



Wrap up MAP

+ Convolutional Neural Networks (CNNs)

+ Recurrent Neural Networks (RNNs)

Pat Virtue & Matt Gormley

Lecture 17

Mar. 16, 2026

REMINDER: MAP

EXAMPLE: BETA PRIOR FOR BERNOULLI LIKELIHOOD

The MAP Estimation Objective

MLE:

$$p(\mathcal{D} | \theta)$$

likelihood

posterior

MAP:

$$p(\theta | \mathcal{D})$$

$$= \frac{p(\mathcal{D} | \theta)p(\theta)}{p(\mathcal{D})}$$

prior

Bayes Rule

$$\theta_{MAP} = \operatorname{argmax}_{\theta} p(\theta | \mathcal{D})$$

$$= \operatorname{argmax}_{\theta} \frac{p(\mathcal{D} | \theta)p(\theta)}{p(\mathcal{D})}$$

$$= \operatorname{argmax}_{\theta} p(\mathcal{D} | \theta)p(\theta)$$

$$= \operatorname{argmax}_{\theta} \log p(\mathcal{D} | \theta) + \log p(\theta)$$

$\ell_{MAP}(\theta)$

Recipe for Closed-form MLE

1. Assume data was generated iid from some model, i.e., write the *generative story*

$$x^{(i)} \sim p(x|\boldsymbol{\theta})$$

2. Write the log-likelihood

$$\ell(\boldsymbol{\theta}) = \log p(x^{(1)}|\boldsymbol{\theta}) + \dots + \log p(x^{(N)}|\boldsymbol{\theta})$$

3. Compute partial derivatives, i.e., the gradient

$$\partial \ell(\boldsymbol{\theta}) / \partial \theta_1 = \dots$$

...

$$\partial \ell(\boldsymbol{\theta}) / \partial \theta_M = \dots$$

4. Set derivatives equal to zero and solve for $\boldsymbol{\theta}$

$$\partial \ell(\boldsymbol{\theta}) / \partial \theta_m = 0 \text{ for all } m \in \{1, \dots, M\}$$

$\boldsymbol{\theta}^{\text{MLE}}$ = solution to system of M equations and M variables

5. Compute the second derivative and check that $\ell(\boldsymbol{\theta})$ is concave down at $\boldsymbol{\theta}^{\text{MLE}}$

Recipe for Closed-form MAP

1. Assume data was generated iid from some model, i.e., write the *generative story*

$$\boldsymbol{\theta} \sim p(\boldsymbol{\theta}) \text{ and then for all } i: x^{(i)} \sim p(x|\boldsymbol{\theta})$$

2. Write the log posterior

$$\ell_{\text{MAP}}(\boldsymbol{\theta}) = \log p(x^{(1)}|\boldsymbol{\theta}) + \dots + \log p(x^{(N)}|\boldsymbol{\theta}) + \log p(\boldsymbol{\theta})$$

3. Compute partial derivatives, i.e., the gradient

$$\partial \ell_{\text{MAP}}(\boldsymbol{\theta}) / \partial \theta_1 = \dots$$

...

$$\partial \ell_{\text{MAP}}(\boldsymbol{\theta}) / \partial \theta_M = \dots$$

4. Set derivatives to equal zero and solve for $\boldsymbol{\theta}$

$$\partial \ell_{\text{MAP}}(\boldsymbol{\theta}) / \partial \theta_m = 0 \text{ for all } m \in \{1, \dots, M\}$$

$\boldsymbol{\theta}^{\text{MAP}}$ = solution to system of M equations and M variables

5. Compute the second derivative and check that $\ell_{\text{MAP}}(\boldsymbol{\theta})$ is concave down at $\boldsymbol{\theta}^{\text{MAP}}$

MAP

EXAMPLE: BETA PRIOR FOR BERNOULLI LIKELIHOOD

The Prior Distribution

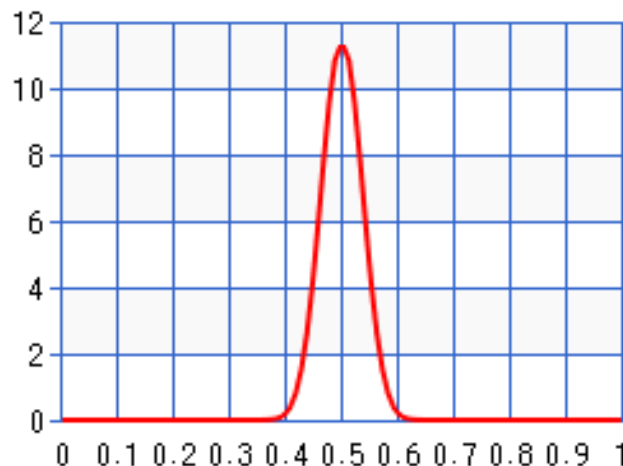
$$\pi_p(x^{(i)} | \phi)$$

- The prior distribution encodes domain knowledge about the problem.
- Question: Why do we use the Beta distribution as the prior for the Bernoulli?
- Reason #1: It has the right support, i.e. $[0,1]$.

$$\phi \sim \text{Beta}(\alpha, \beta)$$

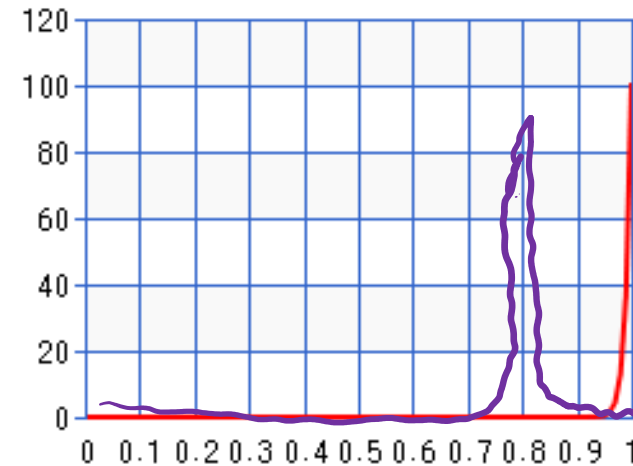
$$= \frac{1}{B(\alpha, \beta)} \phi^{\alpha-1} (1-\phi)^{\beta-1}$$

Example: Beta prior “fair coin”



$$f(\phi | \alpha = 101, \beta = 101)$$

Example: Beta prior “unfair coin”



$$f(\phi | \alpha = 101, \beta = 1)$$

$$N_1 = \alpha - 1$$

$$N_0 = \beta - 1$$

The Prior Distribution

- The prior distribution encodes domain knowledge about the problem.
- Question: Why do we use the Beta distribution as the prior for the Bernoulli?
- Reason #2: The Beta is a conjugate prior for the Bernoulli.
- Definition: A distribution is the **conjugate prior** of a likelihood if the form of the posterior is the same as the form of the prior.



Posterior $p(\theta D)$	Likelihood $p(D \theta)$	Prior $p(\theta)$	Conjugate?
Beta	<u>Bernoulli</u>	<u>Beta</u>	yes
Dirichlet	Multinomial	Multinomial	yes
Gaussian	Guassian	Guassian	yes
Gamma	Exponential	Gamma	yes
??	Multinomial	Logistic Normal	no

MLE of Bernoulli Model

1. Model: $\mathbf{x}^{(i)} \sim \text{Bernoulli}(\phi)$ for $i = 1, \dots, N$

$$\begin{aligned} p(x^{(i)} | \phi) &= \begin{cases} \phi & \text{if } x^{(i)} = 1 \\ (1 - \phi) & \text{if } x^{(i)} = 0 \end{cases} \\ &= \phi^{x^{(i)}} (1 - \phi)^{1-x^{(i)}} \end{aligned}$$

$$\begin{aligned} N_1 &= \#(x^{(i)} = 1) \\ N_0 &= \#(x^{(i)} = 0) \end{aligned}$$

2. Log-likelihood:

$$\begin{aligned} \ell_{\text{MLE}}(\phi) &= \log p(\mathcal{D} | \phi) \\ &= \log \prod_{i=1}^N p(x^{(i)} | \phi) \\ &= \log \prod_{i=1}^N \phi^{x^{(i)}} (1 - \phi)^{1-x^{(i)}} \\ &= \log (\phi^{N_1} (1 - \phi)^{N_0}) \\ &= N_1 \log(\phi) + N_0 \log(1 - \phi) \end{aligned}$$

3. Derivative:

$$\begin{aligned} \frac{\partial \ell_{\text{MLE}}(\phi)}{\partial \phi} &= \frac{\partial}{\partial \phi} (N_1 \log(\phi) + N_0 \log(1 - \phi)) \\ &= \frac{N_1}{\phi} - \frac{N_0}{1 - \phi} \end{aligned}$$

4. Set to zero and solve:

$$\begin{aligned} \frac{N_1}{\phi} - \frac{N_0}{1 - \phi} &= 0 \\ \Rightarrow \phi_{\text{MLE}} &= \frac{N_1}{N_1 + N_0} = \frac{N_1}{N} \end{aligned}$$

MAP of Beta-Bernoulli Model

1. Model: $\phi \sim \text{Beta}(\alpha, \beta)$
 $\mathbf{x}^{(i)} \sim \text{Bernoulli}(\phi)$ for $i = 1, \dots, N$

2. Log-posterior:

$$\begin{aligned}\ell_{\text{MAP}}(\phi) &= \log [p(\mathcal{D} \mid \phi) f(\phi \mid \alpha, \beta)] \\&= \log \left[(\phi^{N_1} (1 - \phi)^{N_0}) \left(\frac{1}{B(\alpha, \beta)} \phi^{(\alpha-1)} (1 - \phi)^{(\beta-1)} \right) \right] \\&= \log \left[\phi^{(N_1 + \alpha - 1)} (1 - \phi)^{(N_0 + \beta - 1)} \frac{1}{B(\alpha, \beta)} \right] \\&= (N_1 + \alpha - 1) \log(\phi) + (N_0 + \beta - 1) \log(1 - \phi) - \log B(\alpha, \beta) \\&= N'_1 \log(\phi) + N'_0 \log(1 - \phi) - \log B(\alpha, \beta)\end{aligned}$$

$$N_1 = \#(x^{(i)} = 1)$$

$$N_0 = \#(x^{(i)} = 0)$$

3. Derivative: $\frac{\partial \ell_{\text{MAP}}(\phi)}{\partial \phi} = \frac{N'_1}{\phi} - \frac{N'_0}{1 - \phi}$

$$N'_1 = N_1 + \alpha - 1$$

$$N'_0 = N_0 + \beta - 1$$

4. Set to zero and solve: $\phi_{\text{MAP}} = \frac{N'_1}{N'_1 + N'_0} = \frac{N_1 + \alpha - 1}{N_1 + \alpha - 1 + N_0 + \beta - 1}$

MAP of Beta-Bernoulli Model

- Model: $\phi \sim \text{Beta}(\alpha, \beta)$
 $\mathbf{x}^{(i)} \sim \text{Bernoulli}(\phi)$ for $i = 1, \dots, N$

$$N_1 = \#(x^{(i)} = 1)$$

$$N_0 = \#(x^{(i)} = 0)$$

- Log-posterior:

$$\ell_{\text{MAP}}(\phi) = \log p(\mathcal{D} | \phi) + \log p(\phi)$$

$$= \log \left[\prod_{i=1}^N p(x^{(i)} | \phi) \right] + \log p(\phi)$$

$$= \log \prod_{i=1}^N \phi^{x^{(i)}} (1 - \phi)^{1-x^{(i)}} + \log \left[\frac{1}{B(\alpha, \beta)} \phi^{\alpha-1} (1-\phi)^{\beta-1} \right]$$

$$= \log (\phi^{N_1} (1 - \phi)^{N_0})$$

$$= N_1 \log(\phi) + N_0 \log(1 - \phi) + (\alpha - 1) \log \phi + (\beta - 1) \log(1 - \phi) + \log \frac{1}{B(\alpha, \beta)}$$

$$= (N_1 + \alpha - 1) \log \phi + (N_0 + \beta - 1) \log(1 - \phi)$$

MAP of Beta-Bernoulli Model

1. Model: $\phi \sim \text{Beta}(\alpha, \beta)$

$$\mathbf{x}^{(i)} \sim \text{Bernoulli}(\phi) \text{ for } i = 1, \dots, N$$

2. Log-posterior:

$$\ell_{\text{MAP}}(\phi) = \log [p(\mathcal{D} \mid \phi) f(\phi \mid \alpha, \beta)]$$

$$= \log \left[(\phi^{N_1} (1 - \phi)^{N_0}) \left(\frac{1}{B(\alpha, \beta)} \phi^{(\alpha-1)} (1 - \phi)^{(\beta-1)} \right) \right]$$

$$= \log \left[\phi^{(N_1 + \alpha - 1)} (1 - \phi)^{(N_0 + \beta - 1)} \frac{1}{B(\alpha, \beta)} \right]$$

$$= (N_1 + \alpha - 1) \log(\phi) + (N_0 + \beta - 1) \log(1 - \phi) - \log B(\alpha, \beta)$$

→ $= N'_1 \log(\phi) + N'_0 \log(1 - \phi) - \log B(\alpha, \beta)$

$$N_1 = \#(x^{(i)} = 1)$$

$$N_0 = \#(x^{(i)} = 0)$$

$$N'_1 = N_1 + \alpha - 1$$

$$N'_0 = N_0 + \beta - 1$$

3. Derivative: $\frac{\partial \ell_{\text{MAP}}(\phi)}{\partial \phi} = \frac{N'_1}{\phi} - \frac{N'_0}{1 - \phi}$

4. Set to zero and solve: $\phi_{\text{MAP}} = \frac{N'_1}{N'_1 + N'_0} = \frac{N_1 + \alpha - 1}{N_1 + \alpha - 1 + N_0 + \beta - 1}$

MAP of Beta-Bernoulli Model

Example 1 (MLE) Suppose $D = \{8H, 2T\}$

$$\phi_{\text{MLE}} = \frac{8}{10} = 0.8$$

Example 2 (MAP) Same dataset, but $\phi \sim \text{Beta}(\alpha = 101, \beta = 101)$

$$\phi_{\text{MAP}} = \frac{8 + 101 - 1}{8 + 101 - 1 + 2 + 101 - 1} = \frac{108}{108 + 102} \approx 0.5$$

“fair
coin”
prior

Example 3 (MAP) Same dataset, but $\phi \sim \text{Beta}(\alpha = 101, \beta = 1)$

$$\phi_{\text{MAP}} = \frac{108}{108 + 2} \approx 1.0$$

“unfair
coin”
prior

Example 4 (MLE) Suppose $D = \{108H, 102T\}$

$$\phi_{\text{MLE}} = \frac{108}{108 + 102} \approx 0.5$$

MLE for Linear Regression

$y \sim \mathcal{N}(\mu, \sigma^2)$

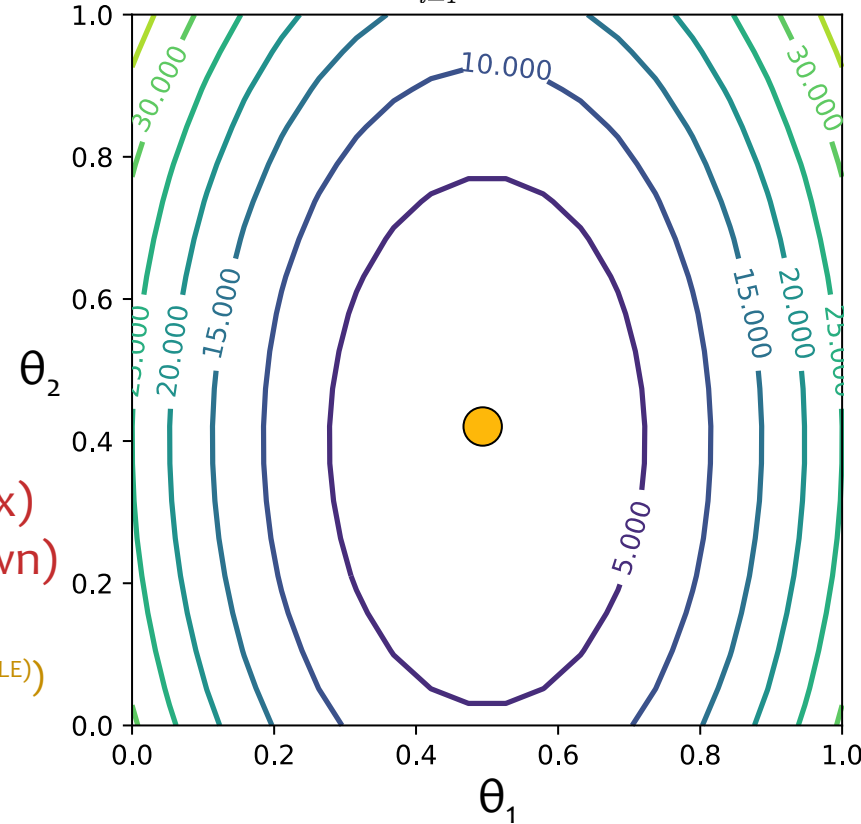
Optimization Method #2: Closed Form

1. Evaluate

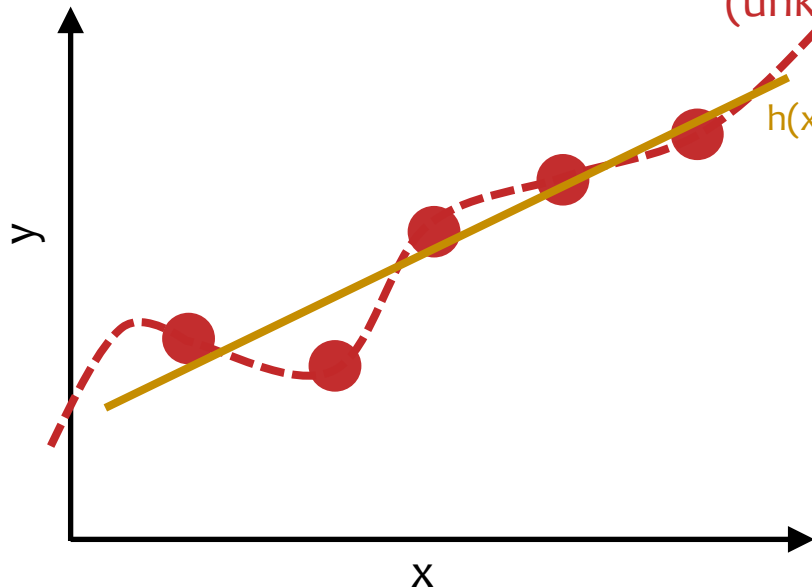
$$\theta^{\text{MLE}} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}$$

2. Return θ^{MLE}

$$J(\theta) = J(\theta_1, \theta_2) = \frac{1}{N} \sum_{i=1}^N (y^{(i)} - \theta^T \mathbf{x}^{(i)})^2$$



$y = h^*(x)$
(unknown)
 $h(x; \theta^{\text{MLE}})$



t	θ_1	θ_2	$J(\theta_1, \theta_2)$
MLE	0.59	0.43	0.2

1. You'll work through the view of linear regression as a probabilistic model in the homework!

2. You'll also see how L1 and L2 regularization is equivalent MAP estimation

Takeaways

- One view of what ML is trying to accomplish is **function approximation**
 - The principle of **maximum likelihood estimation** provides an alternate view of learning
- 
- **Synthetic data** can help **debug** ML algorithms
 - Probability distributions can be used to **model** real data that occurs in the world
(don't worry we'll make our distributions more interesting soon!)

Learning Objectives

MLE / MAP

You should be able to...

1. Recall probability basics, including but not limited to: discrete and continuous random variables, probability mass functions, probability density functions, events vs. random variables, expectation and variance, joint probability distributions, marginal probabilities, conditional probabilities, independence, conditional independence
2. Describe common probability distributions such as the Beta, Dirichlet, Multinomial, Categorical, Gaussian, Exponential, etc.
3. State the principle of maximum likelihood estimation and explain what it tries to accomplish
4. State the principle of maximum a posteriori estimation and explain why we use it
5. Derive the MLE or MAP parameters of a simple model in closed form

THE BIG PICTURE

ML Big Picture

Learning Paradigms:

What data is available and when? What form of prediction?

- supervised learning
- unsupervised learning
- semi-supervised learning
- reinforcement learning
- active learning
- imitation learning
- domain adaptation
- online learning
- density estimation
- recommender systems
- feature learning
- manifold learning
- dimensionality reduction
- ensemble learning
- distant supervision
- hyperparameter optimization

Theoretical Foundations:

What principles guide learning?

- ☐ probabilistic
- ☐ information theoretic
- ☐ evolutionary search
- ☐ ML as optimization

Problem Formulation:

What is the structure of our output prediction?

boolean	Binary Classification
categorical	Multiclass Classification
ordinal	Ordinal Classification
real	Regression
ordering	Ranking
multiple discrete	Structured Prediction
multiple continuous	(e.g. dynamical systems)
both discrete & cont.	(e.g. mixed graphical models)

Facets of Building ML Systems:

How to build systems that are robust, efficient, adaptive, effective?

1. Data prep
2. Model selection
3. Training (optimization / search)
4. Hyperparameter tuning on validation data
5. (Blind) Assessment on test data

Big Ideas in ML:

Which are the ideas driving development of the field?

