

Fundamentals of Deep Learning

Raymond Ptucha,
Rochester Institute of Technology
NVIDIA, Deep Learning Institute



Electronic Imaging 2020: SC20
January 28, 2019, 8:30am-12:45pm



IS&T International Symposium on
Electronic Imaging 2020
SCIENCE AND TECHNOLOGY

26-30 January 2020
Burlingame, California USA

1

Fair Use Agreement

This agreement covers the use of all slides in this document, please read carefully.

- You may freely use these slides, if:
 - You send me an email telling me the conference/venue/company name in advance, and which slides you wish to use.
 - You receive a positive confirmation email back from me.
 - My name (R. Ptucha) appears on each slide you use.

(c) Raymond Ptucha, rwpeec@rit.edu

R. Ptucha '20

2

2

Agenda

- Part I- Intuition and Theory
 - 8:35-9:15pm: Introduction
 - 9:15-10:00pm: Convolutional Neural Networks
 - 10:00-10:40pm: Recurrent Neural Networks
- 10:40-11:00pm: Break
- Part II- Hands on
 - 11:00am-12:45pm: Hands-on exercises

R. Ptucha '20

4

4

Two Most Important Deep Learning Fields

- Convolutional Neural Networks (CNN)
 - Examine high dimensional input, learn features and classifier simultaneously
- Recurrent Neural Networks (RNN)
 - Learn temporal signals, remember both short and long sequences

R. Ptucha '20

5

5

Two Most Important Deep Learning Fields

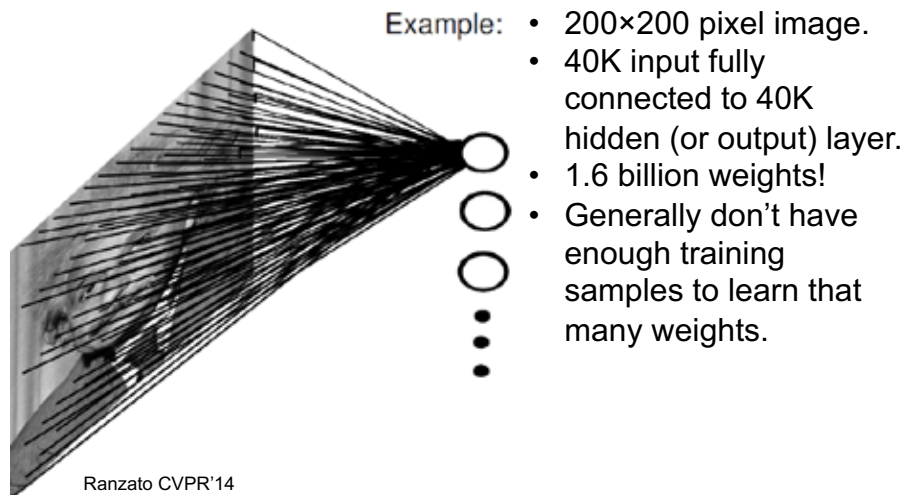
- Convolutional Neural Networks (CNN)
 - Examine high dimensional input, learn features and classifier simultaneously
- Recurrent Neural Networks (RNN)
 - Learn temporal signals, remember both short and long sequences

R. Ptucha '20

6

6

Fully Connected Layers?

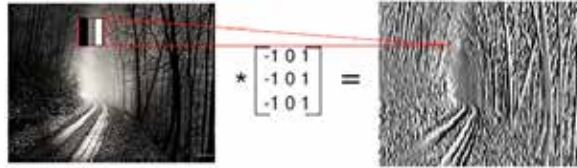


R. Ptucha '20

7

7

Convolution Filter



Ranzato CVPR'14

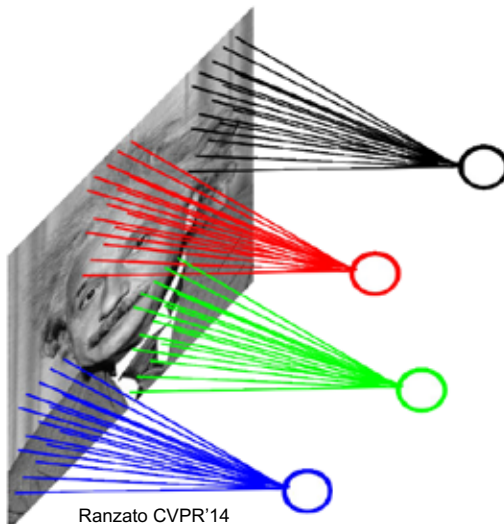
- Convolution filters apply a transform to an image.
- The above filter detects vertical edges.

