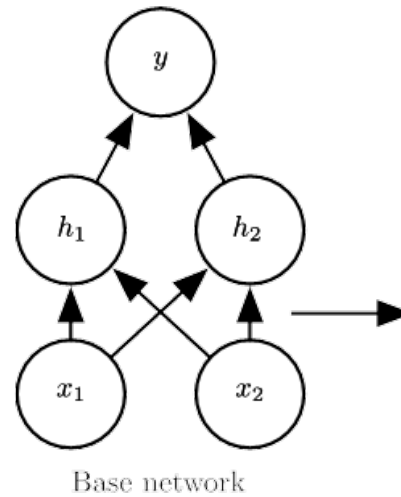
# Bagging

- Average multiple models trained on subsets of the data
- First subset: learns top loop, Second subset: bottom loop

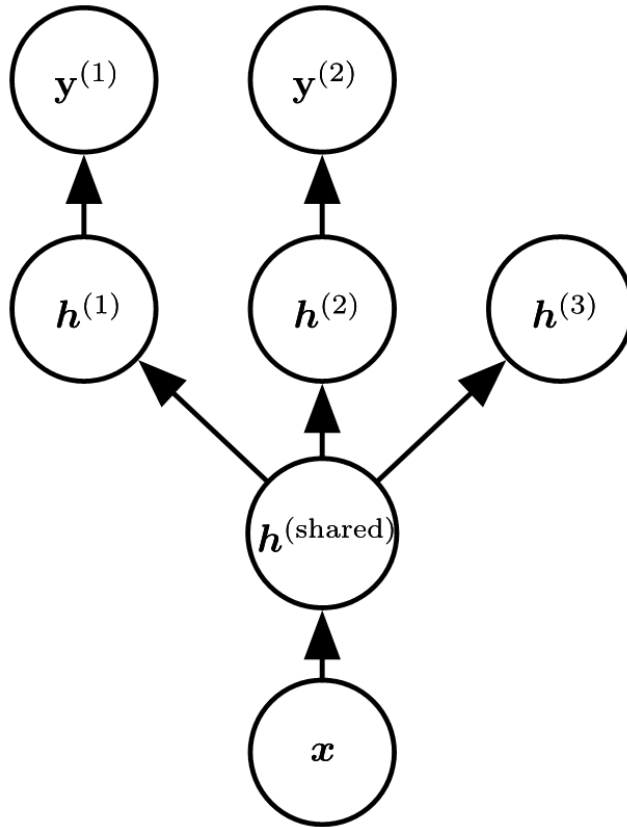


# Dropout

- Random sample of connection weights is set to zero
- Train different network model each time
- Learn more robust, generalizable features



# Multitask learning



- Shared parameters are trained with more data
- Improved generalization error due to increased statistical strength