- inductive bias
- generalization / overfitting
- bias-variance decomposition
- generative vs. discriminative
- deep nets, graphical models
- PAC learning
- distant rewards

Application Areas

Key challenges?

NLP, Speech, Computer Vision, Robotics, Medicine, Search

Classification and Regression: The Big Picture

Recipe for Machine Learning

1. Given data $\mathcal{D} = \{\mathbf{x}^{(i)}, y^{(i)}\}_{i=1}^N$
2. (a) Choose a decision function $h_{\theta}(\mathbf{x}) = \dots$
(parameterized by θ)
(b) Choose an objective function $J_{\mathcal{D}}(\theta) = \dots$
(relies on data)
3. Learn by choosing parameters that optimize the objective $J_{\mathcal{D}}(\theta)$

$$\hat{\theta} \approx \underset{\theta}{\operatorname{argmin}} J_{\mathcal{D}}(\theta)$$

4. Predict on new test example $\mathbf{x}_{\text{new}}$ using $h_{\theta}(\cdot)$

$$\hat{y} = h_{\theta}(\mathbf{x}_{\text{new}})$$

Optimization Method

- Gradient Descent: $\theta \rightarrow \theta - \gamma \nabla_{\theta} J(\theta)$
- SGD: $\theta \rightarrow \theta - \gamma \nabla_{\theta} J^{(i)}(\theta)$
for $i \sim \text{Uniform}(1, \dots, N)$
where $J(\theta) = \frac{1}{N} \sum_{i=1}^N J^{(i)}(\theta)$
- mini-batch SGD
- closed form
 1. compute partial derivatives
 2. set equal to zero and solve

Decision Functions

- Perceptron: $h_{\theta}(\mathbf{x}) = \text{sign}(\theta^T \mathbf{x})$
- Linear Regression: $h_{\theta}(\mathbf{x}) = \theta^T \mathbf{x}$
- Discriminative Models: $h_{\theta}(\mathbf{x}) = \underset{y}{\operatorname{argmax}} p_{\theta}(y | \mathbf{x})$
 - Logistic Regression: $p_{\theta}(y = 1 | \mathbf{x}) = \sigma(\theta^T \mathbf{x})$
 - Neural Net (classification):
 $p_{\theta}(y = 1 | \mathbf{x}) = \sigma((\mathbf{W}^{(2)})^T \sigma((\mathbf{W}^{(1)})^T \mathbf{x} + \mathbf{b}^{(1)}) + \mathbf{b}^{(2)})$
- Generative Models: $h_{\theta}(\mathbf{x}) = \underset{y}{\operatorname{argmax}} p_{\theta}(\mathbf{x}, y)$
 - Naive Bayes: $p_{\theta}(\mathbf{x}, y) = p_{\theta}(y) \prod_{m=1}^M p_{\theta}(x_m | y)$

Objective Function

- MLE: $J(\theta) = -\sum_{i=1}^N \log p(\mathbf{x}^{(i)}, \mathbf{y}^{(i)})$
- MCLE: $J(\theta) = -\sum_{i=1}^N \log p(\mathbf{y}^{(i)} | \mathbf{x}^{(i)})$
- L2 Regularized: $J'(\theta) = J(\theta) + \lambda \|\theta\|_2^2$
(same as Gaussian prior $p(\theta)$ over parameters)
- L1 Regularized: $J'(\theta) = J(\theta) + \lambda \|\theta\|_1$
(same as Laplace prior $p(\theta)$ over parameters)

Backpropagation and Deep Learning

Convolutional neural networks (CNNs) and **recurrent neural networks** (RNNs) are simply fancy computation graphs (aka. hypotheses or decision functions).

Our recipe also applies to these models and (again) relies on the **backpropagation algorithm** to compute the necessary gradients.

BACKGROUND: COMPUTER VISION

Example: Image Classification

- ImageNet LSVRC-2011 contest:
 - **Dataset:** 1.2 million labeled images, 1000 classes
 - **Task:** Given a new image, label it with the correct class
 - **Multiclass** classification problem
- Examples from <http://image-net.org/>

Bird

Warm-blooded egg-laying vertebrates characterized by feathers and forelimbs modified as wings

2126
pictures

92.85%
Popularity
Percentile

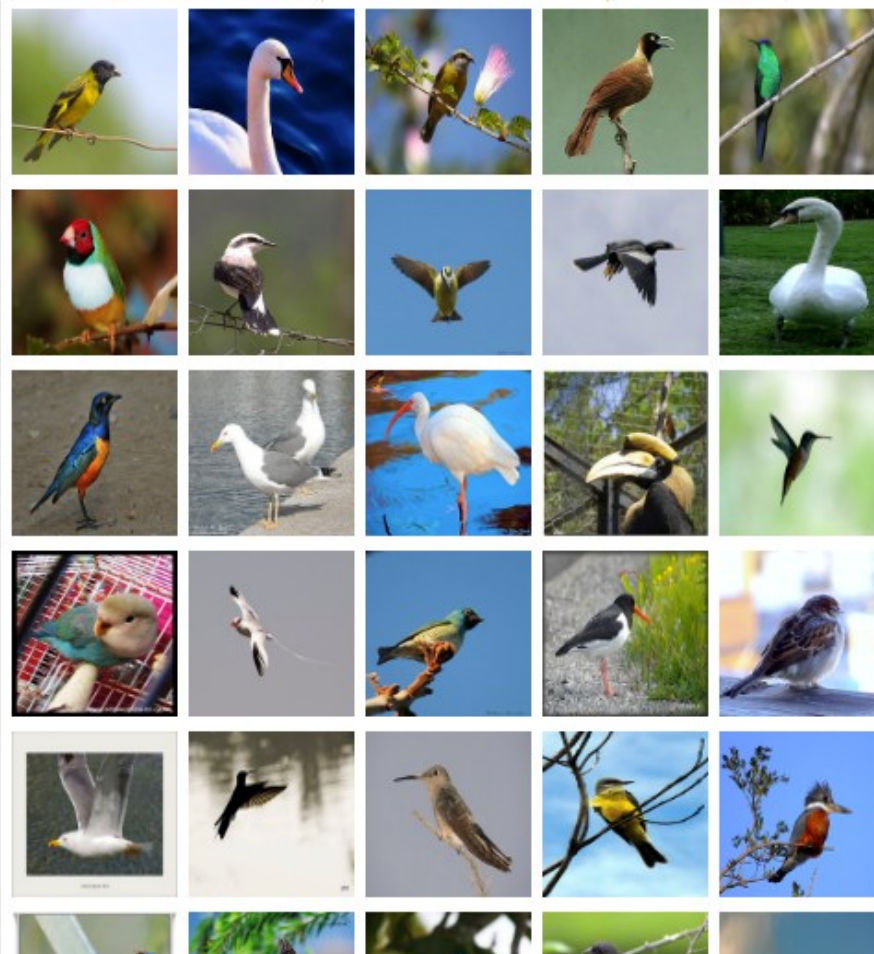


- marine animal, marine creature, sea animal, sea creature (1)
- scavenger (1)
- biped (0)
- predator, predatory animal (1)
- larva (49)
- acrodont (0)
- feeder (0)
- stunt (0)
- chordate (3087)
 - tunicate, urochordate, urochord (6)
 - cephalochordate (1)
 - vertebrate, craniate (3077)
 - mammal, mammalian (1169)
 - bird (871)
 - dickeybird, dickey-bird, dickybird, dicky-bird (0)
 - cock (1)
 - hen (0)
 - nester (0)
 - night bird (1)
 - bird of passage (0)
 - protoavis (0)
 - archaeopteryx, archeopteryx, Archaeopteryx lithographi Sinornis (0)
 - Ibero-mesornis (0)
 - archaeornis (0)
 - ratite, ratite bird, flightless bird (10)
 - carinate, carinate bird, flying bird (0)
 - passerine, passeriform bird (279)
 - nonpasserine bird (0)
 - bird of prey, raptor, raptorial bird (80)
 - gallinaceous bird, gallinacean (114)

Treemap Visualization

Images of the Synset

Downloads



German iris, *Iris kochii*Iris of northern Italy having deep blue-purple flowers; similar to but smaller than *Iris germanica*469
pictures49.6%
Popularity
Percentile

- ... halophyte (0)
- ▶ succulent (39)
- ... cultivar (0)
- ... cultivated plant (0)
- ▶ weed (54)
- ... evergreen, evergreen plant (0)
- ... deciduous plant (0)
- ▶ vine (272)
- ... creeper (0)
- ▶ woody plant, ligneous plant (1868)
- ... geophyte (0)
- ▶ desert plant, xerophyte, xerophytic plant, xerophile, xerophilic
- ... mesophyte, mesophytic plant (0)
- ▶ aquatic plant, water plant, hydrophyte, hydrophytic plant (11)
- ... tuberous plant (0)
- ▶ bulbous plant (179)
- ▶ iridaceous plant (27)
 - ▶ iris, flag, fleur-de-lis, sword lily (19)
 - ▶ bearded iris (4)
 - ... Florentine iris, orris, *Iris germanica florentina*, Iris
 - ... German iris, *Iris germanica* (0)
 - ▶ German iris, *Iris kochii* (0)
 - ... Dalmatian iris, *Iris pallida* (0)
 - ▶ beardless iris (4)
 - ... bulbous iris (0)
 - ... dwarf iris, *Iris cristata* (0)
 - ... stinking iris, gladdon, gladdon iris, stinking gladdyn,
 - ... Persian iris, *Iris persica* (0)
 - ... yellow iris, yellow flag, yellow water flag, *Iris pseudacorus*
 - ... dwarf iris, vernal iris, *Iris verna* (0)
 - ... blue flag, *Iris versicolor* (0)

Treemap Visualization

Images of the Synset

Downloads



Court, courtyard

An area wholly or partly surrounded by walls or buildings; "the house was built around an inner court"

165
pictures

92.61%
Popularity
Percentile



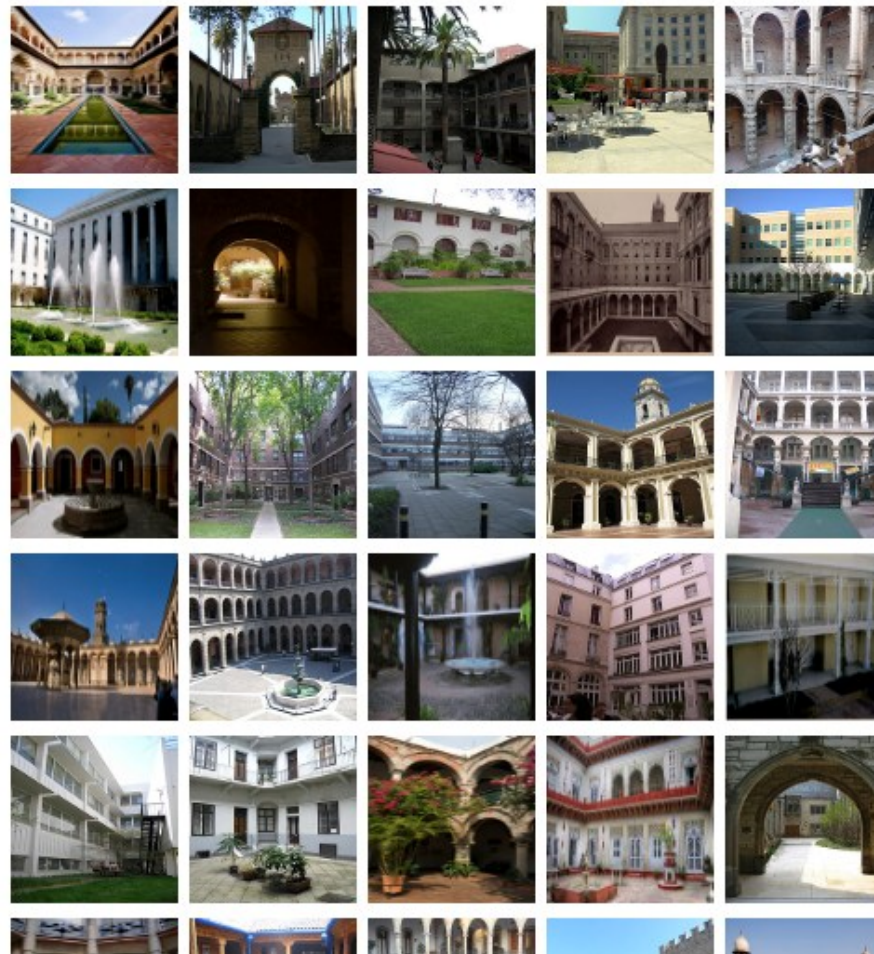
Numbers in brackets: (the number of synsets in the subtree).

- ImageNet 2011 Fall Release (32326)
 - plant, flora, plant life (4486)
 - geological formation, formation (175)
 - natural object (1112)
 - sport, athletics (176)
 - artifact, artefact (10504)
 - instrumentality, instrumentation (5494)
 - structure, construction (1405)
 - airdock, hangar, repair shed (0)
 - altar (1)
 - arcade, colonnade (1)
 - arch (31)
 - area (344)
 - aisle (0)
 - auditorium (1)
 - baggage claim (0)
 - box (1)
 - breakfast area, breakfast nook (0)
 - bullpen (0)
 - chancel, sanctuary, bema (0)
 - choir (0)
 - corner, nook (2)
 - court, courtyard (6)
 - atrium (0)
 - bailey (0)
 - cloister (0)
 - food court (0)
 - forecourt (0)
 - narvis (0)

Treemap Visualization

Images of the Synset

Downloads

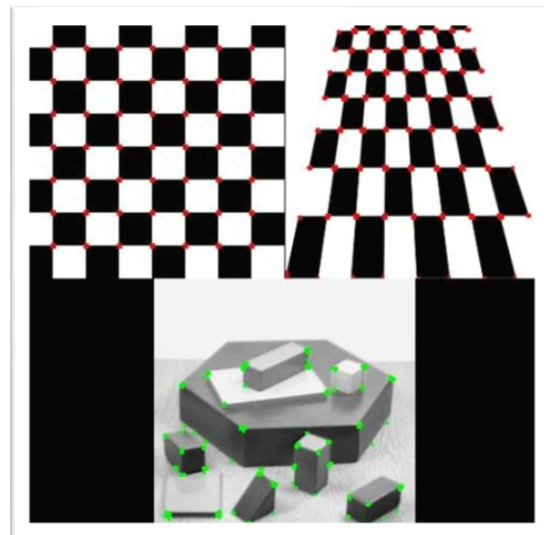


Feature Engineering for CV

Edge detection (Canny)



Corner Detection (Harris)



Figures from <http://opencv.org>

Scale Invariant Feature Transform (SIFT)



Figure 3: Model images of planar objects are shown in the top row. Recognition results below show model outlines and match keys used for matching.

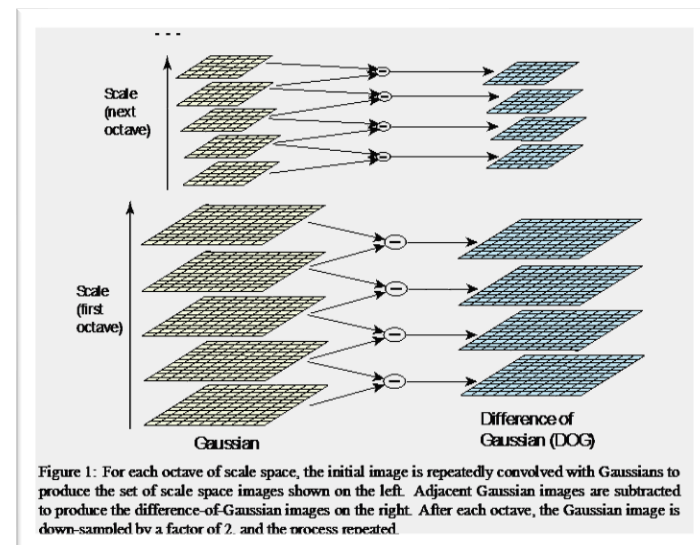


Figure from Lowe (1999) and Lowe (2004)

Example: Image Classification

1998
LeNet-5
Yann LeCun

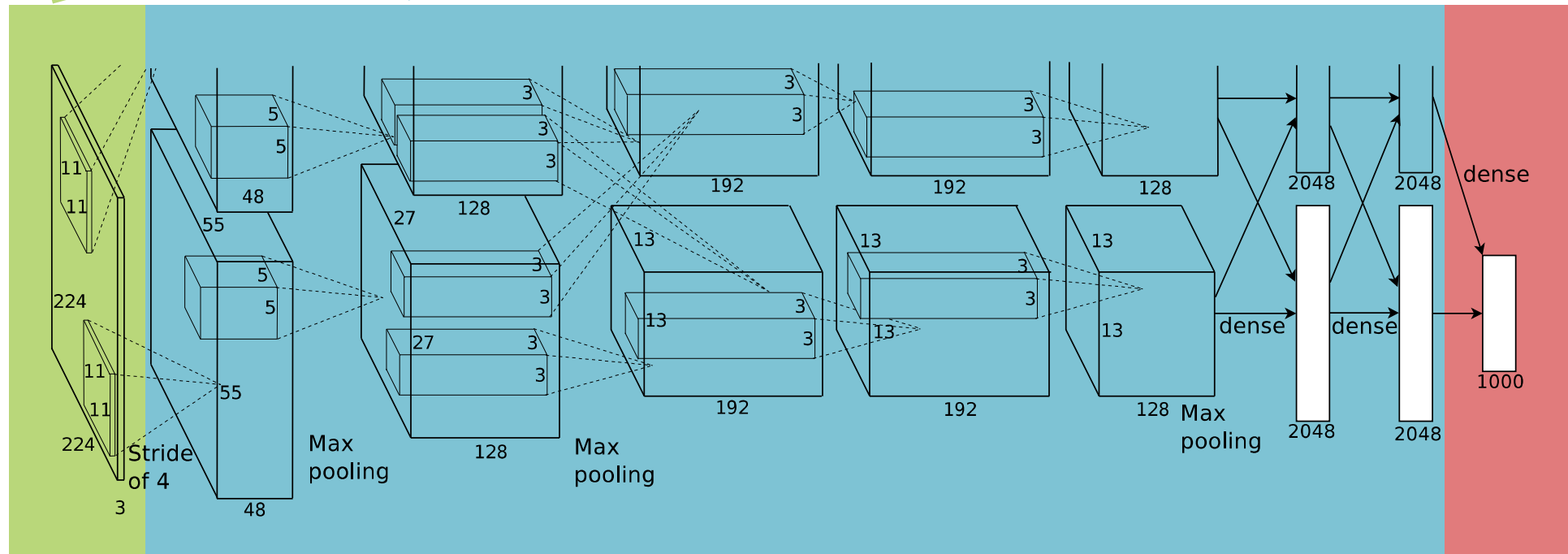
2012

AlexNet – a CNN for Image Classification
(Krizhevsky, Sutskever & Hinton, 2012)
15.3% error on ImageNet LSVRC-2012 contest

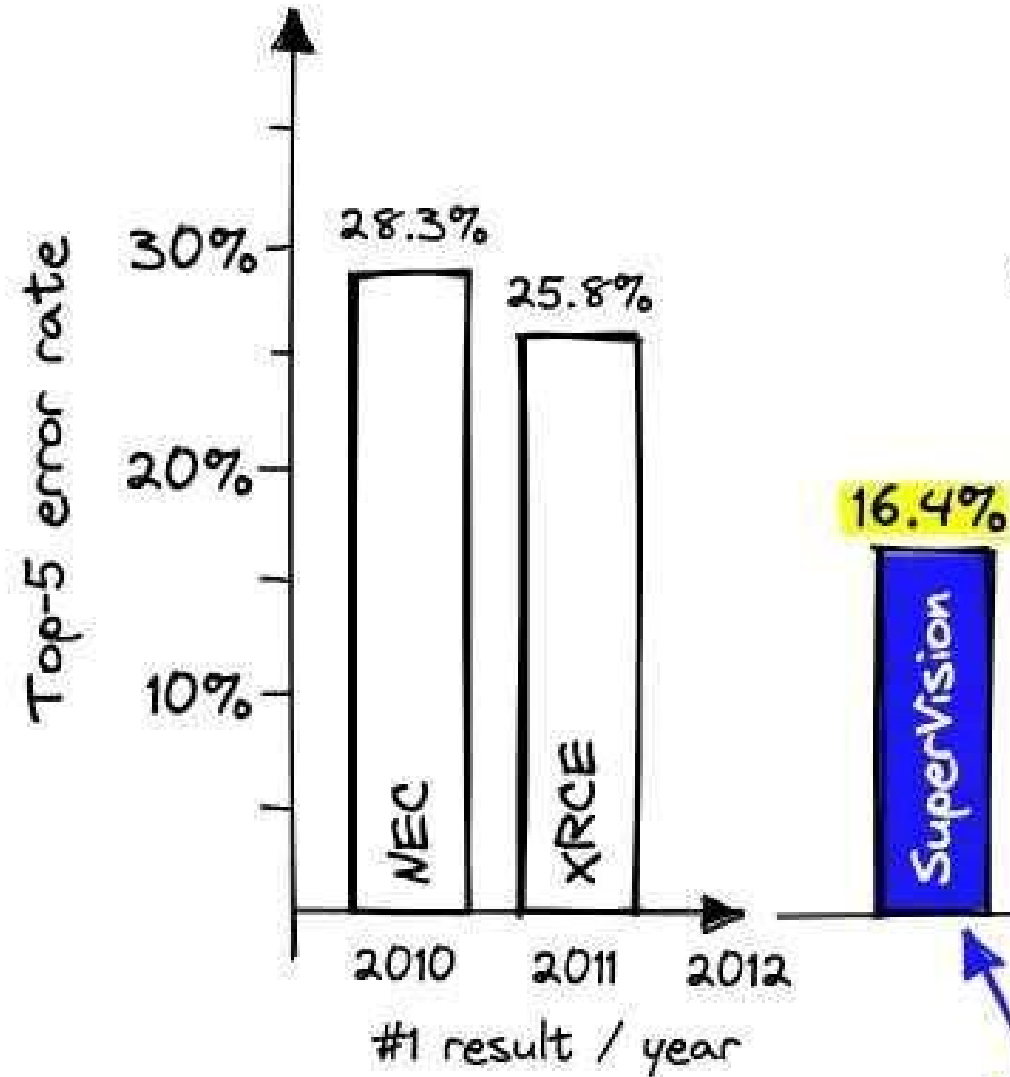
Input
image
(pixels)