R. Ptucha '20

8

8

Locally Connected Layer



Ranzato CVPR'14

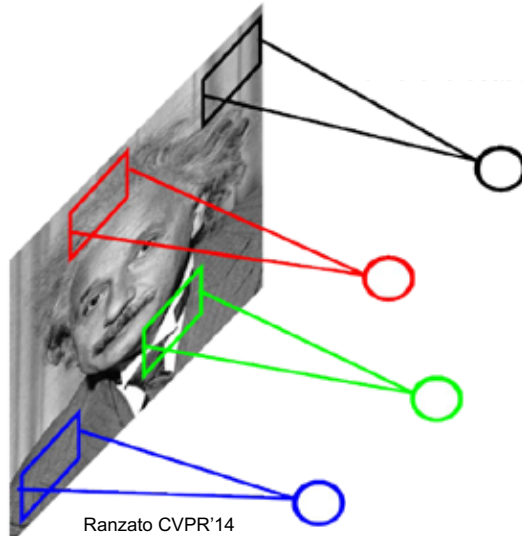
- 200×200 pixel image.
- 40K input.
- Four 10×10 filters, each fully connected
- 40K×10×10×4=16M weights....getting better!

R. Ptucha '20

9

9

Locally Connected Layer



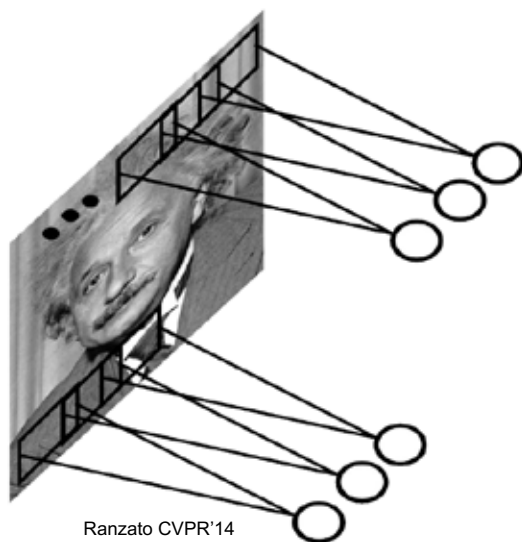
- 200×200 pixel image.
- 40K input.
- Four 10×10 filters, each fully connected
- $40K \times 10 \times 10 \times 4 = 16M$ weights....getting better!
- Can we formulate so each filter has similar statistics across all locations?

R. Ptucha '20

10

10

Convolution Layer



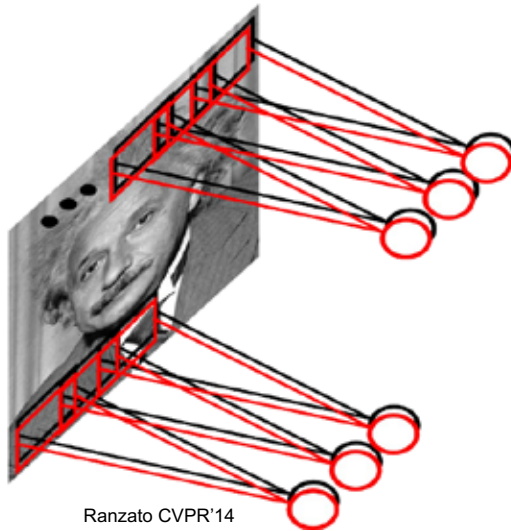
- 200×200 pixel image.
- 40K input.
- Four 10×10 filters, each fully connected
- $40K \times 10 \times 10 \times 4 = 16M$ weights....getting better!
- Require each filter has same statistics across all locations.
- Learn filters.