# **Components of popular architectures**



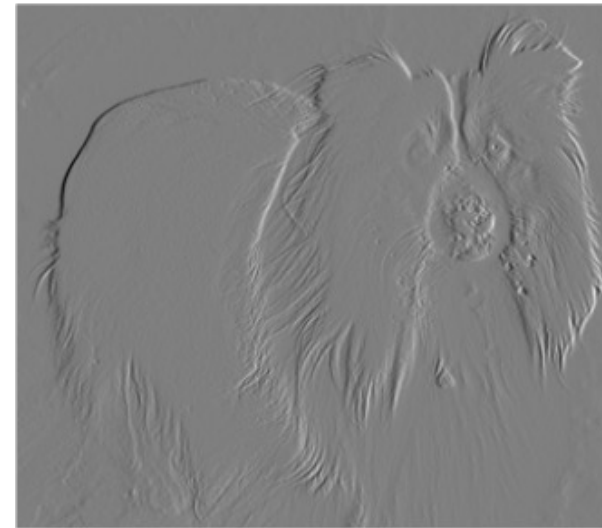
# Convolution as edge detector



Input

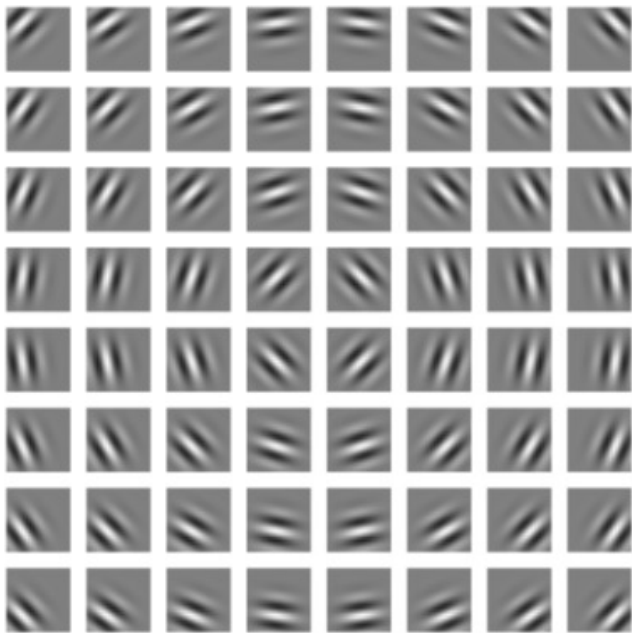
|   |    |
|---|----|
| 1 | -1 |
|---|----|

Kernel

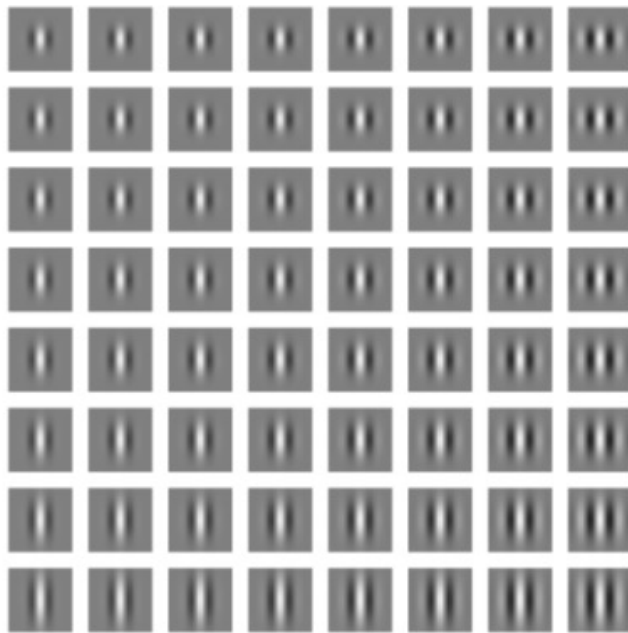


Output

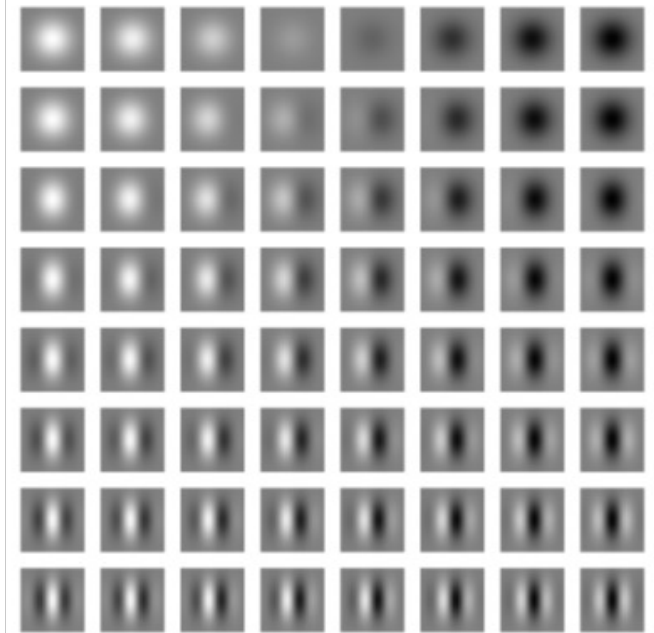
# Gabor wavelets (kernels)