- Five convolutional layers (w/max-pooling)
- Three fully connected layers

1000-way
softmax



CNNs for Image Recognition



2 3 4 5

6

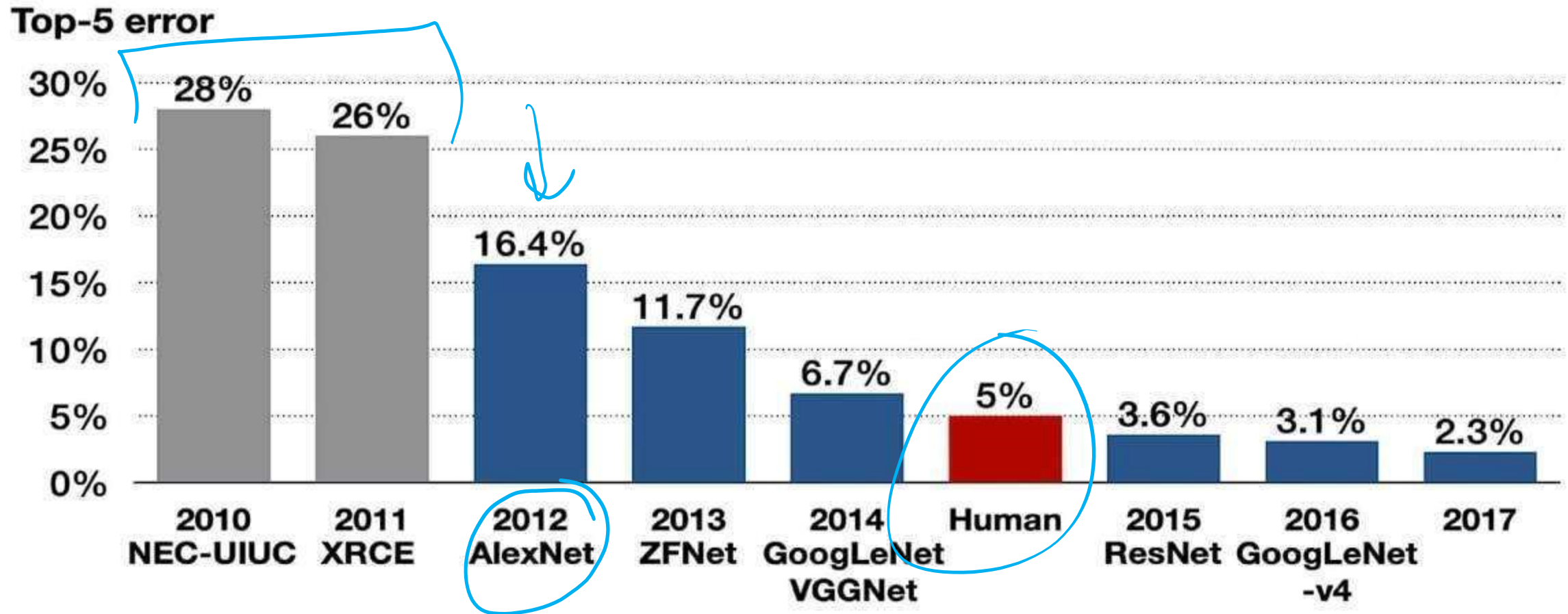
CNNs for Image Recognition



CNNs for Image Recognition



CNNs for Image Recognition



Feed-forward Neural Networks for Computer Vision

Feed-forward Neural Networks for Computer Vision

CONVOLUTIONAL NEURAL NETS

A Recipe for Machine Learning

1. Given training data:

$$\mathcal{D} = \{\mathbf{x}^{(i)}, y^{(i)}\}_{i=1}^N$$

2. Choose each of these:

- Decision function

$$\hat{y} = h_{\theta}(\mathbf{x})$$

- Loss function

$$\ell(\hat{y}, y) \in \mathbb{R}$$

3. Define goal:

$$\hat{\theta} \approx \operatorname{argmin}_{\theta} \sum_{i=1}^N \ell(h_{\theta}(\mathbf{x}^{(i)}), y^{(i)})$$

4. Train with SGD:

(take small steps
opposite the gradient)

$$\theta^{(t+1)} = \theta^{(t)} - \eta_t \nabla \ell(h_{\theta}(\mathbf{x}^{(i)}), y^{(i)})$$

A Recipe for Machine Learning

1. • Convolutional Neural Networks (CNNs) provide another form of **decision function**
• Let's see what they look like...

2. Choose each of these:

– Decision function

$$\hat{y} = h_{\theta}(\mathbf{x})$$

– Loss function

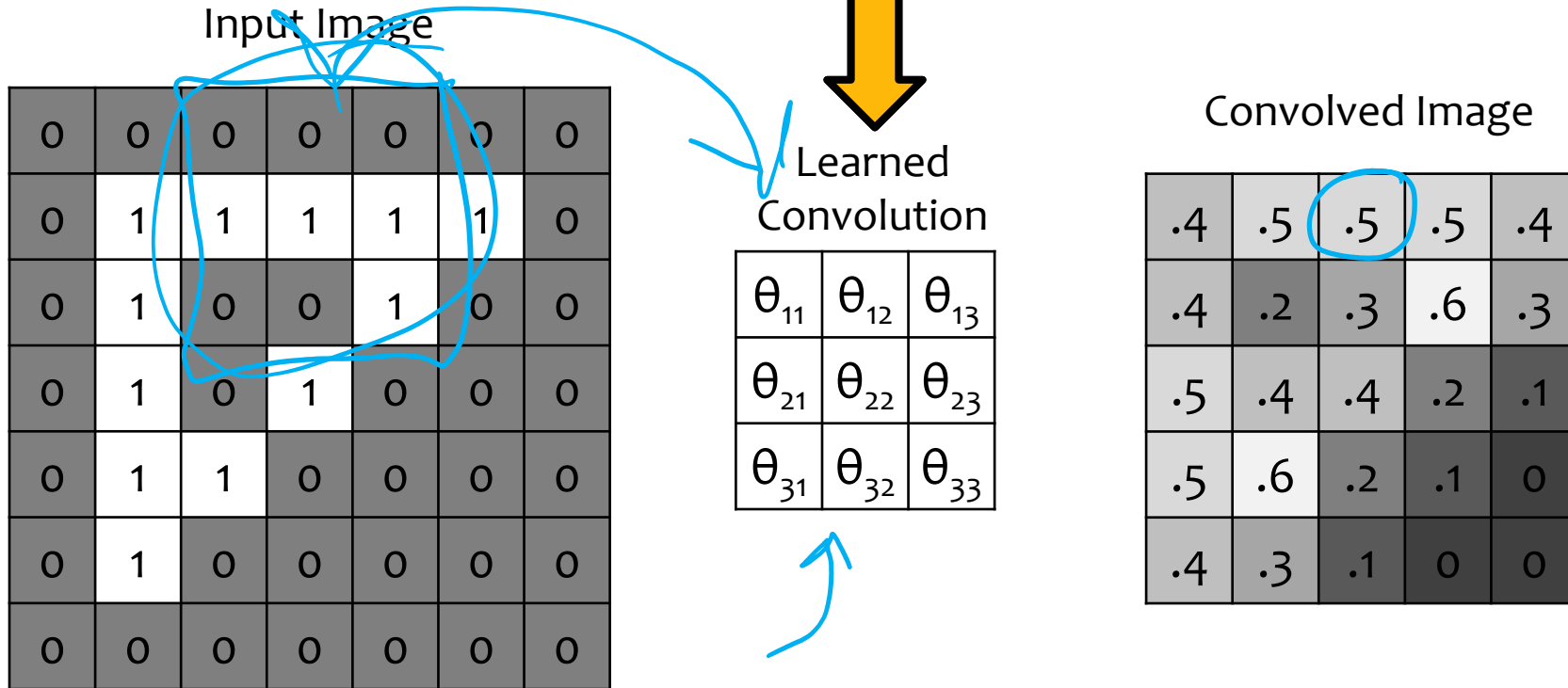
$$\ell(\hat{y}, y) \in \mathbb{R}$$

4. Train with SGD:
(take small steps opposite the gradient)

$$\theta^{(t+1)} = \theta^{(t)} - \eta \nabla \ell(h_{\theta}(\mathbf{x}^{(t)}), y^{(t)})$$

Convolutional Layer

CNN key idea:
Treat convolution matrix as
parameters and learn them!

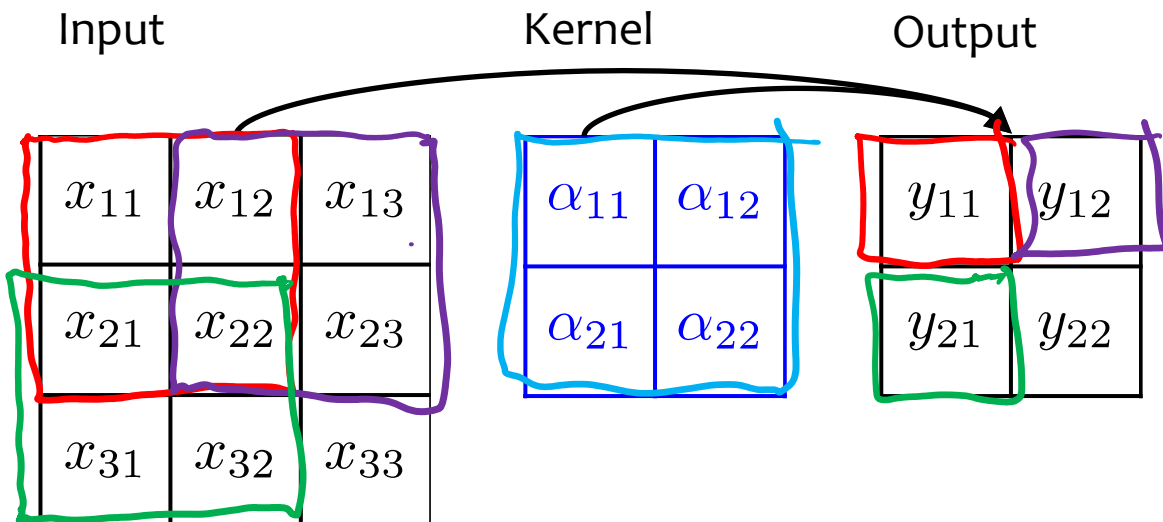


CONVOLUTION

2D Convolution

- Basic idea:
 - Pick a 2x2 matrix F of weights (called a kernel or convolution matrix)
 - Slide this over an image and compute the “inner product” (similarity) of F and the corresponding field of the image, and replace the pixel in the center of the field with the output of the inner product operation
- Key point:
 - Different convolutions extract different types of low-level “features” from an image
 - All that we need to vary to generate these different features is the weights of F

Example: 1 input channel, 1 output channel



$$\begin{aligned} y_{11} &= \alpha_{11}x_{11} + \alpha_{12}x_{12} + \alpha_{21}x_{21} + \alpha_{22}x_{22} + \alpha_0 \\ y_{12} &= \alpha_{11}x_{12} + \alpha_{12}x_{13} + \alpha_{21}x_{22} + \alpha_{22}x_{23} + \alpha_0 \\ y_{21} &= \alpha_{11}x_{21} + \alpha_{12}x_{22} + \alpha_{21}x_{31} + \alpha_{22}x_{32} + \alpha_0 \\ y_{22} &= \alpha_{11}x_{22} + \alpha_{12}x_{23} + \alpha_{21}x_{32} + \alpha_{22}x_{33} + \alpha_0 \end{aligned}$$

2D Convolution

- Pick a 2x2 matrix F of weights (called a kernel or convolution matrix)
- Slide this over an image and compute the “inner product” (similarity) of F and the corresponding field of the image, and replace the pixel in the center of the field with the output of the inner product operation

Input Image

0	0	0	0	0	0	0
0	1	1	1	1	1	0
0	1	0	0	1	0	0
0	1	0	1	0	0	0
0	1	1	0	0	0	0
0	1	0	0	0	0	0
0	0	0	0	0	0	0

Convolution

0	0	0
0	1	1
0	1	0

Convolved Image

3	2	2	3	1
2	0	2	1	0
2	2	1	0	0
3	1	0	0	0
1	0	0	0	0

2D Convolution

- Pick a 2x2 matrix F of weights (called a kernel or convolution matrix)
- Slide this over an image and compute the “inner product” (similarity) of F and the corresponding field of the image, and replace the pixel in the center of the field with the output of the inner product operation

Input Image

0	0	0	0	0	0	0
0	1	1	1	1	1	0
0	1	0	0	1	0	0
0	1	0	1	0	0	0
0	1	1	0	0	0	0
0	1	0	0	0	0	0
0	0	0	0	0	0	0

Convolution

0	0	0
0	1	1
0	1	0

Convolved Image

3	2	2	3	1
2	0	2	1	0
2	2	1	0	0
3	1	0	0	0
1	0	0	0	0

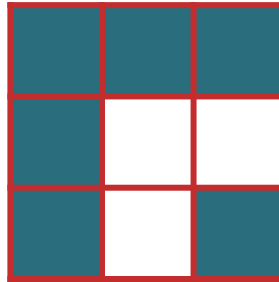
2D Convolution

- Pick a 2x2 matrix F of weights (called a kernel or convolution matrix)
- Slide this over an image and compute the “inner product” (similarity) of F and the corresponding field of the image, and replace the pixel in the center of the field with the output of the inner product operation

Input Image

0	0	0	0	0	0	0
0	1	1	1	1	1	0
0	1	0	0	1	0	0
0	1	0	1	0	0	0
0	1	1	0	0	0	0
0	1	0	0	0	0	0
0	0	0	0	0	0	0

Convolution



Convolved Image

3	2	2	3	1
2	0	2	1	0
2	2	1	0	0
3	1	0	0	0
1	0	0	0	0

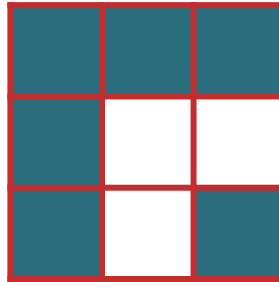
2D Convolution

- Pick a 2x2 matrix F of weights (called a kernel or convolution matrix)
- Slide this over an image and compute the “inner product” (similarity) of F and the corresponding field of the image, and replace the pixel in the center of the field with the output of the inner product operation

Input Image

0	0	0	0	0	0	0
0	1	1	1	1	1	0
0	1	0	0	1	0	0
0	1	0	1	0	0	0
0	1	1	0	0	0	0
0	1	0	0	0	0	0
0	0	0	0	0	0	0

Convolution



Convolved Image

3	2	2	3	1
2	0	2	1	0
2	2	1	0	0
3	1	0	0	0
1	0	0	0	0

2D Convolution

- Pick a 2x2 matrix F of weights (called a kernel or convolution matrix)
- Slide this over an image and compute the “inner product” (similarity) of F and the corresponding field of the image, and replace the pixel in the center of the field with the output of the inner product operation

Input Image

			0	0	0	0
	1	1	1	1	1	0
	1		0	1	0	0
0	1	0	1	0	0	0
0	1	1	0	0	0	0
0	1	0	0	0	0	0
0	0	0	0	0	0	0

Convolution

Convolved Image

3				

2D Convolution

- Pick a 2x2 matrix F of weights (called a kernel or convolution matrix)
- Slide this over an image and compute the “inner product” (similarity) of F and the corresponding field of the image, and replace the pixel in the center of the field with the output of the inner product operation

Input Image

0				0	0	0
0		1	1	1	1	0
0		0		1	0	0
0	1	0	1	0	0	0
0	1	1	0	0	0	0
0	1	0	0	0	0	0
0	0	0	0	0	0	0

Convolution

Convolved Image

3	2			

2D Convolution

- Pick a 2x2 matrix F of weights (called a kernel or convolution matrix)
- Slide this over an image and compute the “inner product” (similarity) of F and the corresponding field of the image, and replace the pixel in the center of the field with the output of the inner product operation

Input Image

0	0				0	0
0	1		1	1	1	0
0	1		0		0	0
0	1	0	1	0	0	0
0	1	1	0	0	0	0
0	1	0	0	0	0	0
0	0	0	0	0	0	0

Convolution

Convolved Image

3	2	2		

2D Convolution

- Pick a 2x2 matrix F of weights (called a kernel or convolution matrix)
- Slide this over an image and compute the “inner product” (similarity) of F and the corresponding field of the image, and replace the pixel in the center of the field with the output of the inner product operation

Input Image

0	0	0				0
0	1	1		1	1	0
0	1	0		1		0
0	1	0	1	0	0	0
0	1	1	0	0	0	0
0	1	0	0	0	0	0
0	0	0	0	0	0	0

Convolution

Convolved Image

3	2	2	3	

2D Convolution

- Pick a 2x2 matrix F of weights (called a kernel or convolution matrix)
- Slide this over an image and compute the “inner product” (similarity) of F and the corresponding field of the image, and replace the pixel in the center of the field with the output of the inner product operation

Input Image

0	0	0	0			
0	1	1	1		1	0
0	1	0	0		0	
0	1	0	1	0	0	0
0	1	1	0	0	0	0
0	1	0	0	0	0	0
0	0	0	0	0	0	0

Convolution

Convolved Image

3	2	2	3	1

2D Convolution

- Pick a 2x2 matrix F of weights (called a kernel or convolution matrix)
- Slide this over an image and compute the “inner product” (similarity) of F and the corresponding field of the image, and replace the pixel in the center of the field with the output of the inner product operation

Input Image

0	0	0	0	0	0	0
0	0	0	1	1	1	0
0	1	0	0	1	0	0
0	1	1	1	0	0	0
0	1	0	0	0	0	0
0	0	0	0	0	0	0

Convolution

0	0	0
0	1	1
0	1	0

Convolved Image

3	2	2	3	1
2	0	0	0	0
0	0	0	0	0
0	0	0	0	0
0	0	0	0	0

2D Convolution

- Pick a 2x2 matrix F of weights (called a kernel or convolution matrix)
- Slide this over an image and compute the “inner product” (similarity) of F and the corresponding field of the image, and replace the pixel in the center of the field with the output of the inner product operation

Input Image

0	0	0	0	0	0	0
0				1	1	0
0			0	0	1	0
0			0		0	0
0	1	1	0	0	0	0
0	1	0	0	0	0	0
0	0	0	0	0	0	0

Convolution

Convolved Image

3	2	2	3	1
2	0			

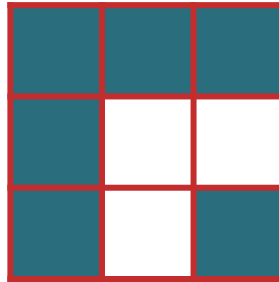
2D Convolution

- Pick a 2x2 matrix F of weights (called a kernel or convolution matrix)
- Slide this over an image and compute the “inner product” (similarity) of F and the corresponding field of the image, and replace the pixel in the center of the field with the output of the inner product operation

Input Image

0	0	0	0	0	0	0
0	1	1	1	1	1	0
0	1	0	0	1	0	0
0	1	0	1	0	0	0
0	1	1	0	0	0	0
0	1	0	0	0	0	0
0	0	0	0	0	0	0

Convolution



Convolved Image

3	2	2	3	1
2	0	2	1	0
2	2	1	0	0
3	1	0	0	0
1	0	0	0	0

PADDING

Padding

Suppose you want to preserve the size of the original input image in your convolved image.

You can accomplish this by padding your input image with zeros.

Input Image

0	0	0	0	0	0	0
0	1	1	1	1	1	0
0	1	0	0	1	0	0
0	1	0	1	0	0	0
0	1	1	0	0	0	0
0	1	0	0	0	0	0
0	0	0	0	0	0	0

Identity
Convolution

0	0	0
0	1	0
0	0	0

Convolved Image

1	1	1	1	1
1	0	0	1	0
1	0	1	0	0
1	1	0	0	0
1	0	0	0	0

Padding

Suppose you want to preserve the size of the original input image in your convolved image.