R. Ptucha '20

11

11

Convolution Layer



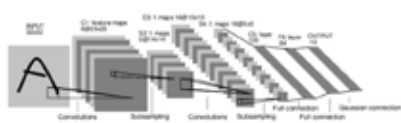
- 200×200 pixel image.
- 40K input.
- Four 10×10 filters, each fully connected
- $40K \times 10 \times 10 \times 4 = 16M$ weights....getting better!
- Require each filter has same statistics across all locations.
- Learn filters.
- To learn four filters we have $4 \times 10 \times 10 = 400$ parameters- great!

R. Ptucha '20

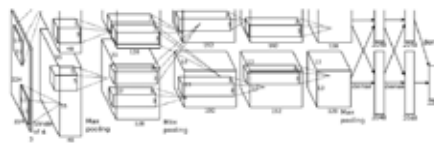
12

12

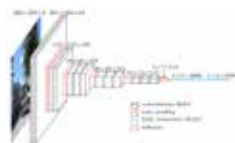
Many Flavors of CNNs...



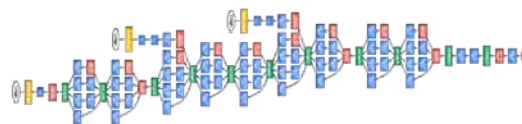
LeNet-5, LeCun 1989



AlexNet, Krizhevsky 2012



VGGNet, Simonyan 2014



GoogLeNet (Inception), Szegedy 2014



ResNet, He 2015



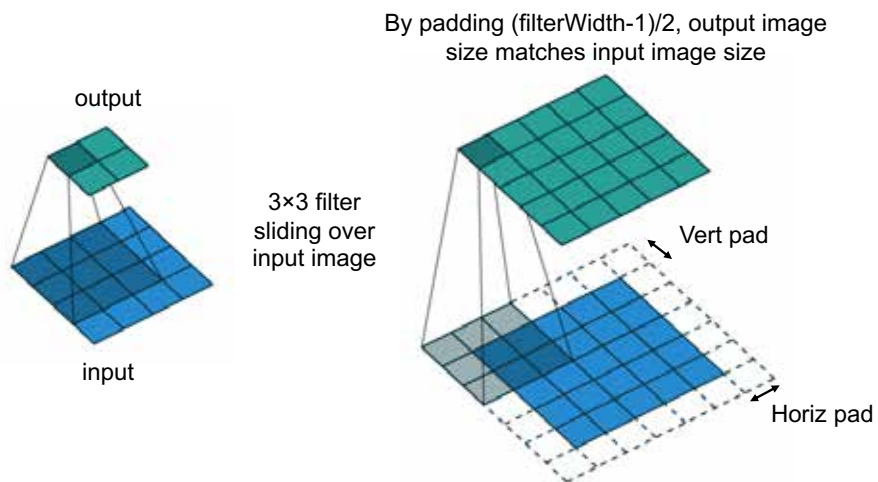
DenseNet, Huang 2017

R. Ptucha '20

13

13

Image Convolution



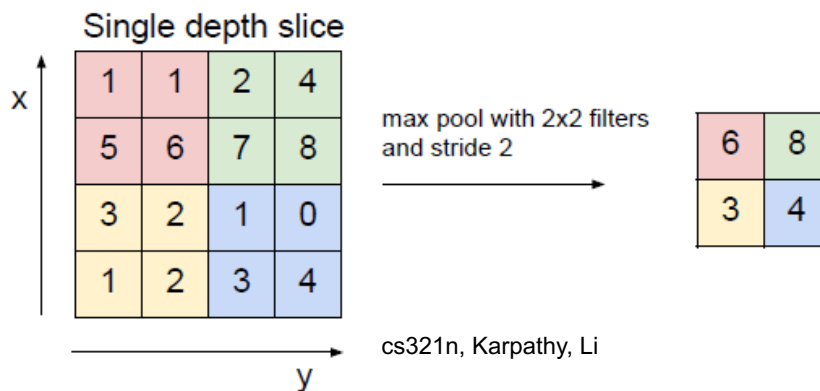
https://github.com/vdumoulin/conv_arithmetic