Directional second  
derivative

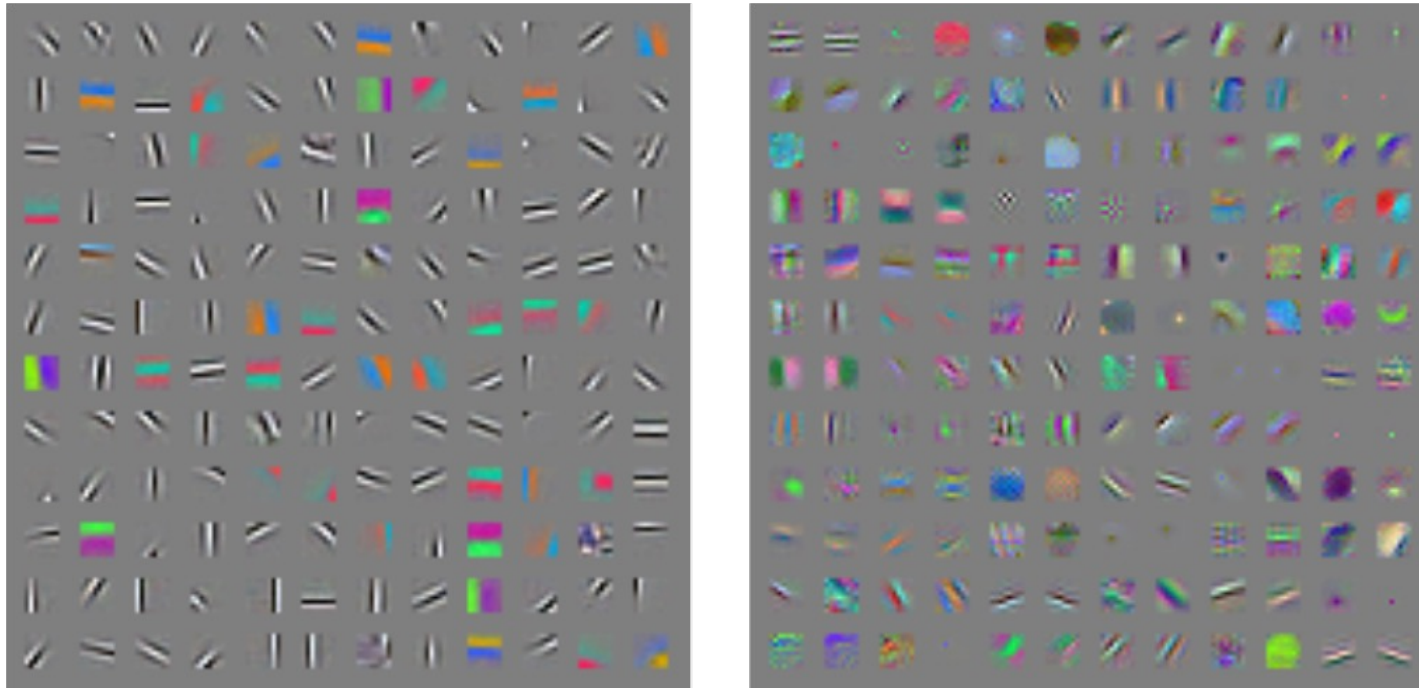


Second derivative  
(curvature)