You can accomplish this by padding your input image with zeros.

Input Image

0	0	0	0	0	0	0	0	0
0	0	0	0	0	0	0	0	0
0	0	1	1	1	1	1	0	0
0	0	1	0	0	1	0	0	0
0	0	1	0	1	0	0	0	0
0	0	1	1	0	0	0	0	0
0	0	1	0	0	0	0	0	0
0	0	0	0	0	0	0	0	0
0	0	0	0	0	0	0	0	0

Identity
Convolution

0	0	0
0	1	0
0	0	0

Convolved Image

0	0	0	0	0	0	0
0	1	1	1	1	1	0
0	1	0	0	1	0	0
0	1	0	1	0	0	0
0	1	1	0	0	0	0
0	1	0	0	0	0	0
0	0	0	0	0	0	0

Kernels for Image Processing

A **convolution matrix** (aka. kernel) is used in image processing for tasks such as edge detection, blurring, sharpening, etc.

Input Image

0	0	0	0	0	0	0	0	0
0	0	0	0	0	0	0	0	0
0	0	1	1	1	1	1	0	0
0	0	1	0	0	1	0	0	0
0	0	1	0	1	0	0	0	0
0	0	1	1	0	0	0	0	0
0	0	1	0	0	0	0	0	0
0	0	0	0	0	0	0	0	0
0	0	0	0	0	0	0	0	0

Identity
Convolution

0	0	0
0	1	0
0	0	0

Convolved Image

0	0	0	0	0	0	0
0	1	1	1	1	1	0
0	1	0	0	1	0	0
0	1	0	1	0	0	0
0	1	1	0	0	0	0
0	1	0	0	0	0	0
0	0	0	0	0	0	0

Kernels for Image Processing

A **convolution matrix** (aka. kernel) is used in image processing for tasks such as edge detection, blurring, sharpening, etc.

Input Image

0	0	0	0	0	0	0	0	0
0	0	0	0	0	0	0	0	0
0	0	1	1	1	1	1	0	0
0	0	1	0	0	1	0	0	0
0	0	1	0	1	0	0	0	0
0	0	1	1	0	0	0	0	0
0	0	1	0	0	0	0	0	0
0	0	0	0	0	0	0	0	0
0	0	0	0	0	0	0	0	0

Blurring
Convolution

.1	.1	.1
.1	.2	.1
.1	.1	.1

Convolved Image

.1	.2	.3	.3	.3	.2	.1
.2	.4	.5	.5	.5	.4	.1
.3	.4	.2	.3	.6	.3	.1
.3	.5	.4	.4	.2	.1	0
.3	.5	.6	.2	.1	0	0
.2	.4	.3	.1	0	0	0
.1	.1	.1	0	0	0	0

Kernels for Image Processing

A **convolution matrix** (aka. kernel) is used in image processing for tasks such as edge detection, blurring, sharpening, etc.

Input Image

0	0	0	0	0	0	0	0	0
0	0	0	0	0	0	0	0	0
0	0	1	1	1	1	1	0	0
0	0	1	0	0	1	0	0	0
0	0	1	0	1	0	0	0	0
0	0	1	1	0	0	0	0	0
0	0	1	0	0	0	0	0	0
0	0	0	0	0	0	0	0	0
0	0	0	0	0	0	0	0	0

Vertical
Edge
Detector

-1	0	1
-1	0	1
-1	0	1

Convolved Image

-1	-1	0	0	0	1	1
-2	-1	1	-1	0	2	1
-3	-1	1	-1	1	2	1
-3	-1	2	0	1	1	0
-3	-1	2	1	1	0	0
-2	-1	2	1	0	0	0
-1	0	1	0	0	0	0

Kernels for Image Processing

A **convolution matrix** (aka. kernel) is used in image processing for tasks such as edge detection, blurring, sharpening, etc.

Input Image

0	0	0	0	0	0	0	0	0
0	0	0	0	0	0	0	0	0
0	0	1	1	1	1	1	0	0
0	0	1	0	0	1	0	0	0
0	0	1	0	1	0	0	0	0
0	0	1	1	0	0	0	0	0
0	0	1	0	0	0	0	0	0
0	0	0	0	0	0	0	0	0
0	0	0	0	0	0	0	0	0

Horizontal
Edge
Detector

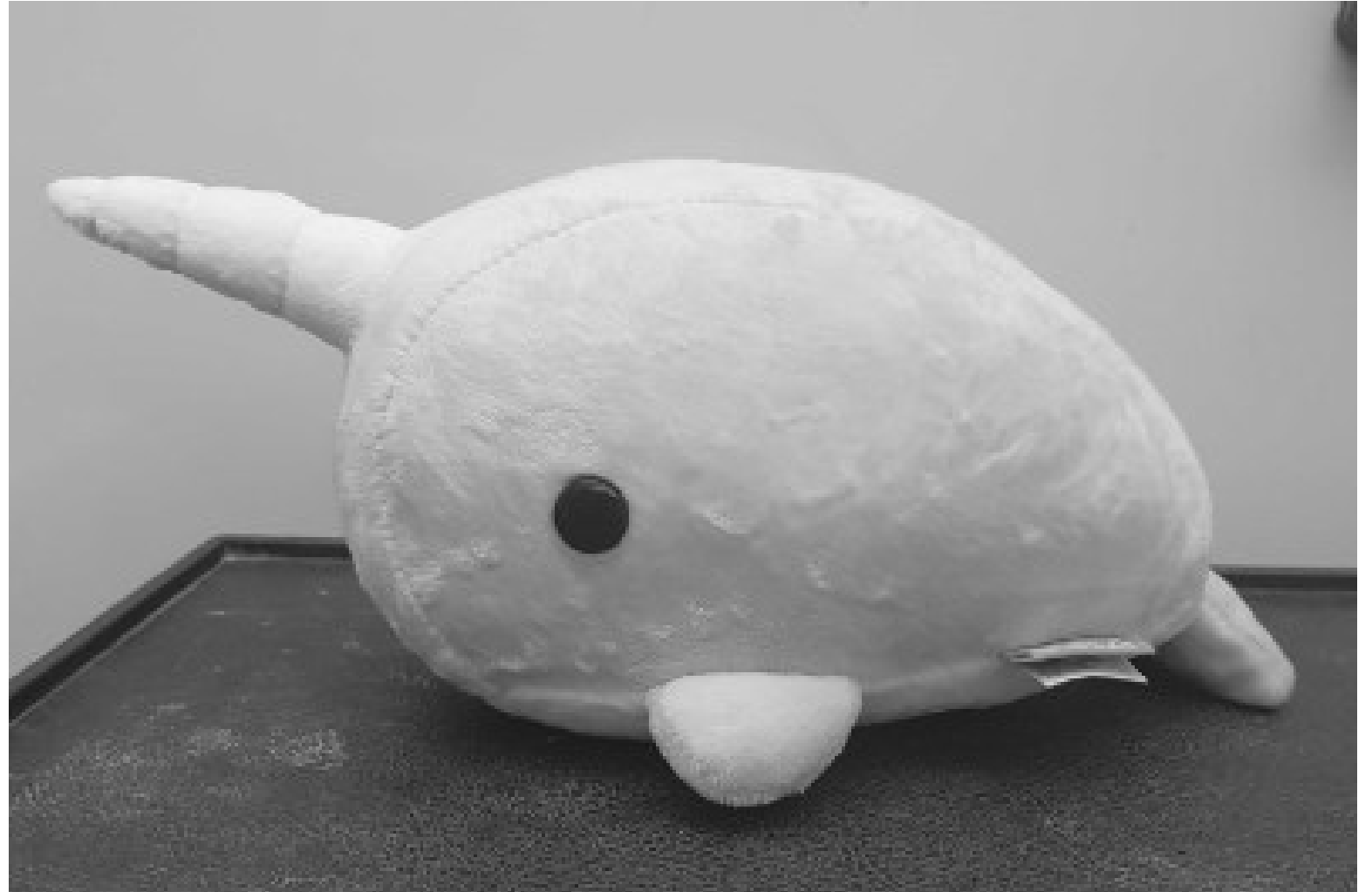
-1	-1	-1
0	0	0
1	1	1

Convolved Image

-1	-2	-3	-3	-3	-2	-1
-1	-1	-1	-1	-1	-1	0
0	1	1	2	2	2	1
0	-1	-1	0	1	1	0
0	0	1	1	1	0	0
1	2	2	1	0	0	0
1	1	1	0	0	0	0

Convolution Examples

Original
Image



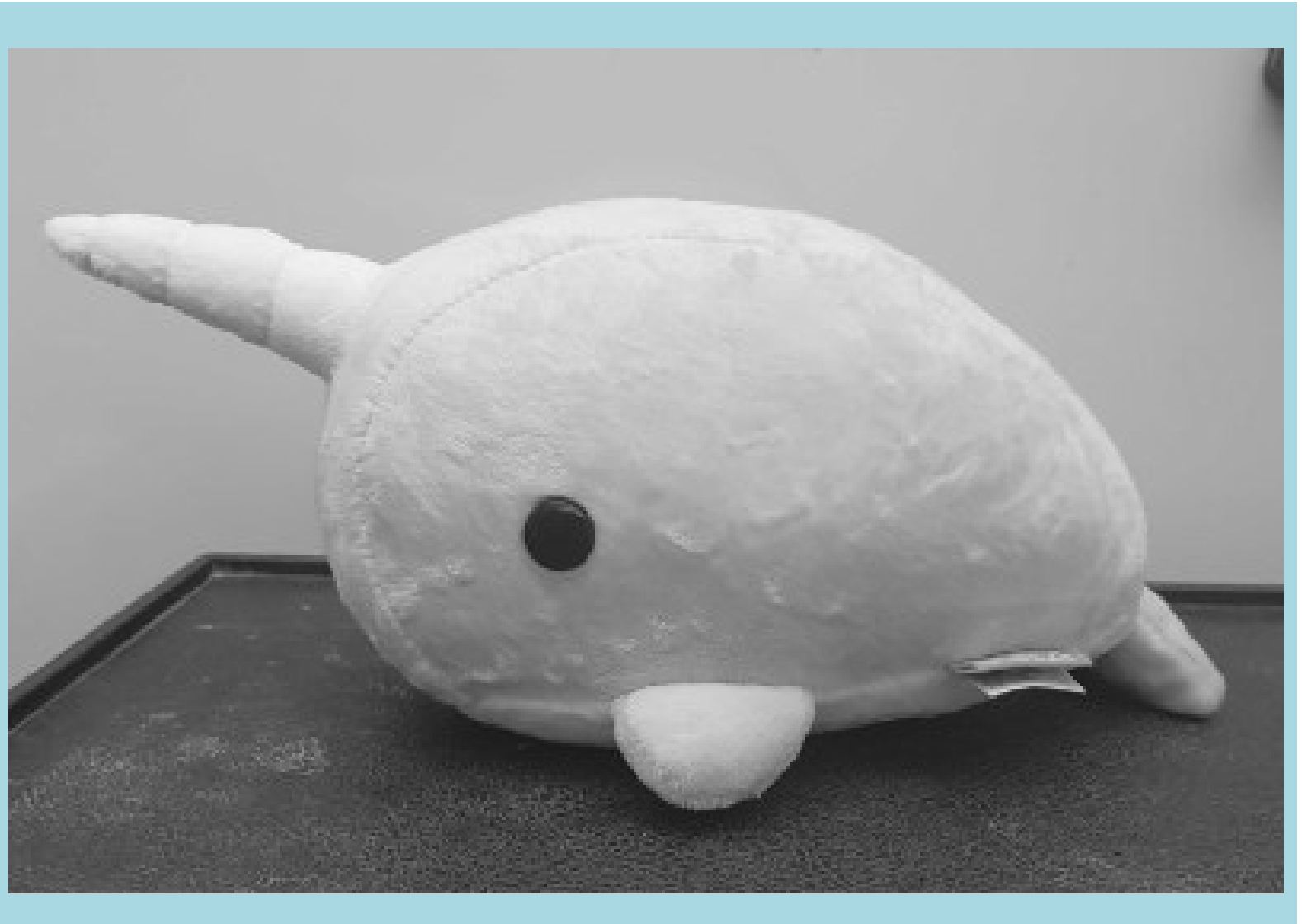
Convolution Examples

Poll Question 1:

What effect do you think the following filter will have on an image?

$1/9$	$1/9$	$1/9$
$1/9$	$1/9$	$1/9$
$1/9$	$1/9$	$1/9$

- A. Sharpen the image
- B. Blur the image
- C. Shift the image left
- D. Rotate the image clockwise
- E. Detect edges
- F. Nothing (TOXIC)



Convolution Examples

Gaussian
Blur

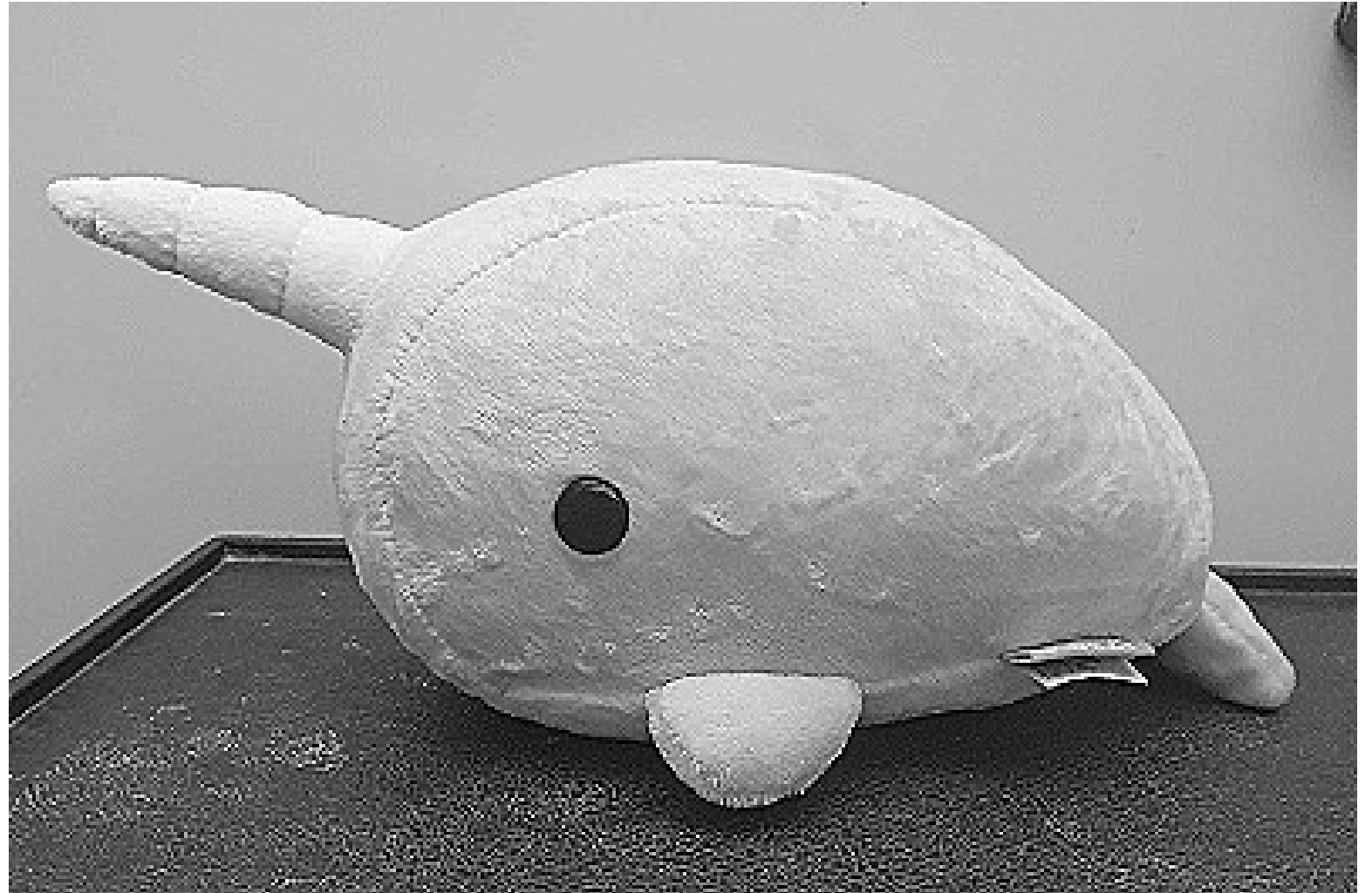
.01	.04	.06	.04	.01
.04	.19	.25	.19	.04
.06	.25	.37	.25	.06
.04	.19	.25	.19	.04
.01	.04	.06	.04	.01



Convolution Examples

Sharpening
Kernel

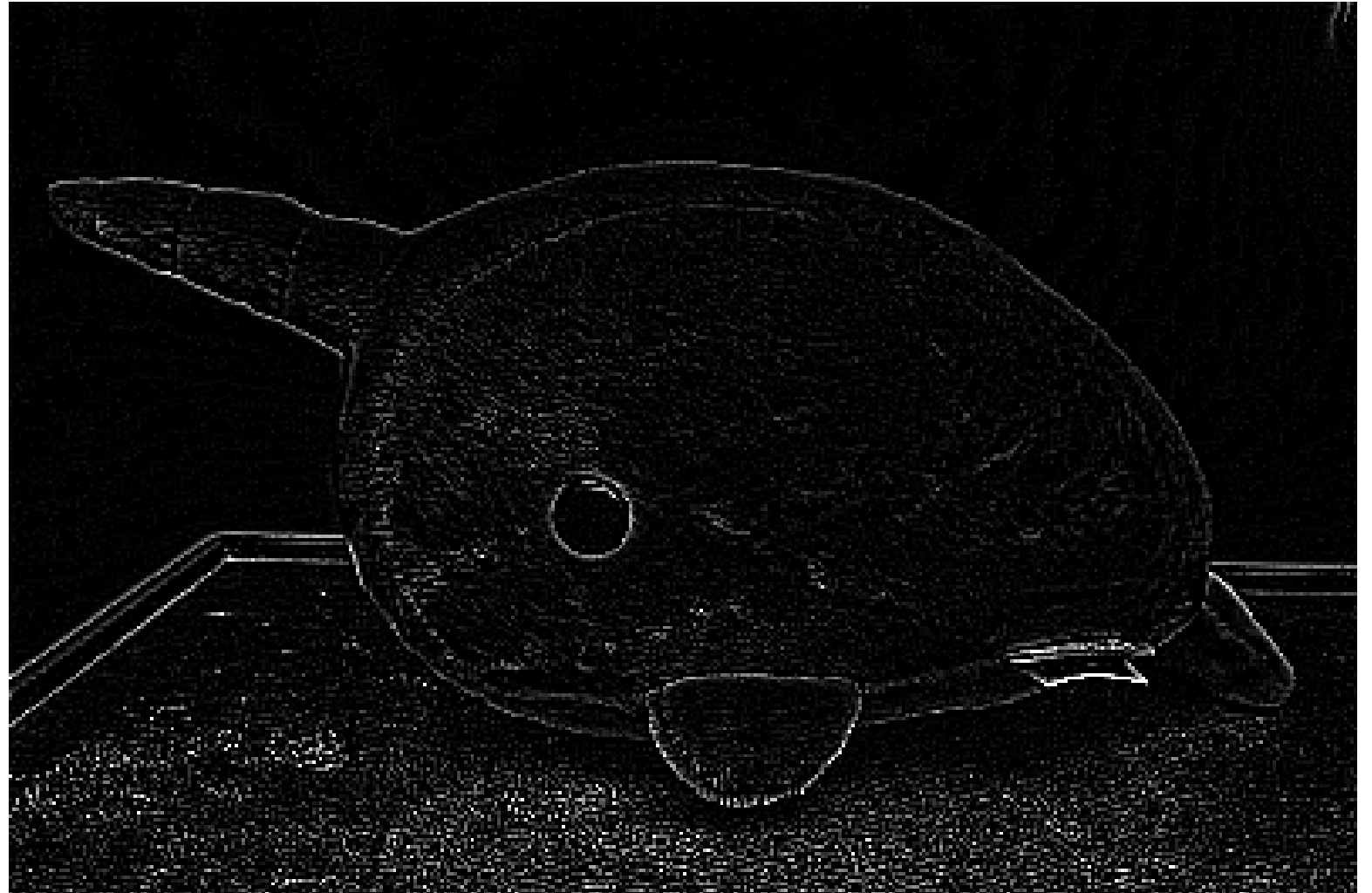
0	-1	0
-1	5	-1
0	-1	0



Convolution Examples

Edge
Detector

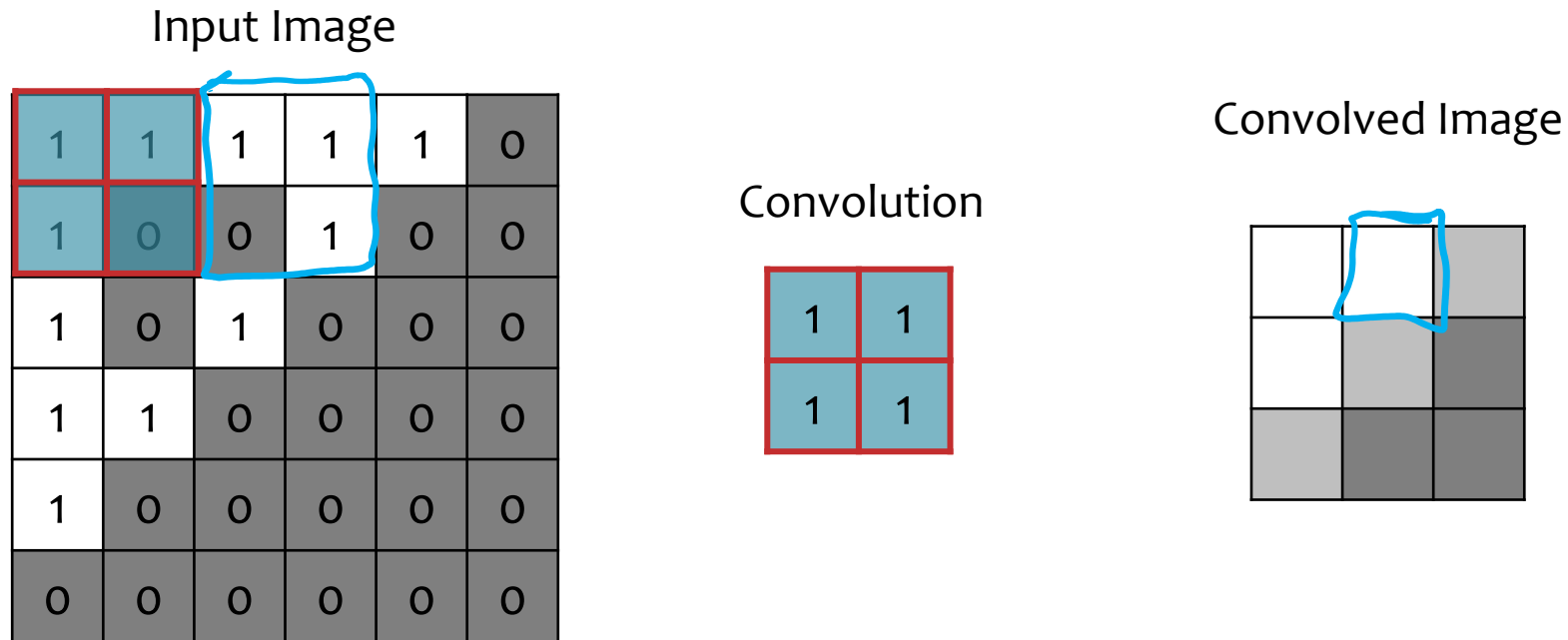
-1	-1	-1
-1	8	-1
-1	-1	-1



STRIDE AND DOWNSAMPLING

Stride and Downsampling

- Suppose we use a convolution with stride 2
- Only 9 patches visited in input, so only 9 pixels in output