R. Ptucha '20

14

14

Max Pooling- Reducing the Size of an Image

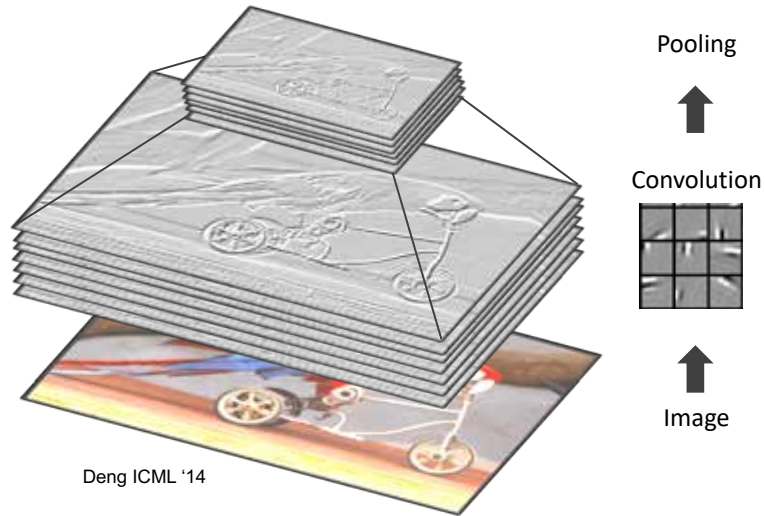


R. Ptucha '20

15

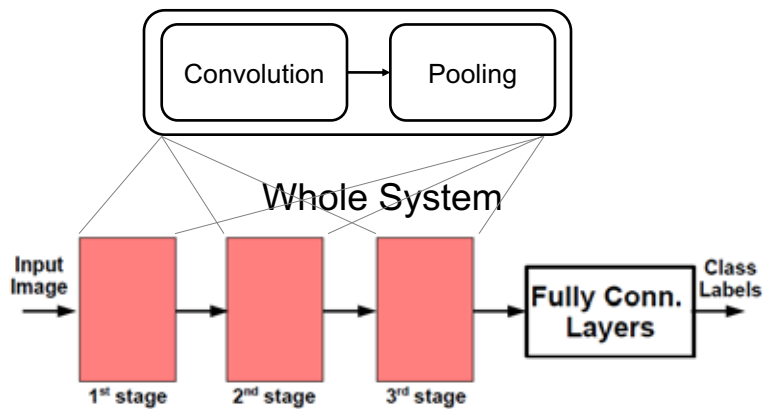
15

Convolution Neural Network (CNN) Building Block



16

Putting it All Together

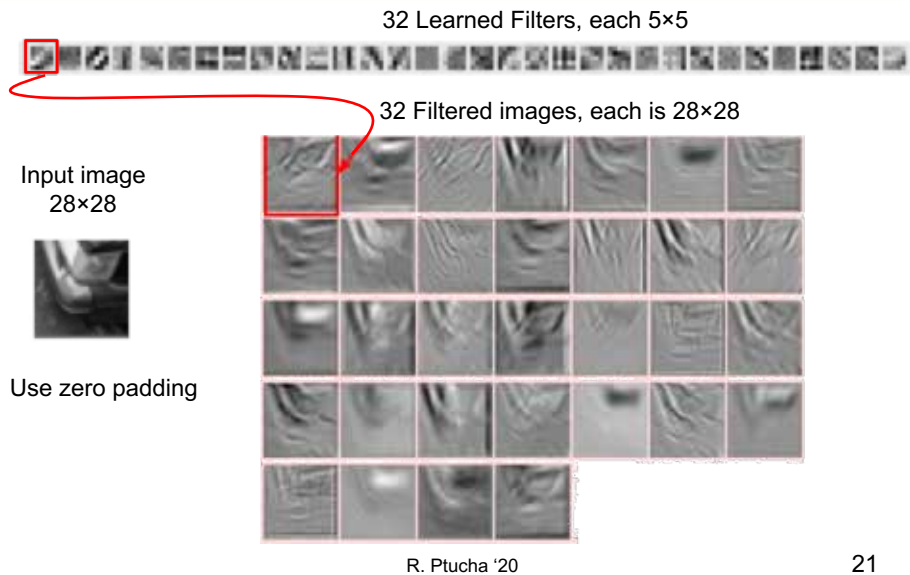


R. Ptucha '20

17

17

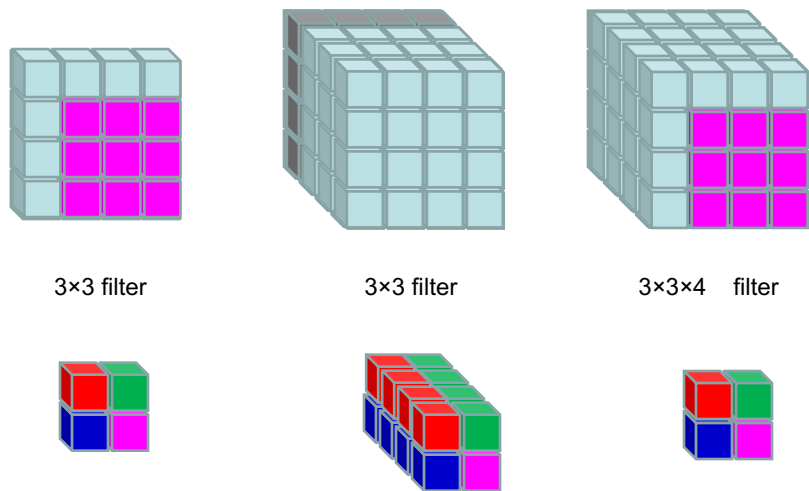
Learning Filters



21

21

Filters

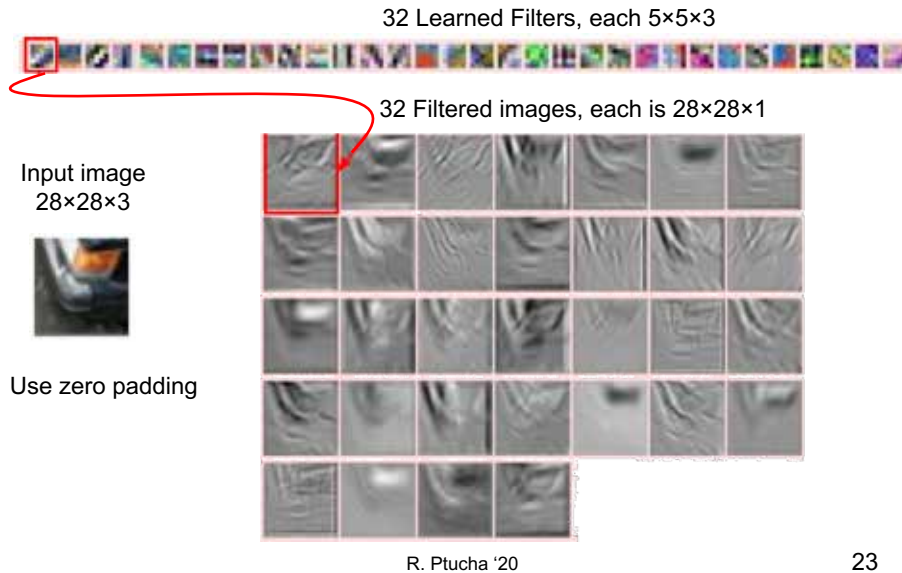


R. Ptucha '20

22

22

Learning Filters



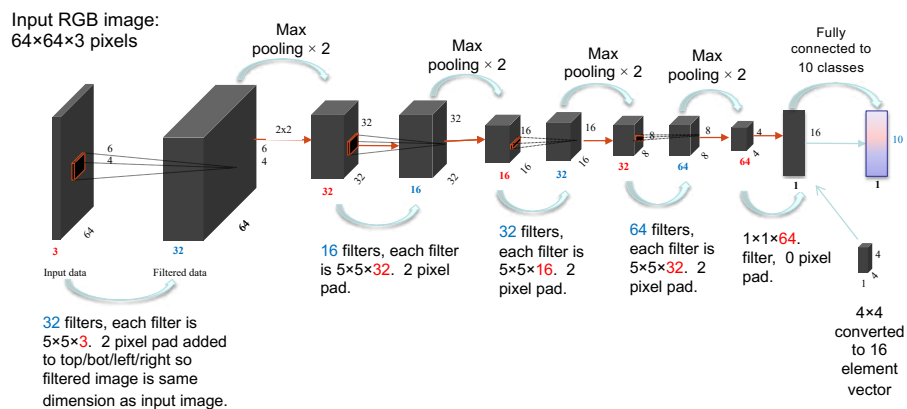
23

23

CNN Architecture

(Not so) Toy Example

Output:
prediction of 1 of
10 categories

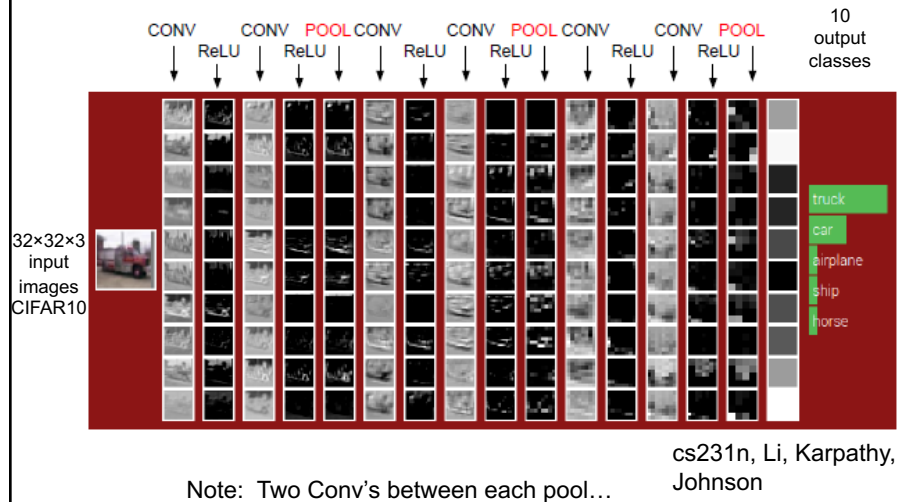


R. Ptucha '20

24

24

CNN Example