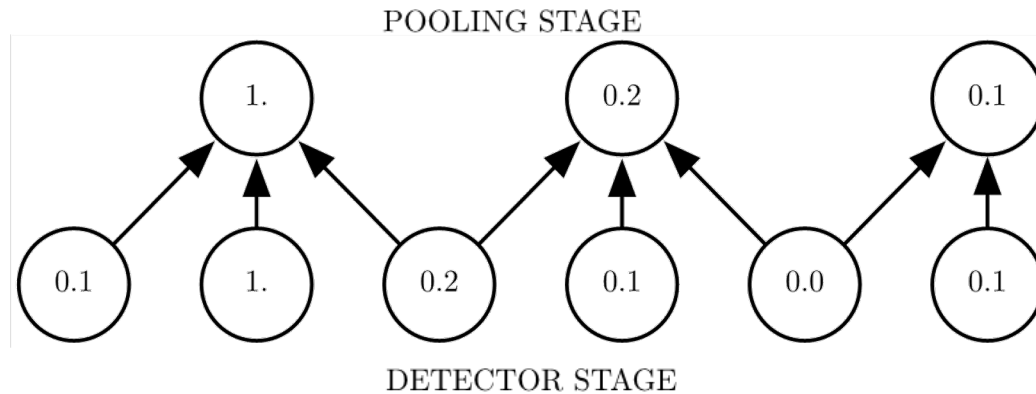
Local average, first  
derivative

# Gabor-like learned kernels

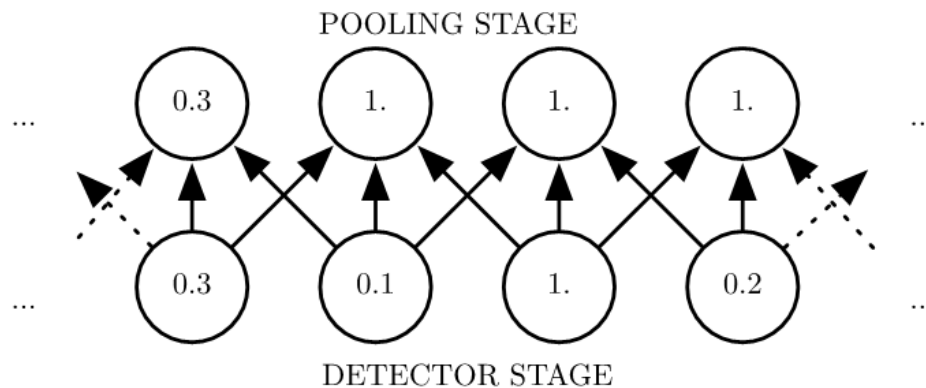


- Features extractors provided by pretrained networks

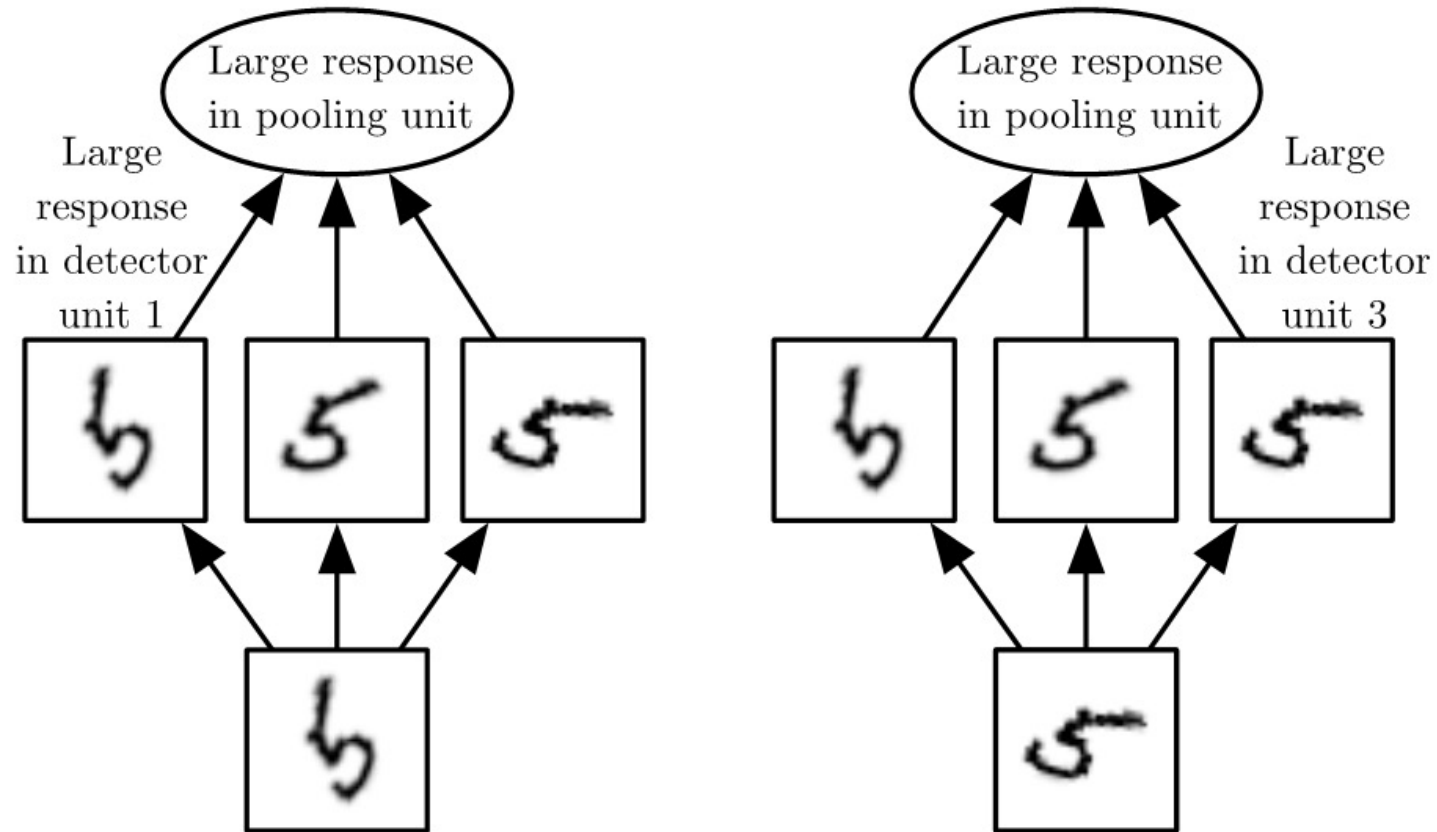
# Max pooling translation invariance



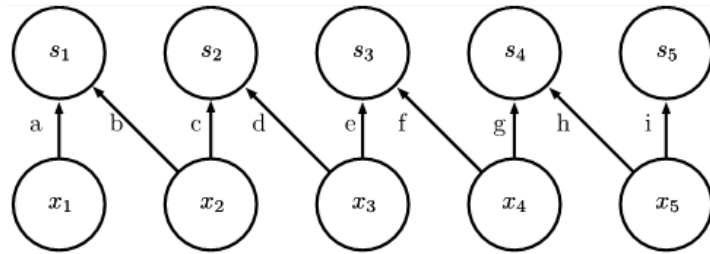
- Take max of certain neighbourhood
- Often combined, followed by downsampling



# Max pooling transform invariance



# Types of connectivity



Local connection:  
like convolution,  
but no sharing

# Choosing architecture family

- No structure → fully connected
- Spatial structure → convolutional
- Sequential structure → recurrent

# **Software for Deep Learning**



# Current Frameworks

- Tensorflow / Keras
- PyTorch
- DL4J
- Caffe (superseded by Caffe2, which is merged into PyTorch)
- And many more
- Most have CPU-only mode but much faster on NVIDIA GPU

# Development strategy

- Identify needs: High accuracy or low accuracy?
- Choose metric
  - Accuracy (% of examples correct), Coverage (% examples processed)
  - Precision  $TP/(TP+FP)$ , Recall  $TP/(TP+FN)$
  - Amount of error in case of regression
- Build end-to-end system
  - Start from baseline, e.g. initialize with pre-trained network
- Refine driven by data

# Sources

- I. Goodfellow, Y. Bengio, A. Courville “Deep Learning” MIT Press 2016 [[link](#)]