Stride and Downsampling

- Suppose we use a convolution with stride 2
- Only 9 patches visited in input, so only 9 pixels in output

Input Image

1	1	1	1	1	0
1	0	0	1	0	0
1	0	1	0	0	0
1	1	0	0	0	0
1	0	0	0	0	0
0	0	0	0	0	0

Convolution

1	1
1	1

Convolved Image

3		

Stride and Downsampling

- Suppose we use a convolution with stride 2
- Only 9 patches visited in input, so only 9 pixels in output

Input Image

1	1	1	1	1	0
1	0	0	1	0	0
1	0	1	0	0	0
1	1	0	0	0	0
1	0	0	0	0	0
0	0	0	0	0	0

Convolution

1	1
1	1

Convolved Image

3	3	

Stride and Downsampling

- Suppose we use a convolution with stride 2
- Only 9 patches visited in input, so only 9 pixels in output

Input Image

1	1	1	1	1	0
1	0	0	1	0	0
1	0	1	0	0	0
1	1	0	0	0	0
1	0	0	0	0	0
0	0	0	0	0	0

Convolution

1	1
1	1

Convolved Image

3	3	1

Stride and Downsampling

- Suppose we use a convolution with stride 2
- Only 9 patches visited in input, so only 9 pixels in output

Input Image

1	1	1	1	1	0
1	0	0	1	0	0
1	0	1	0	0	0
1	1	0	0	0	0
1	0	0	0	0	0
0	0	0	0	0	0

Convolution

1	1
1	1

Convolved Image

3	3	1
3		

Stride and Downsampling

- Suppose we use a convolution with stride 2
- Only 9 patches visited in input, so only 9 pixels in output

Input Image

1	1	1	1	1	0
1	0	0	1	0	0
1	0	1	0	0	0
1	1	0	0	0	0
1	0	0	0	0	0
0	0	0	0	0	0

Convolution

1	1
1	1

Convolved Image

3	3	1
3	1	

Stride and Downsampling

- Suppose we use a convolution with stride 2
- Only 9 patches visited in input, so only 9 pixels in output

Input Image

1	1	1	1	1	0
1	0	0	1	0	0
1	0	1	0	0	0
1	1	0	0	0	0
1	0	0	0	0	0
0	0	0	0	0	0

Convolution

1	1
1	1

Convolved Image

3	3	1
3	1	0

Stride and Downsampling

- Suppose we use a convolution with stride 2
- Only 9 patches visited in input, so only 9 pixels in output

Input Image

1	1	1	1	1	0
1	0	0	1	0	0
1	0	1	0	0	0
1	1	0	0	0	0
1	0	0	0	0	0
0	0	0	0	0	0

Convolution

1	1
1	1

Convolved Image

3	3	1
3	1	0
1		

Stride and Downsampling

- Suppose we use a convolution with stride 2
- Only 9 patches visited in input, so only 9 pixels in output

Input Image

1	1	1	1	1	0
1	0	0	1	0	0
1	0	1	0	0	0
1	1	0	0	0	0
1	0	0	0	0	0
0	0	0	0	0	0

Convolution

1	1
1	1

Convolved Image

3	3	1
3	1	0
1	0	

Stride and Downsampling

- Suppose we use a convolution with stride 2
- Only 9 patches visited in input, so only 9 pixels in output

Input Image

1	1	1	1	1	0
1	0	0	1	0	0
1	0	1	0	0	0
1	1	0	0	0	0
1	0	0	0	0	0
0	0	0	0	0	0

Convolution

1	1
1	1

Convolved Image

3	3	1
3	1	0
1	0	0

Downsampling by Averaging

- Downsampling by averaging is a special case of convolution where the weights are fixed to a uniform distribution
- The example below uses a stride of 2

Input Image

1	1	1	1	1	0
1	0	0	1	0	0
1	0	1	0	0	0
1	1	0	0	0	0
1	0	0	0	0	0
0	0	0	0	0	0

Convolution

$\frac{1}{4}$	$\frac{1}{4}$
$\frac{1}{4}$	$\frac{1}{4}$

Convolved Image

$\frac{3}{4}$	$\frac{3}{4}$	$\frac{1}{4}$
$\frac{3}{4}$	$\frac{1}{4}$	0
$\frac{1}{4}$	0	0

Max-Pooling

- Max-pooling with a stride > 1 is another form of downsampling
- Instead of averaging, we take the max value within the same range as the equivalently-sized convolution
- The example below uses a stride of 2

Input Image

1	1	1	1	1	0
1	0	0	1	0	0
1	0	1	0	0	0
1	1	0	0	0	0
1	0	0	0	0	0
0	0	0	0	0	0

Max-
pooling

$x_{i,j}$	$x_{i,j+1}$
$x_{i+1,j}$	$x_{i+1,j+1}$

$$y_{ij} = \max(x_{ij}, \\ x_{i,j+1}, \\ x_{i+1,j}, \\ x_{i+1,j+1})$$

Max-Pooled
Image

1	1	1
1	1	0
1	0	0

TRAINING CNNS

A Recipe for Machine Learning

1. Given training data:

$$\mathcal{D} = \{\mathbf{x}^{(i)}, y^{(i)}\}_{i=1}^N$$

2. Choose each of these:

- Decision function

$$\hat{y} = h_{\theta}(\mathbf{x})$$

- Loss function

$$\ell(\hat{y}, y) \in \mathbb{R}$$

3. Define goal:

$$\hat{\theta} \approx \operatorname{argmin}_{\theta} \sum_{i=1}^N \ell(h_{\theta}(\mathbf{x}^{(i)}), y^{(i)})$$

4. Train with SGD:

(take small steps
opposite the gradient)

$$\theta^{(t+1)} = \theta^{(t)} - \eta_t \nabla \ell(h_{\theta}(\mathbf{x}^{(i)}), y^{(i)})$$

A Recipe for Machine Learning

1. Given training data:

$$\mathcal{D} = \{x^{(i)}, y^{(i)}\}_{i=1}^N$$

2. Choose each of these:

– Decision function

$$\hat{y} = h_{\theta}(x)$$

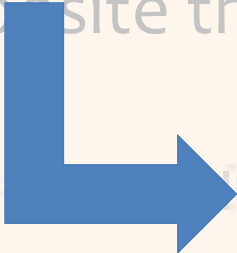
– Loss function

$$\ell(\hat{y}, y) \in \mathbb{R}$$

3. Define goal:

- Q: Now that we have the CNN as a decision function, how do we compute the gradient?
- A: Backpropagation of course!

opposite the gradient)


$$\theta^{(i)} \leftarrow \theta^{(i)} - \eta \nabla \ell(h_{\theta}(x^{(i)}), y^{(i)})$$

SGD for CNNs

Example: Simple CNN Architecture

Given $\mathbf{x}$, $\mathbf{y}^*$ and parameters $\boldsymbol{\theta} = [\boldsymbol{\alpha}, \boldsymbol{\beta}, \mathbf{W}]$

$$J = \ell(\mathbf{y}, \mathbf{y}^*)$$

$$\mathbf{y} = \text{softmax}(\mathbf{z}^{(5)})$$

$$\mathbf{z}^{(5)} = \text{linear}(\mathbf{z}^{(4)}, \mathbf{W})$$

$$\mathbf{z}^{(4)} = \text{relu}(\mathbf{z}^{(3)})$$

$$\mathbf{z}^{(3)} = \text{conv}(\mathbf{z}^{(2)}, \boldsymbol{\beta})$$

$$\mathbf{z}^{(2)} = \text{max-pool}(\mathbf{z}^{(1)})$$

$$\mathbf{z}^{(1)} = \text{conv}(\mathbf{x}, \boldsymbol{\alpha})$$

Algorithm 1 Stochastic Gradient Descent (SGD)

- 1: Initialize $\boldsymbol{\theta}$
 - 2: **while** not converged **do**
 - 3: Sample $i \in \{1, \dots, N\}$
 - 4: Forward: $\mathbf{y} = h_{\boldsymbol{\theta}}(\mathbf{x}^{(i)})$,
 - 5: $J(\boldsymbol{\theta}) = \ell(\mathbf{y}, \mathbf{y}^{(i)})$
 - 6: Backward: Compute $\nabla_{\boldsymbol{\theta}} J(\boldsymbol{\theta})$
 - 7: Update: $\boldsymbol{\theta} \leftarrow \boldsymbol{\theta} - \eta \nabla_{\boldsymbol{\theta}} J(\boldsymbol{\theta})$
-

LAYERS OF A CNN

Convolutional Neural Network (CNN)

- Typical layers include:
 - Convolutional layer
 - Max-pooling layer
 - Fully-connected (Linear) layer
 - ReLU layer (or some other nonlinear activation function)
 - Softmax
- These can be arranged into arbitrarily deep topologies

Architecture #1: LeNet-5

PROC. OF THE IEEE, NOVEMBER 1998

7

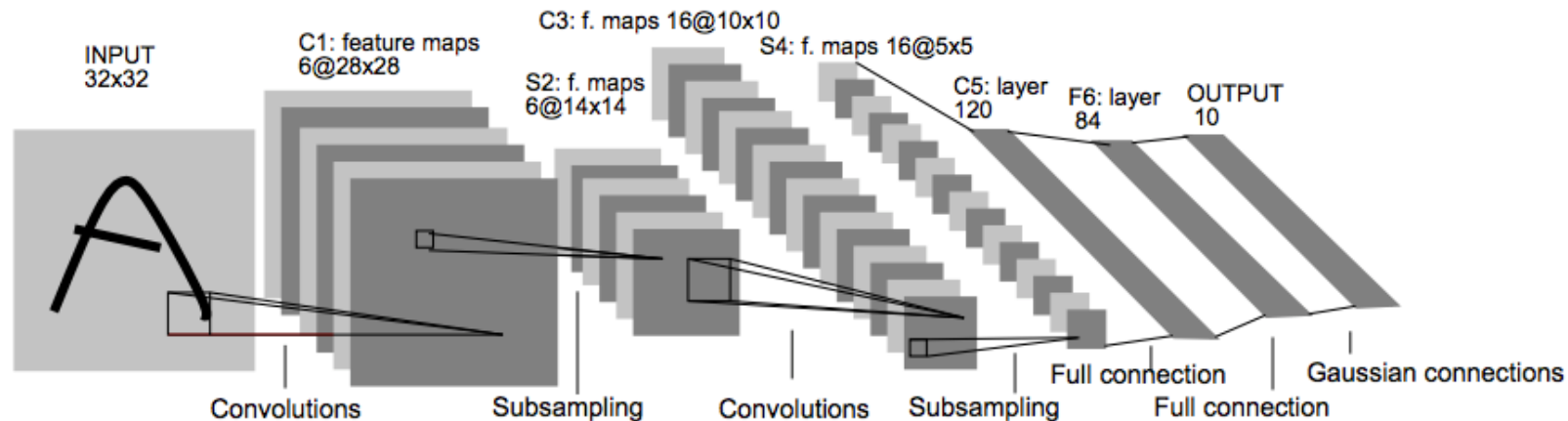


Fig. 2. Architecture of LeNet-5, a Convolutional Neural Network, here for digits recognition. Each plane is a feature map, i.e. a set of units whose weights are constrained to be identical.

ReLU Layer

Output: $y \in \mathbb{R}^K$

Forward:

$$y = \sigma(x), \text{ element-wise} \\ \sigma(a) = \max(0, a)$$

Input: $x \in \mathbb{R}^K$

Input: $\frac{\partial J}{\partial y} \in \mathbb{R}^K$

Backward: for each j ,

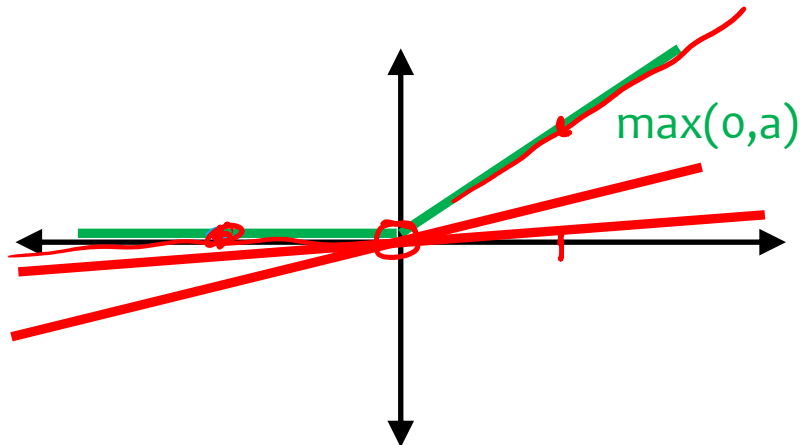
$$\frac{\partial J}{\partial x_j} = \frac{\partial J}{\partial y_j} \frac{\partial y_j}{\partial x_j}$$

where

$$\frac{\partial y_j}{\partial x_j} = \begin{cases} 1 & \text{if } x_j > 0 \\ 0 & \text{otherwise} \end{cases}$$

subderivative

Output: $\frac{\partial J}{\partial x} \in \mathbb{R}^K$



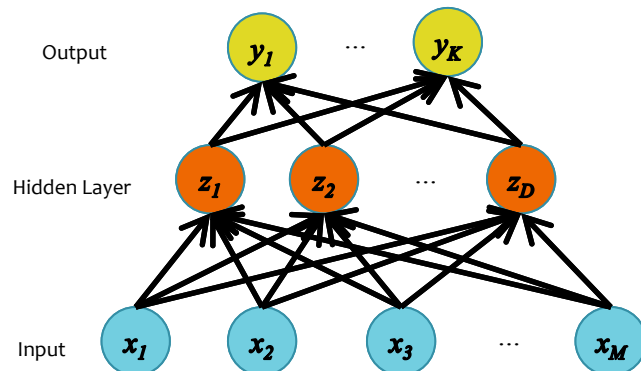
Softmax Layer

Output: $y \in \mathbb{R}^K$

Forward: for each i ,

$$y_i = \frac{\exp(x_i)}{\sum_{k=1}^K \exp(x_k)}$$

Input: $x \in \mathbb{R}^K$



Input: $\frac{\partial J}{\partial y} \in \mathbb{R}^K$

Backward: for each j ,

$$\frac{\partial J}{\partial x_j} = \sum_{i=1}^K \frac{\partial J}{\partial y_i} \frac{\partial y_i}{\partial x_j}$$

where

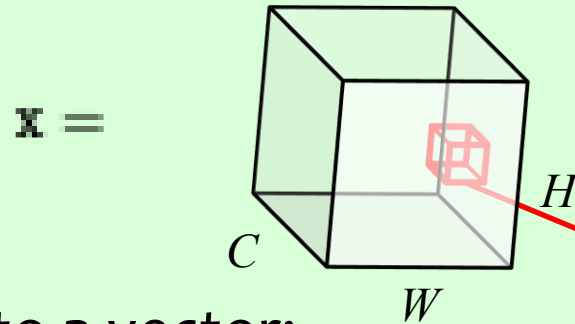
$$\frac{\partial y_i}{\partial x_j} = \begin{cases} y_i(1 - y_i) & \text{if } i = j \\ -y_i y_j & \text{otherwise} \end{cases}$$

Output: $\frac{\partial J}{\partial x} \in \mathbb{R}^K$

Fully-Connected Layer (3D input)

Forward:

1. suppose input is a 3D tensor:



2. flatten out tensor into a vector:

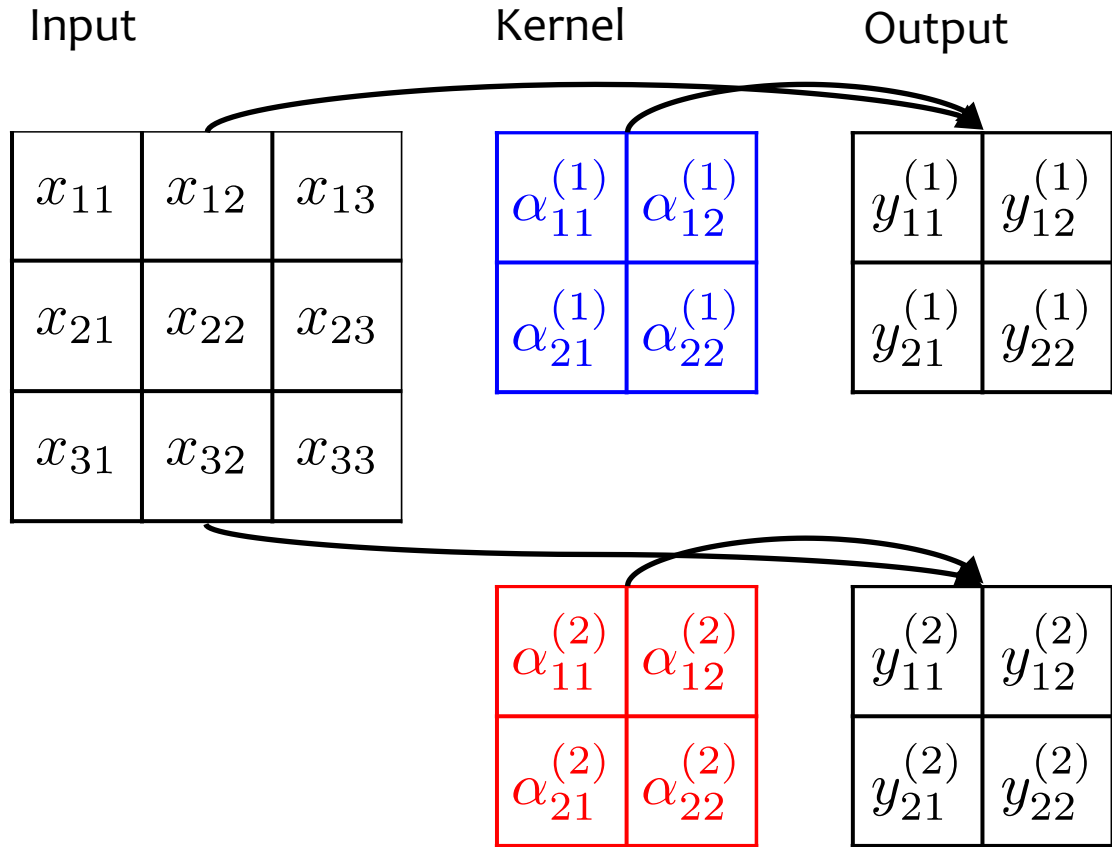
$$\hat{\mathbf{x}} = [\hat{x}_1, \dots, \hat{x}_{(C \times H \times W)}] \quad \text{where } \hat{x}_{(H \times W \times i + W \times j + k)} = x_{i,j,k}$$