R. Ptucha '20

25

25

CNN Example

- Input [32x32x3]
- CONV with ten 3x3 filters, stride 1, pad 1:
 - Parameters: $(3*3*3)*10 + 10 = 280$
 - Memory: $32*32*10$
- CONV with ten 3x3 filters, stride 1, pad 1:
 - Parameters: $(3*3*10)*10 + 10 = 910$
 - Memory: $32*32*10$
- Pool with 2x2 filters, stride 2:
 - Parameters: 0
 - Memory: $16*16*10$

R. Ptucha '20

26

26



CNN Example

Input: $16 \times 16 \times 10$

- CONV with five 3×3 filters, stride 1, pad 1:
 - Parameters: $(3 \times 3 \times 10) \times 5 + 5 = 455$
 - Memory: $16 \times 16 \times 5$
- CONV with twenty 3×3 filters, stride 1, pad 1:
 - Parameters: $(3 \times 3 \times 5) \times 20 + 20 = 920$
 - Memory: $16 \times 16 \times 20$
- Pool with 2×2 filters, stride 2:
 - Parameters: 0
 - Memory: $8 \times 8 \times 20$

R. Ptucha '20

27

27



CNN Example

Input: $8 \times 8 \times 20$

- CONV with ten 3×3 filters, stride 1, pad 1:
 - Parameters: $(3 \times 3 \times 20) \times 10 + 10 = 1810$
 - Memory: $8 \times 8 \times 10$
- CONV with twenty 5×5 filters, stride 1, pad 2:
 - Parameters: $(5 \times 5 \times 10) \times 20 + 20 = 5020$
 - Memory: $8 \times 8 \times 20$
- Pool with 2×2 filters, stride 2:
 - Parameters: 0
 - Memory: $4 \times 4 \times 20$

R. Ptucha '20

28

28



CNN Example

- Fully connect (FC) 4x4x20 to 10 output classes
- Parameters: $(4*4*20)*10 + 10 = 3210$
- Memory: 10
- Done!