3. then push that vector through a standard linear layer:

$$\mathbf{y} = \boldsymbol{\alpha}^T \hat{\mathbf{x}} + \alpha_0 \quad \text{where } \boldsymbol{\alpha} \in \mathbb{R}^{A \times B}, \quad \alpha_0 \in \mathbb{R}^B$$
$$|\hat{\mathbf{x}}| \in \mathbb{R}^A, \quad |\mathbf{y}| \in \mathbb{R}^B$$

2D Convolution

Example: 1 input channel, 2 output channels



$$y_{11}^{(1)} = \alpha_{11}^{(1)} x_{11} + \alpha_{12}^{(1)} x_{12} + \alpha_{21}^{(1)} x_{21} + \alpha_{22}^{(1)} x_{22} + \alpha_0^{(1)}$$

$$y_{12}^{(1)} = \alpha_{11}^{(1)} x_{12} + \alpha_{12}^{(1)} x_{13} + \alpha_{21}^{(1)} x_{22} + \alpha_{22}^{(1)} x_{23} + \alpha_0^{(1)}$$

$$y_{21}^{(1)} = \alpha_{11}^{(1)} x_{21} + \alpha_{12}^{(1)} x_{22} + \alpha_{21}^{(1)} x_{31} + \alpha_{22}^{(1)} x_{32} + \alpha_0^{(1)}$$

$$y_{22}^{(1)} = \alpha_{11}^{(1)} x_{22} + \alpha_{12}^{(1)} x_{23} + \alpha_{21}^{(1)} x_{32} + \alpha_{22}^{(1)} x_{33} + \alpha_0^{(1)}$$

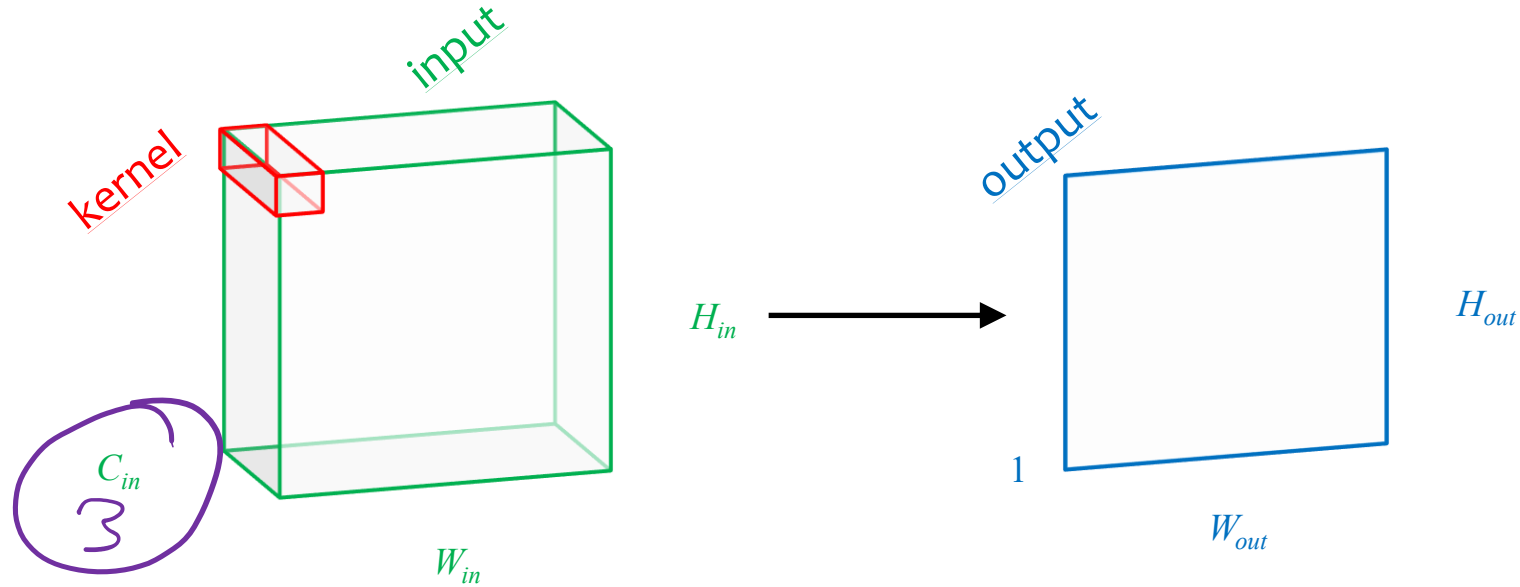
$$y_{11}^{(2)} = \alpha_{11}^{(2)} x_{11} + \alpha_{12}^{(2)} x_{12} + \alpha_{21}^{(2)} x_{21} + \alpha_{22}^{(2)} x_{22} + \alpha_0^{(2)}$$

$$y_{12}^{(2)} = \alpha_{11}^{(2)} x_{12} + \alpha_{12}^{(2)} x_{13} + \alpha_{21}^{(2)} x_{22} + \alpha_{22}^{(2)} x_{23} + \alpha_0^{(2)}$$

$$y_{21}^{(2)} = \alpha_{11}^{(2)} x_{21} + \alpha_{12}^{(2)} x_{22} + \alpha_{21}^{(2)} x_{31} + \alpha_{22}^{(2)} x_{32} + \alpha_0^{(2)}$$

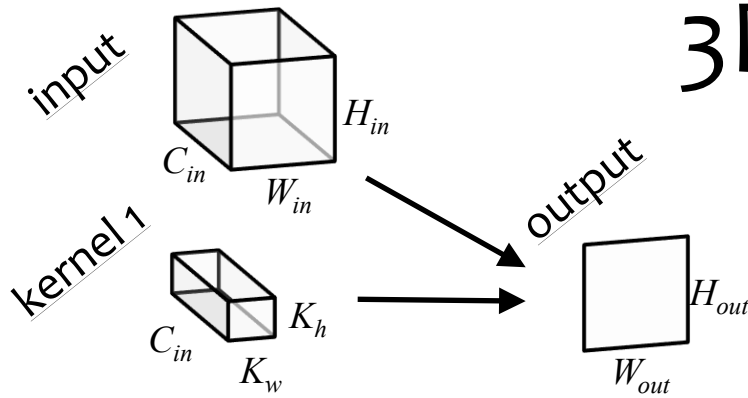
$$y_{22}^{(2)} = \alpha_{11}^{(2)} x_{22} + \alpha_{12}^{(2)} x_{23} + \alpha_{21}^{(2)} x_{32} + \alpha_{22}^{(2)} x_{33} + \alpha_0^{(2)}$$

Convolution of a Color Image

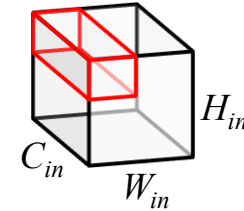


- Color images consist of 3 floats per pixel for RGB (red, green blue) color values
- The kernel must also be 3-dimensional
- input = 3x64x64
- kernel = 3x5x5
- output = 1x64x64 (assuming padding)

3D Convolutional Layer

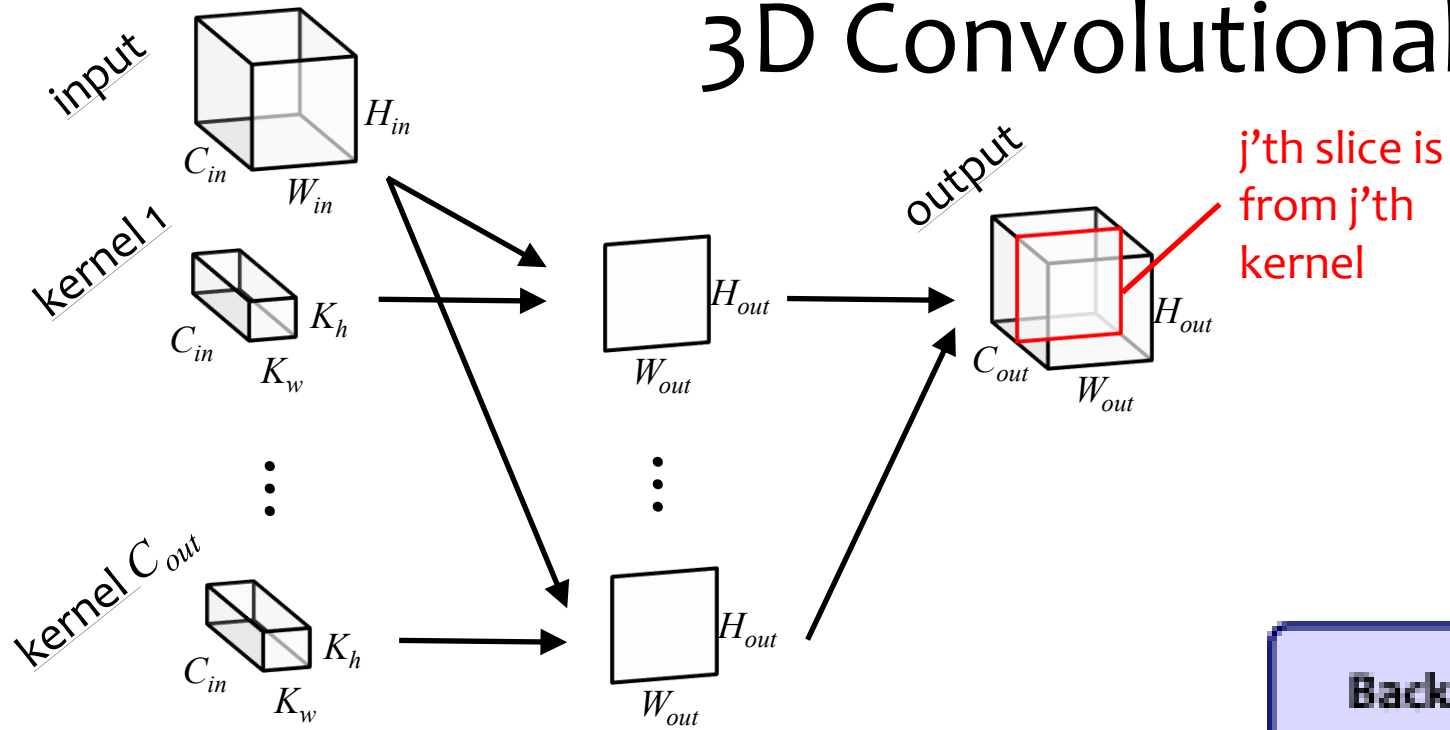


Convolution in 3D

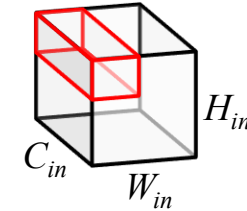


- Color images consist of 3 floats per pixel for RGB (red, green blue) color values
- Convolution must also be 3-dimensional

3D Convolutional Layer



Convolution in 3D



Forward:

$$y_{h',w'}^{(c')} = \beta^{(c')} + \sum_{c=1}^{C_{in}} \sum_{m=1}^{K_h} \sum_{n=1}^{K_w} x_{h'+ms, w'+ns}^{(c)} \cdot \alpha_{m,n}^{(c',c)}$$

$s \in \mathbb{Z}$ (stride)

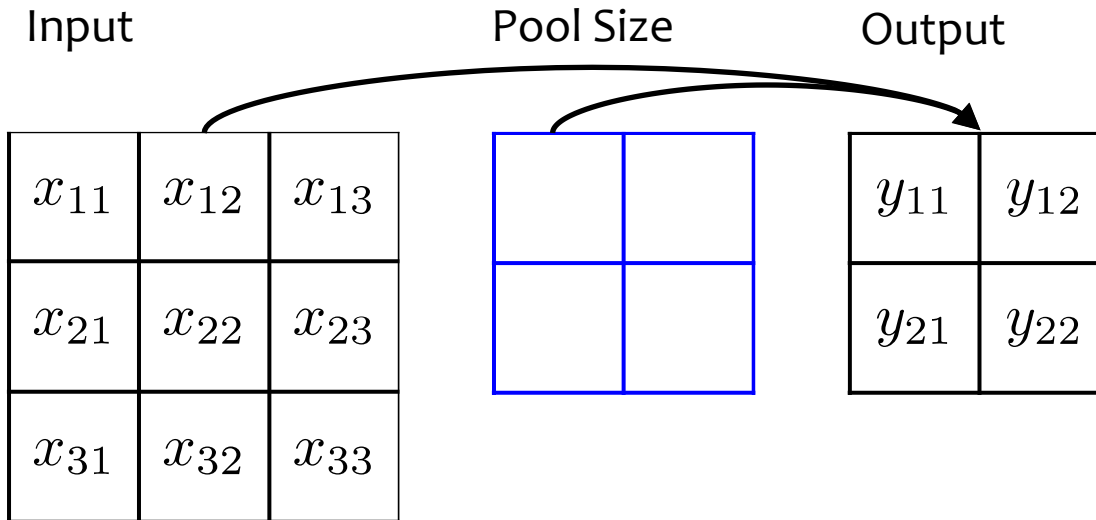
Backward:

$$\frac{\partial J}{\partial \alpha_{m,n}^{(c',c)}} = \sum_{h'=1}^{H_{out}} \sum_{w'=1}^{W_{out}} \frac{\partial J}{\partial y_{h',w'}^{(c')}} \cdot x_{h'+ms, w'+ns}^{(c)}$$

$$\frac{\partial J}{\partial \beta^{(c')}} = \sum_{h'=1}^{H_{out}} \sum_{w'=1}^{W_{out}} \frac{\partial J}{\partial y_{h',w'}^{(c')}} \cdot y_{h',w'}^{(c')}$$

Max-Pooling Layer

Example: 1 input channel, 1 output channel, stride of 1



$$y_{11} = \max(x_{11}, x_{12}, x_{21}, x_{22})$$

$$y_{12} = \max(x_{12}, x_{13}, x_{22}, x_{23})$$

$$y_{21} = \max(x_{21}, x_{22}, x_{31}, x_{32})$$

$$y_{22} = \max(x_{22}, x_{23}, x_{32}, x_{33})$$

3D Max-Pooling Layer

Output: $y \in \mathbb{R}^{C \times H_{out} \times W_{out}}$

Forward:

$$y_{ij}^{(c)} = \max_{q \in \{1, \dots, K_h\}, r \in \{1, \dots, K_w\}} x_{mn}^{(c)}$$

where

$$m = s(i - 1) + q$$

$$n = s(j - 1) + r$$

Input:

$$x \in \mathbb{R}^{C \times H_{in} \times W_{in}}$$

K_h, K_w (kernel size)

s (stride)

Input: $\frac{\partial J}{\partial y} \in \mathbb{R}^{C \times H_{out} \times W_{out}}$

Backward:

$$\frac{\partial J}{\partial x_{mn}^{(c)}} = \sum_i \sum_j \frac{\partial J}{\partial y_{ij}^{(c)}} \frac{\partial y_{ij}^{(c)}}{\partial x_{mn}^{(c)}}$$

Output: $\frac{\partial J}{\partial x} \in \mathbb{R}^{C \times H_{in} \times W_{in}}$

- $\max()$ is not differentiable, but subdifferentiable.
- There are a set of derivatives and we can just choose one for SGD

$$y = \max(a, b)$$

$$\Rightarrow \frac{dy}{da} = \frac{dy}{dy} \frac{dy}{da} \text{ where } \frac{dy}{da} = \begin{cases} 1 & \text{if } a > b \\ 0 & \text{otherwise} \end{cases}$$

CNN ARCHITECTURES

Convolutional Neural Network (CNN)

- Typical layers include:
 - Convolutional layer
 - Max-pooling layer
 - Fully-connected (Linear) layer
 - ReLU layer (or some other nonlinear activation function)
 - Softmax
- These can be arranged into arbitrarily deep topologies

Architecture #1: LeNet-5

PROC. OF THE IEEE, NOVEMBER 1998

7

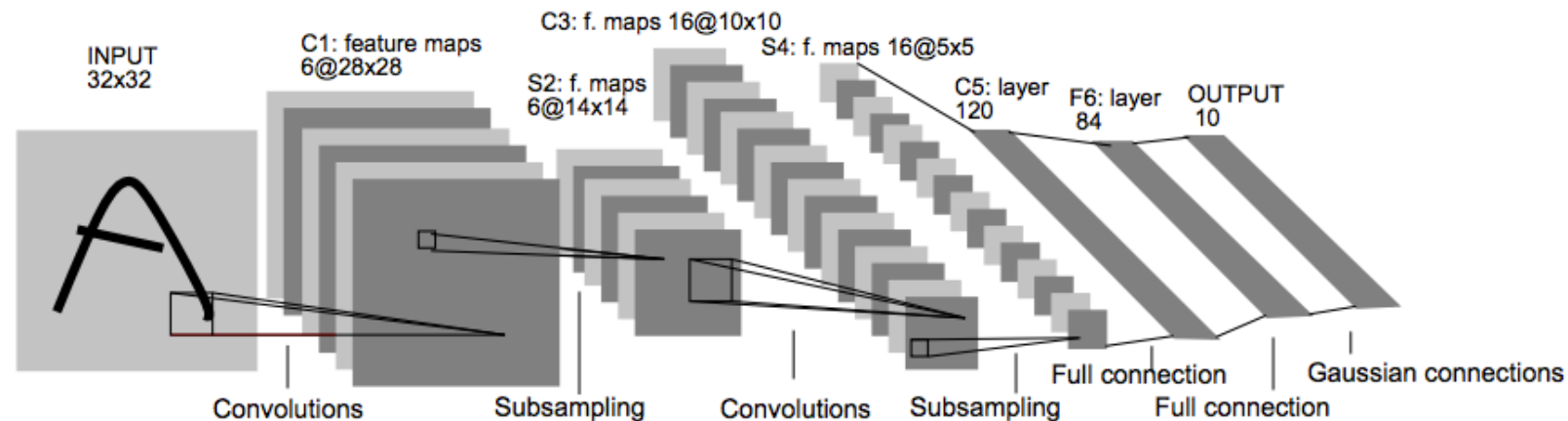


Fig. 2. Architecture of LeNet-5, a Convolutional Neural Network, here for digits recognition. Each plane is a feature map, i.e. a set of units whose weights are constrained to be identical.

Architecture #2: AlexNet

CNN for Image Classification

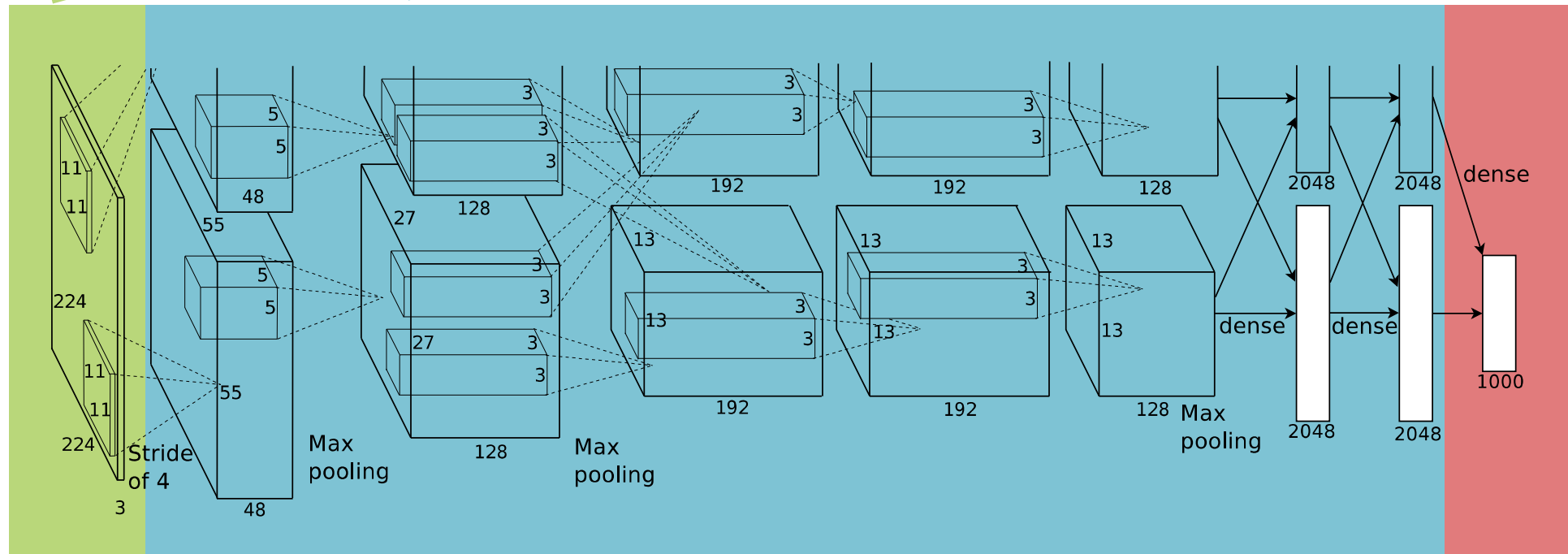
(Krizhevsky, Sutskever & Hinton, 2012)

15.3% error on ImageNet LSVRC-2012 contest

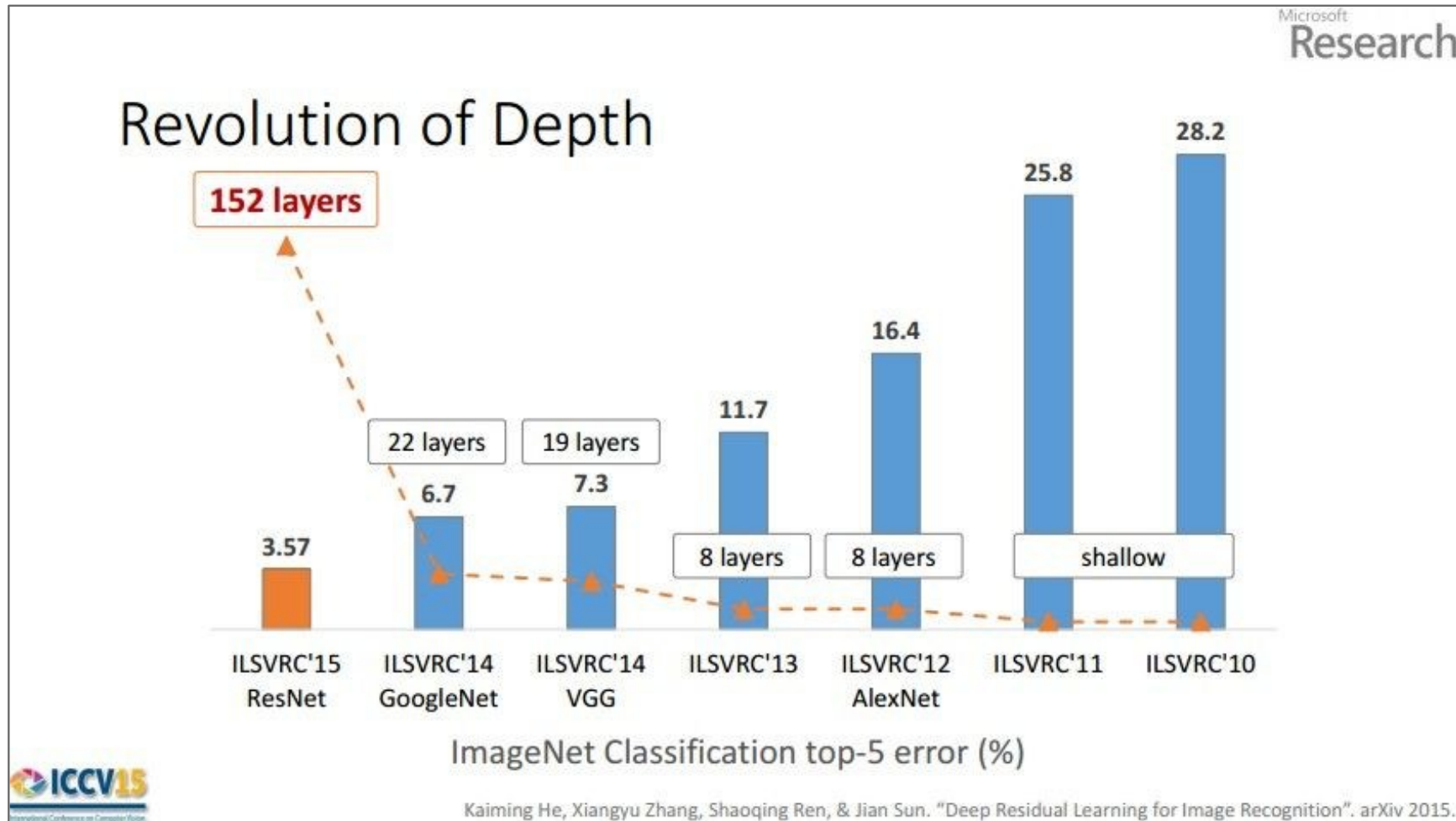
Input
image
(pixels)

- Five convolutional layers (w/max-pooling)
- Three fully connected layers

1000-way
softmax

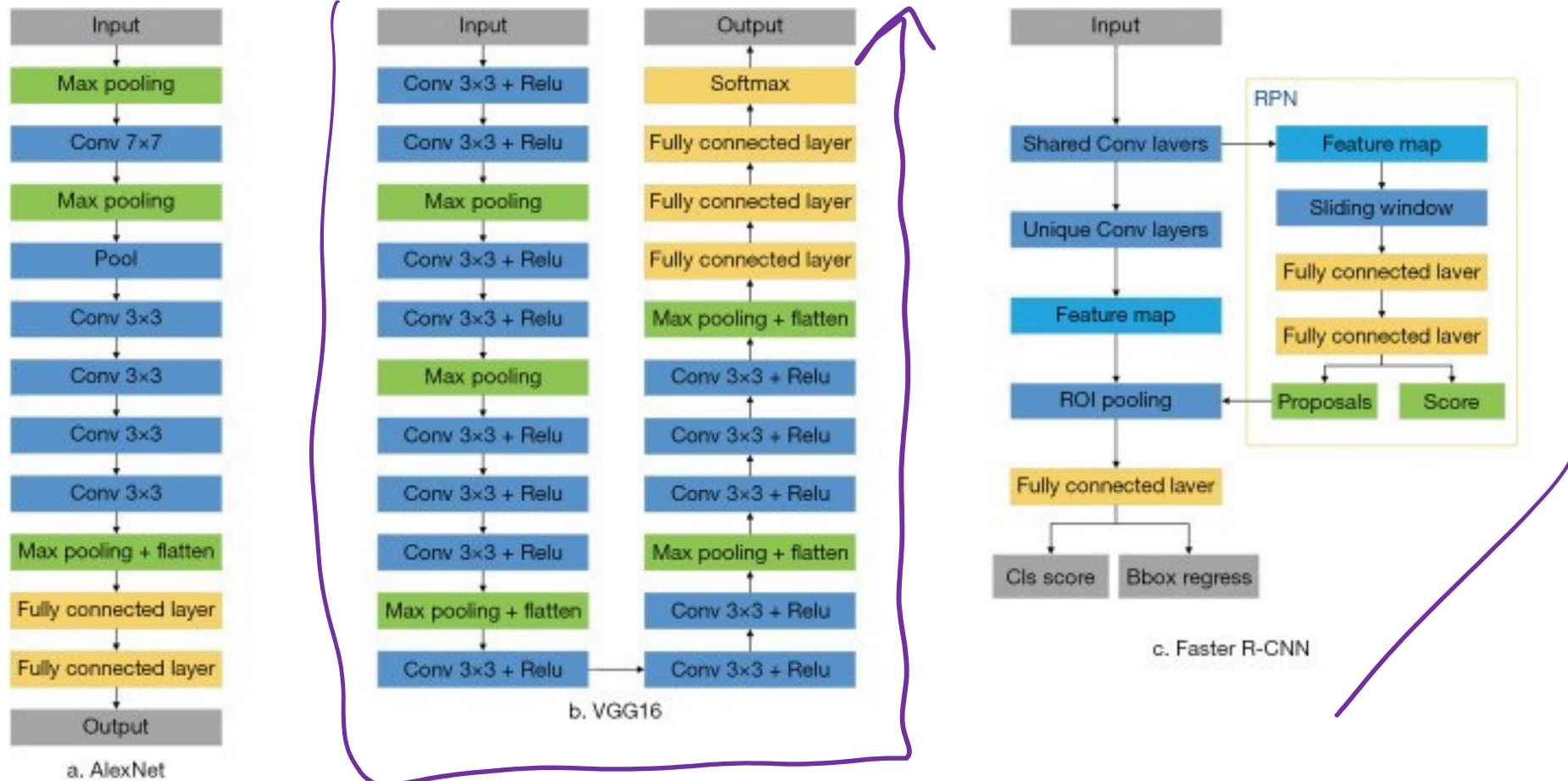


CNNs for Image Recognition



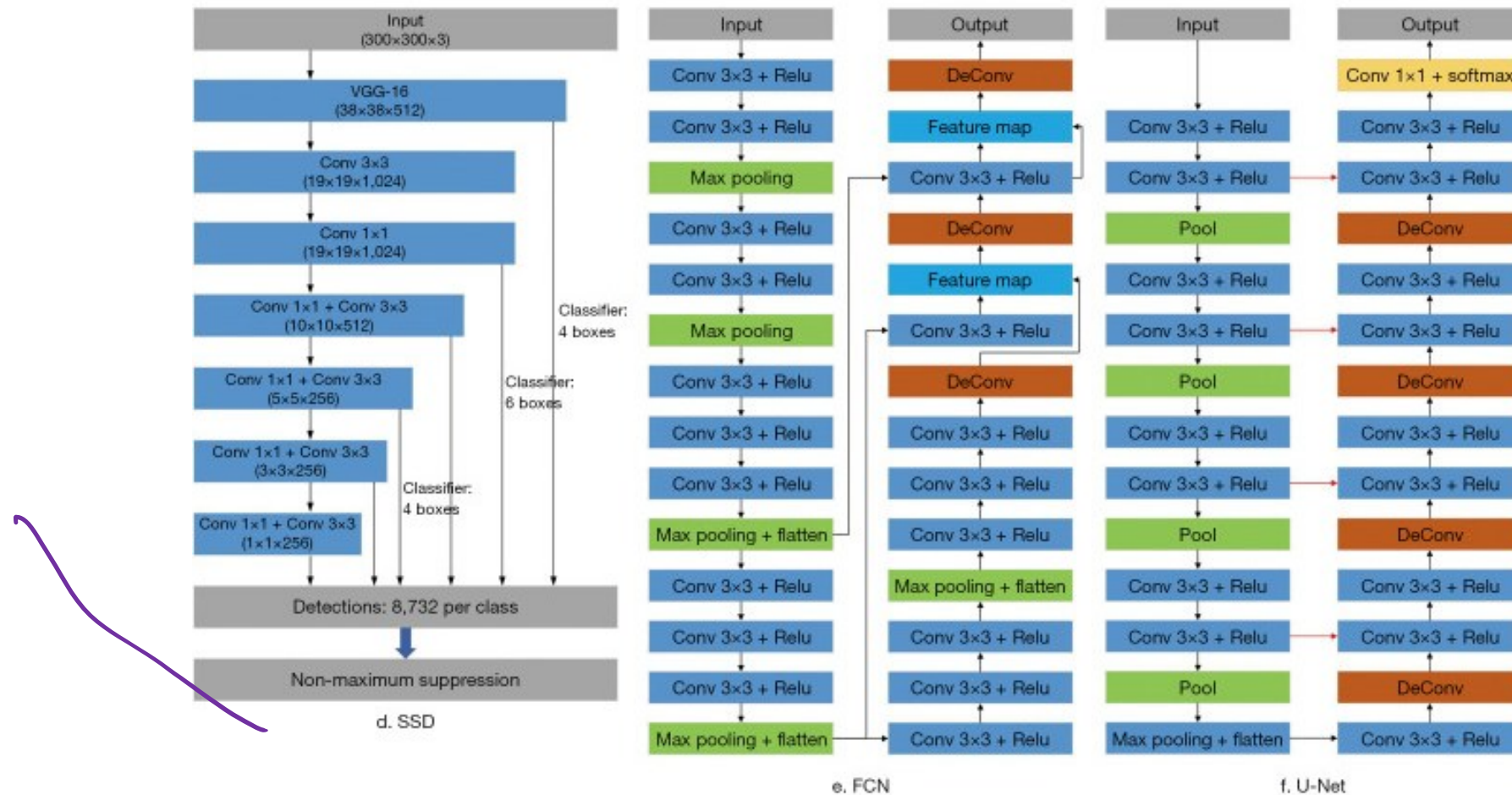
Convolutional Neural Network (CNN)

Typical Architectures



Convolutional Neural Network (CNN)

Typical Architectures



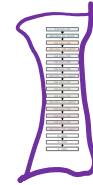
Convolutional Neural Network (CNN)

Typical Architectures

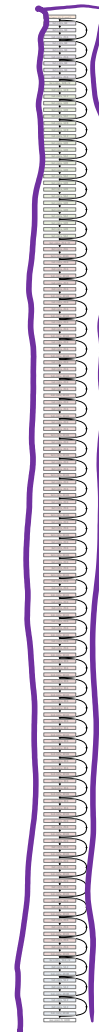
AlexNet, 8 layers
(ILSVRC 2012)



VGG, 19 layers
(ILSVRC 2014)



ResNet, **152 layers**
(ILSVRC 2015)



Microsoft
Research

Location-specific Parameters

Poll Question 2:

Why do many layers used in computer vision *not have* location specific parameters?

Answer:

Convolutional Layer

For a convolutional layer, how do we pick the kernel size (aka. the size of the convolution)?

Input Image

0	0	0	0	0	0	0
0	1	1	1	1	1	0
0	1	0	0	1	0	0
0	1	0	1	0	0	0
0	1	1	0	0	0	0
0	1	0	0	0	0	0
0	0	0	0	0	0	0

2x2
Convolution

θ_{11}	θ_{12}
θ_{21}	θ_{22}

3x3
Convolution

θ_{11}	θ_{12}	θ_{13}
θ_{21}	θ_{22}	θ_{23}
θ_{31}	θ_{32}	θ_{33}

4x4
Convolution

θ_{11}	θ_{12}	θ_{13}	θ_{14}
θ_{21}	θ_{22}	θ_{23}	θ_{24}
θ_{31}	θ_{32}	θ_{33}	θ_{34}
θ_{41}	θ_{42}	θ_{43}	θ_{44}

- A small kernel can only see a very small part of the image, but is fast to compute
- A large kernel can see more of the image, but at the expense of speed

CNN VISUALIZATIONS

Visualization of CNN

https://adamharley.com/nn_vis/cnn/2d.html



MNIST Digit Recognition with CNNs (in your browser)

<https://cs.stanford.edu/people/karpathy/convnetjs/demo/mnist.html>

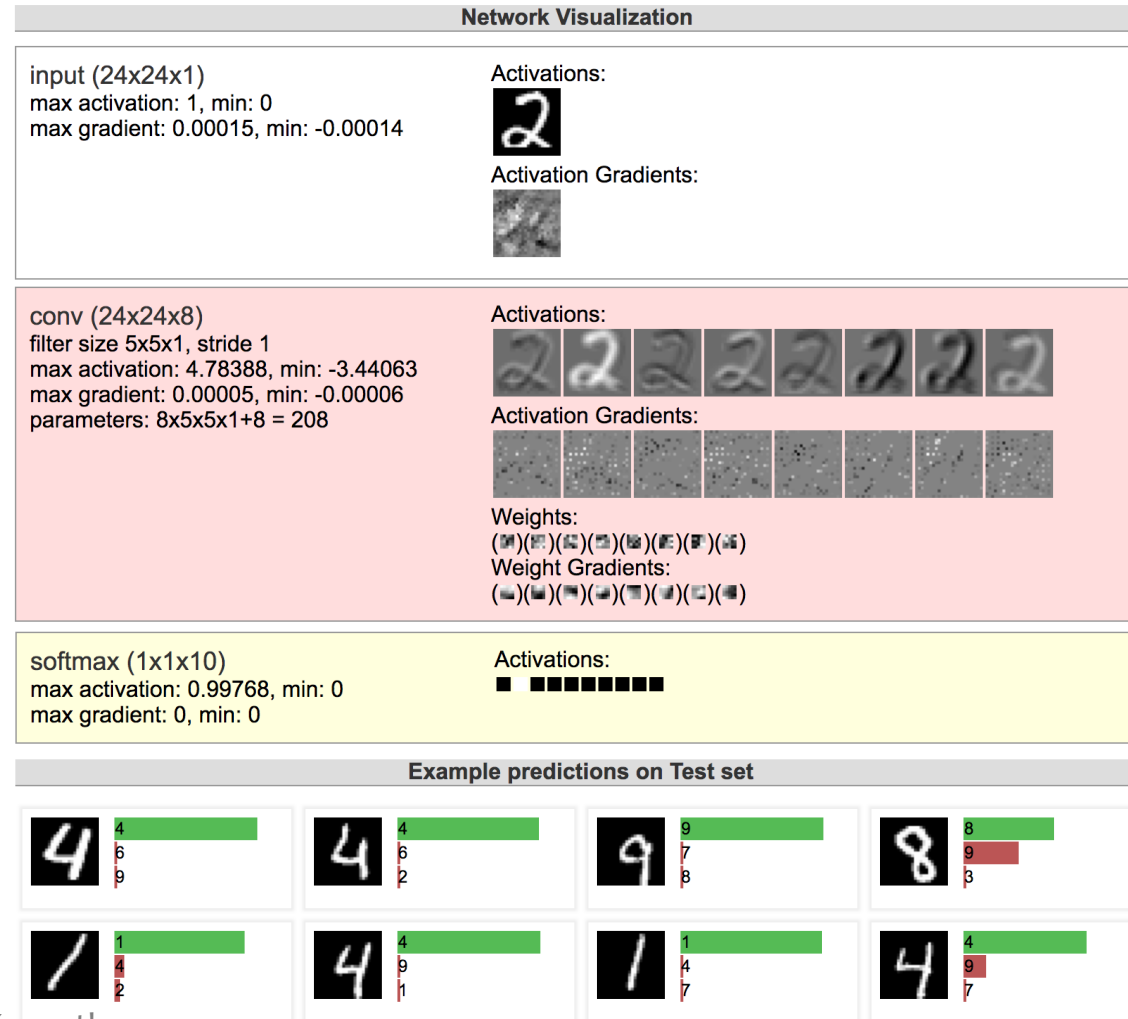


Figure from Andrej Karpathy

CNN Summary

CNNs

- Are used for all aspects of **computer vision**, and have won numerous pattern recognition competitions
- Able learn **interpretable features** at different levels of abstraction
- Typically, consist of **convolution** layers, **pooling** layers, **nonlinearities**, and **fully connected** layers

WORD EMBEDDINGS

Word Embeddings

Key Idea:

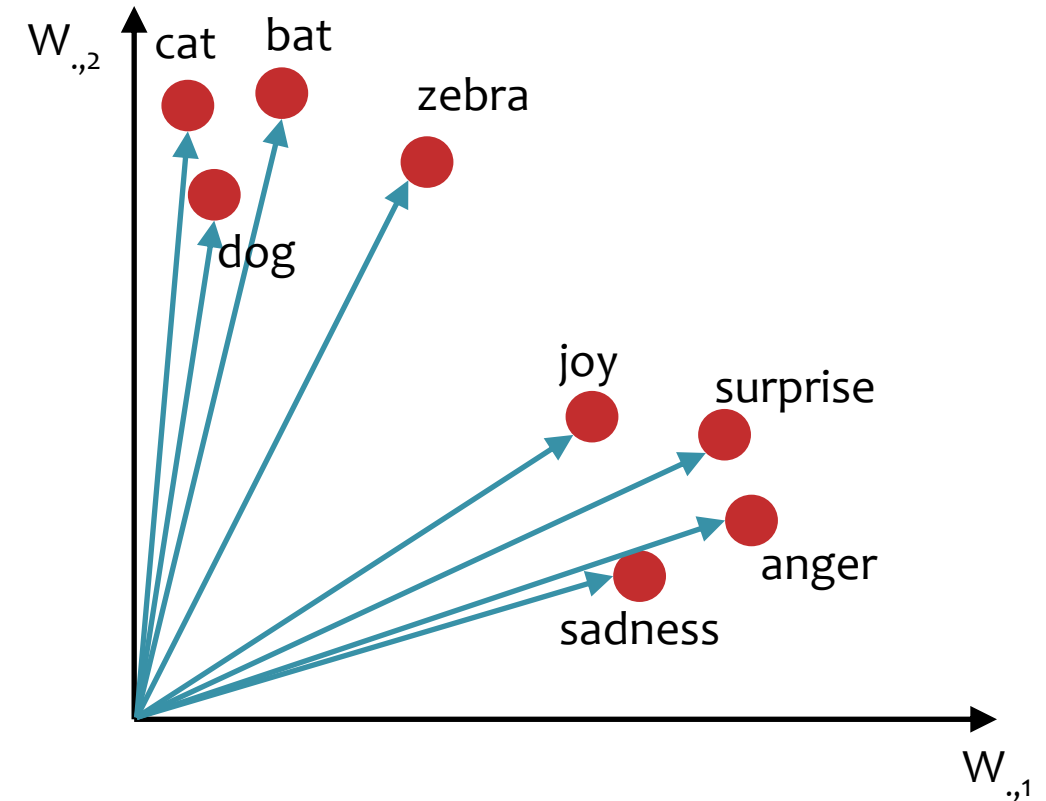
- represent each word in your vocabulary as a vector
- store as a $V \times D$ matrix where:
V = number of words in vocab.
D = vector's dimension

Modeling:

- define a model in which the vectors are parameters
- each copy of the word uses the same parameter vector
- train model so that similar words have high cosine similarity

W

anger	W_{11}	W_{12}
bat	W_{21}	W_{22}
cat	W_{31}	W_{32}
dog	W_{41}	W_{42}
joy	W_{51}	W_{52}
sadness	W_{61}	W_{62}
surprise	W_{71}	W_{72}
zebra	W_{81}	W_{82}



Word Embeddings

Key Idea:

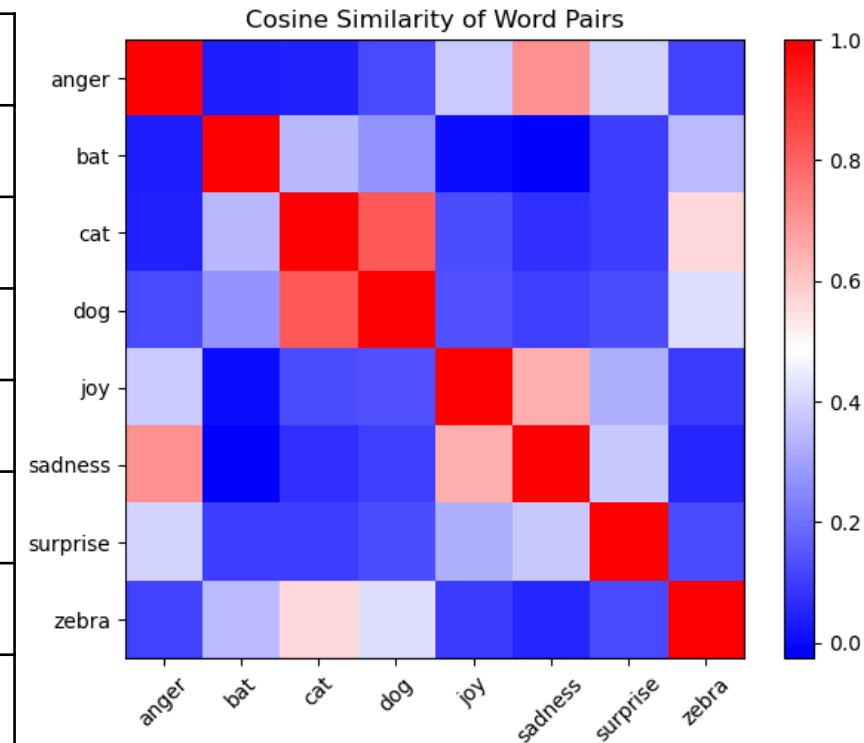
- represent each word in your vocabulary as a vector
- store as a $V \times D$ matrix where:
 V = number of words in vocab.
 D = vector's dimension

Modeling:

- define a model in which the vectors are parameters
- each copy of the word uses the same parameter vector
- train model so that similar words have high cosine similarity

	W				
aardvark	-2.3	0.0	-2.8	...	-4.5
anger	-2.8	-0.9	-1.7	...	-4.3
bat	-4.5	-1.3	0.6	...	-1.7
cat	3.5	-2.0	-2.3	...	-0.4
...				...	
joy	3.0	-0.6	-0.6	...	4.9
...				...	
zebra	-4.7	-4.2	-4.5	...	4.3

in a real use case, the typical embedding dimension is in the hundreds, e.g. $D = 300$

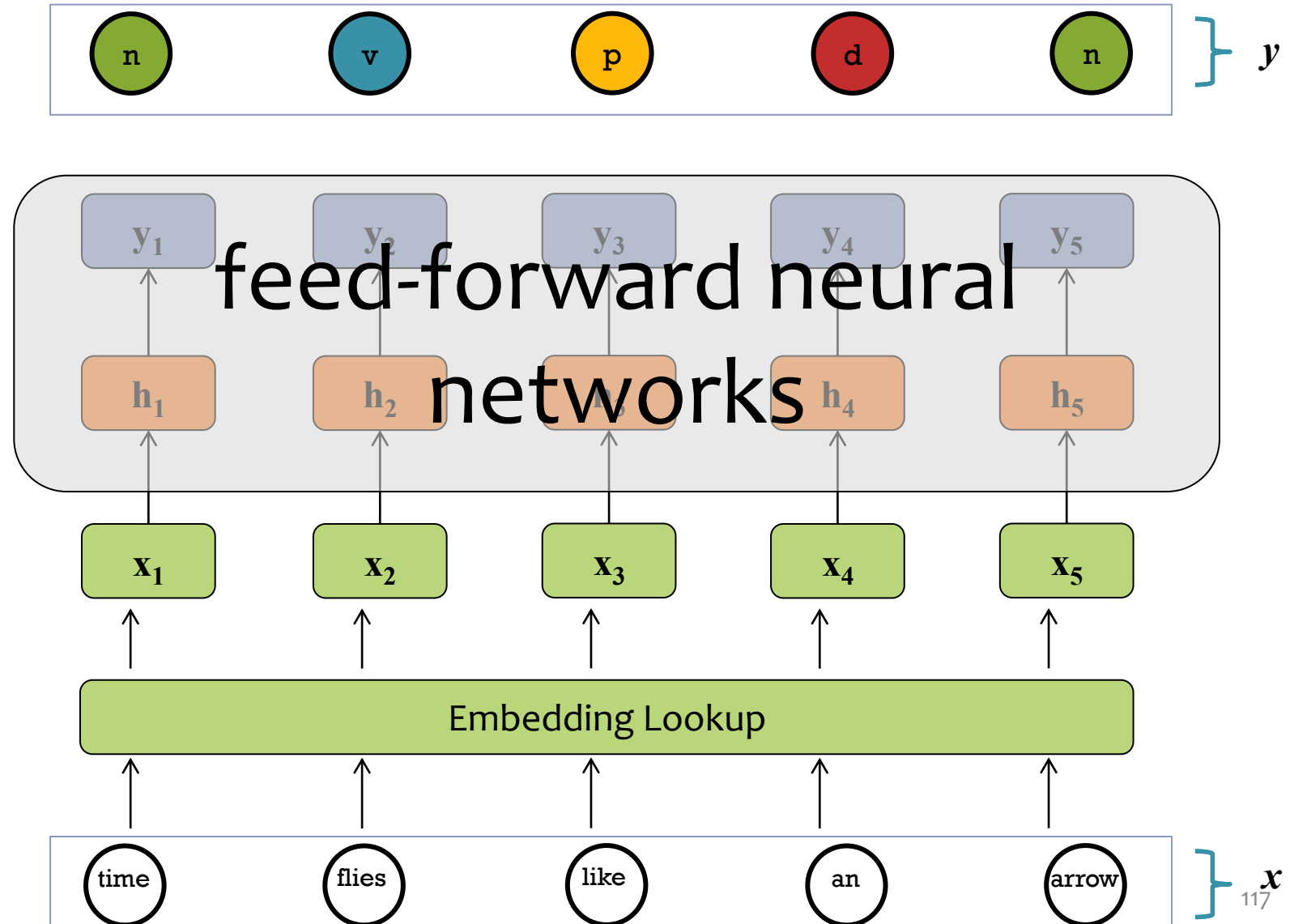


we can't visualize 300 dimensional vectors, but we can inspect their pairwise cosine similarities

Word Embeddings

In all the models we're about to consider (neural networks, RNNs, Transformers) that work with sentences...

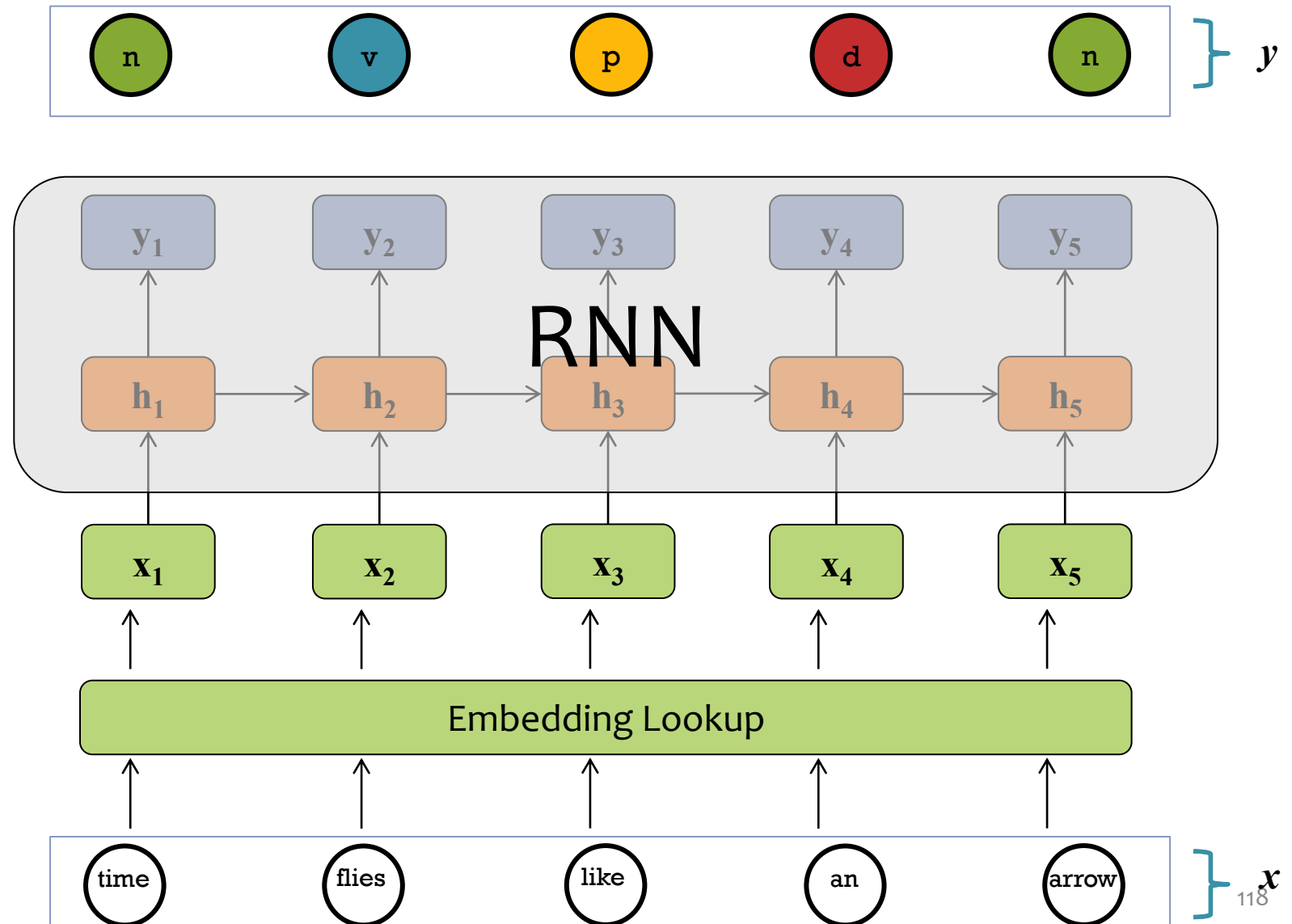
... the first step is always to look up the t 'th word's embedding vector parameters and use said vector for the value of x_t



Word Embeddings

In all the models we're about to consider (neural networks, RNNs, Transformers) that work with sentences...

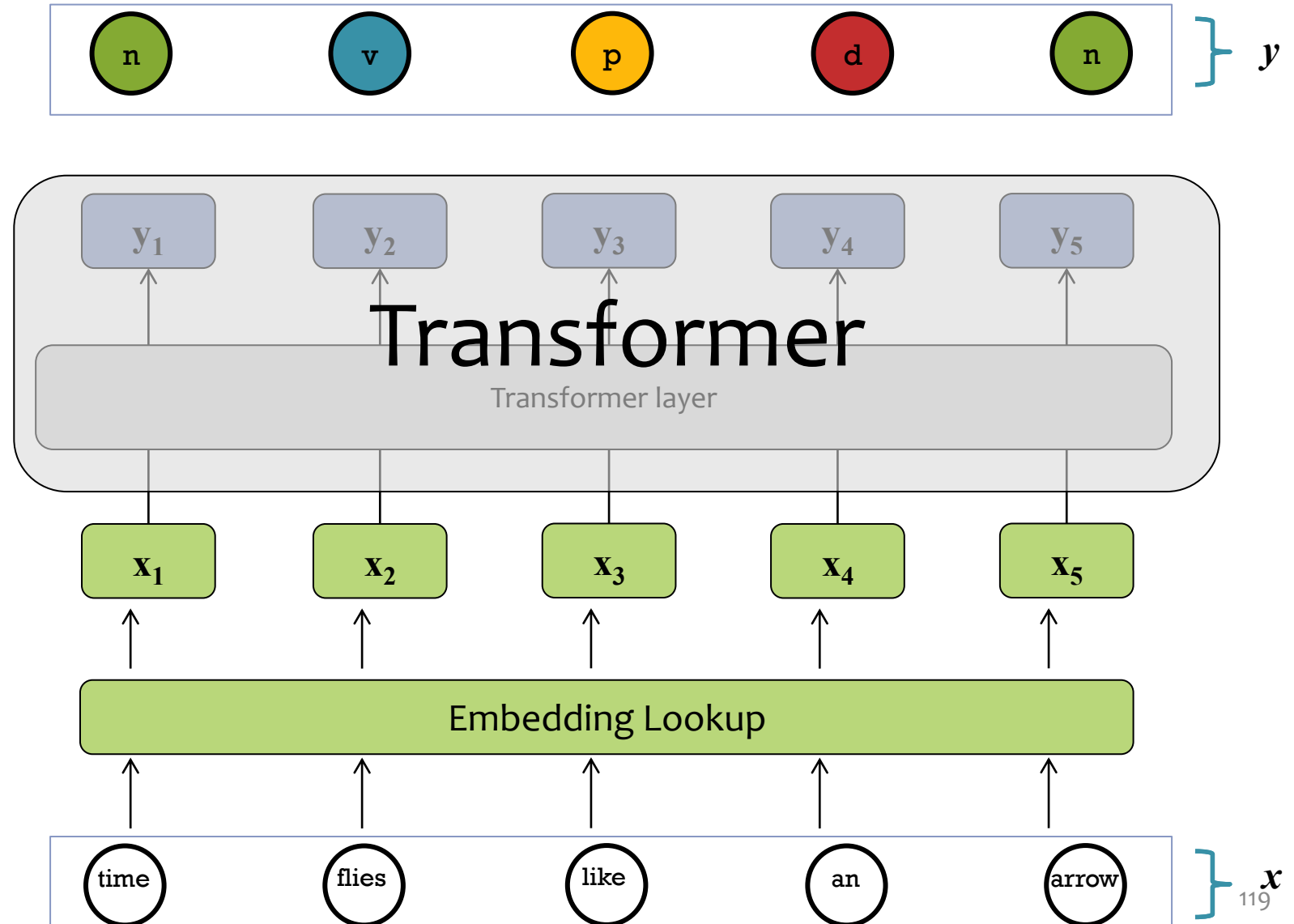
... the first step is always to look up the t 'th word's embedding vector parameters and use said vector for the value of x_t



Word Embeddings

In all the models we're about to consider (neural networks, RNNs, Transformers) that work with sentences...

... the first step is always to look up the t 'th word's embedding vector parameters and use said vector for the value of x_t



SEQUENCE TAGGING

Dataset for Supervised Part-of-Speech (POS) Tagging

Data: $\mathcal{D} = \{\mathbf{x}^{(n)}, \mathbf{y}^{(n)}\}_{n=1}^N$

Sample 1:	n time	v flies	p like	d an	n arrow	} $y^{(1)}$ } $x^{(1)}$
Sample 2:	n time	n flies	v like	d an	n arrow	} $y^{(2)}$ } $x^{(2)}$
Sample 3:	n flies	v fly	p with	n their	n wings	} $y^{(3)}$ } $x^{(3)}$
Sample 4:	p with	n time	n you	v will	v see	} $y^{(4)}$ } $x^{(4)}$

Dataset for Supervised Handwriting Recognition

Data: $\mathcal{D} = \{\mathbf{x}^{(n)}, \mathbf{y}^{(n)}\}_{n=1}^N$



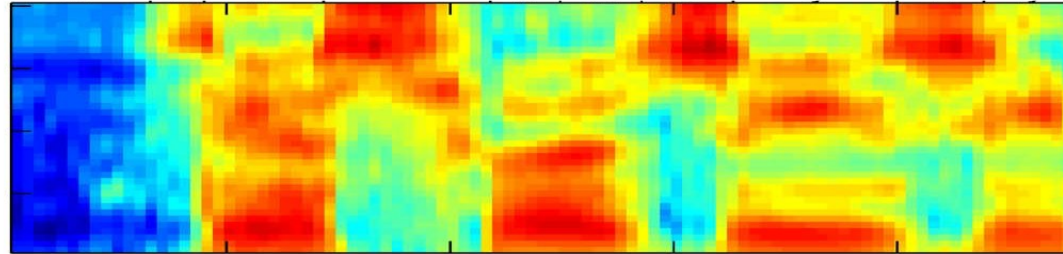
Dataset for Supervised Phoneme (Speech) Recognition

Data: $\mathcal{D} = \{\mathbf{x}^{(n)}, \mathbf{y}^{(n)}\}_{n=1}^N$

Sample 1:



} $\mathbf{y}^{(1)}$

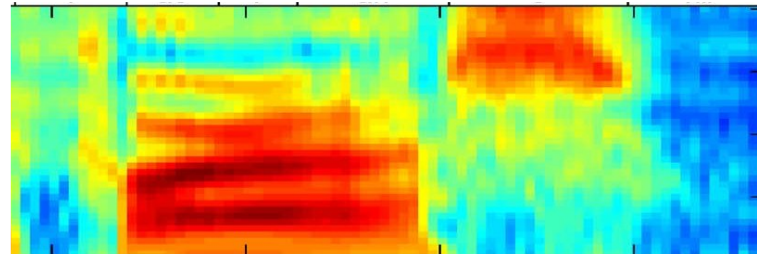


} $\mathbf{x}^{(1)}$

Sample 2:



} $\mathbf{y}^{(2)}$



} $\mathbf{x}^{(2)}$

RECURRENT NEURAL NETWORKS

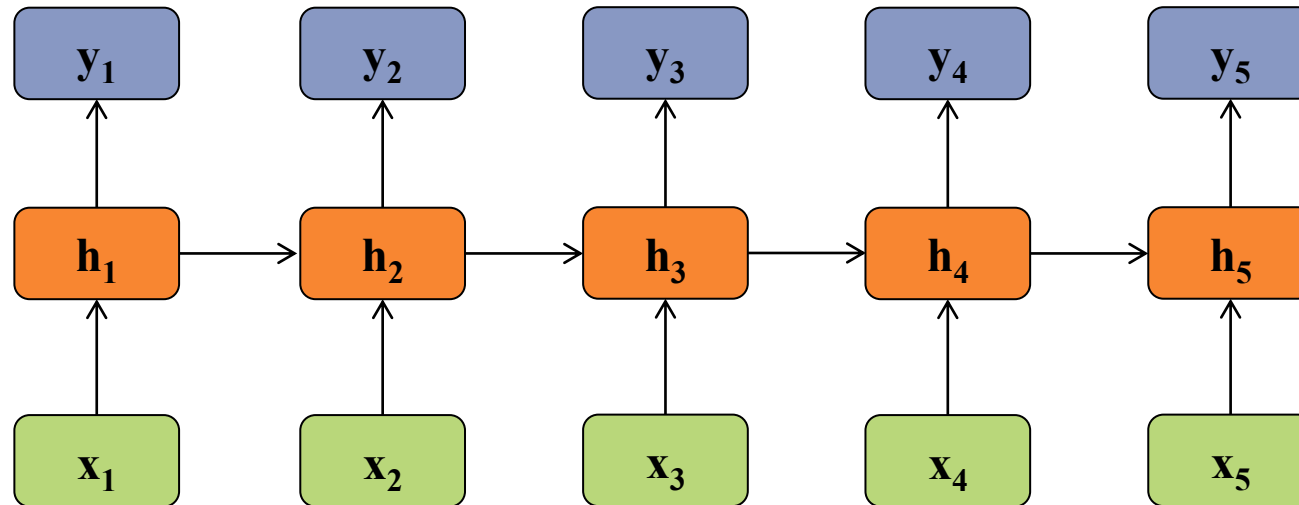
Recurrent Neural Networks (RNNs)

inputs: $\mathbf{x} = (x_1, x_2, \dots, x_T), x_i \in \mathcal{R}^I$
hidden units: $\mathbf{h} = (h_1, h_2, \dots, h_T), h_i \in \mathcal{R}^J$
outputs: $\mathbf{y} = (y_1, y_2, \dots, y_T), y_i \in \mathcal{R}^K$
nonlinearity: $\mathcal{H}$

Definition of the RNN:

$$h_t = \mathcal{H}(W_{xh}x_t + W_{hh}h_{t-1} + b_h)$$

$$y_t = W_{hy}h_t + b_y$$



Recurrent Neural Networks (RNNs)

inputs: $\mathbf{x} = (x_1, x_2, \dots, x_T), x_i \in \mathcal{R}^I$

hidden units: $\mathbf{h} = (h_1, h_2, \dots, h_T), h_i \in \mathcal{R}^J$

outputs: $\mathbf{y} = (y_1, y_2, \dots, y_T), y_i \in \mathcal{R}^K$

nonlinearity: $\mathcal{H}$

Definition of the RNN:

$$h_t = \mathcal{H}(W_{xh}x_t + W_{hh}h_{t-1} + b_h)$$

$$y_t = W_{hy}h_t + b_y$$



This form of RNN is
called an
Elman Network