R. Ptucha '20

29

29

Case Study

Case study: VGGNet / OxfordNet
(runner-up winner of ILSVRC 2014)
[Simonyan and Zisserman]

best model

cs321n, Karpathy, Li

ConvNet Configuration					
A	A+LRN	B	C	D	E
11 weight layers	11 weight layers	13 weight layers	16 weight layers	15 weight layers	19 weight layers
input (224 x 224 RGB image)					
conv3-64	conv3-64 LRN	conv3-64	conv3-64	conv3-64	conv3-64
maxpool					
conv3-128	conv3-128	conv3-128	conv3-128	conv3-128	conv3-128
maxpool					
conv3-256	conv3-256	conv3-256	conv3-256	conv3-256	conv3-256
conv3-256	conv3-256	conv3-256	conv3-256	conv3-256	conv3-256
maxpool					
conv3-512	conv3-512	conv3-512	conv3-512	conv3-512	conv3-512
conv3-512	conv3-512	conv3-512	conv3-512	conv3-512	conv3-512
maxpool					
conv3-512	conv3-512	conv3-512	conv3-512	conv3-512	conv3-512
conv3-512	conv3-512	conv3-512	conv3-512	conv3-512	conv3-512
maxpool					
FC-4096					
FC-4096					
FC-1000					
softmax					

Table 2. Number of parameters (in millions).

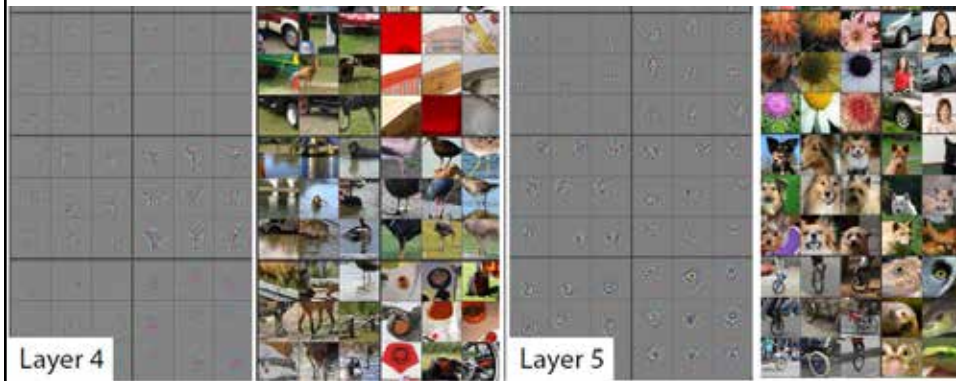
Network	A	A+LRN	B	C	D	E
Number of parameters	133	133	154	138	144	144

R. Ptucha '20

30

30

CNN Visualization



Zeiler, Fergus, 2014

R. Ptucha '20

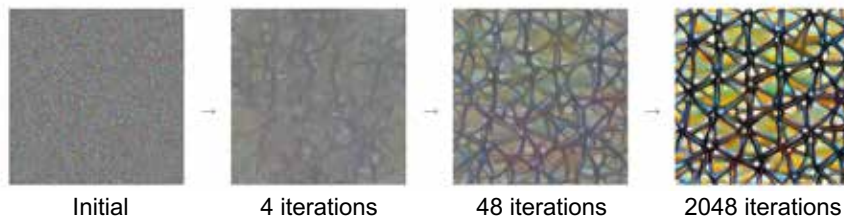
33

33

CNN Visualization

- We can compute the partial derivative of input pixels with respect to a cost.
- Start with random noise image, pretrained network, then iteratively tweak the input as we minimize our cost.
 - Cost is to activate a particular neuron (in this slide).

Google Inception v1: layer mixed4a, unit 11



Olah, et al., "Feature Visualization", Distill, 2017.

R. Ptucha '20

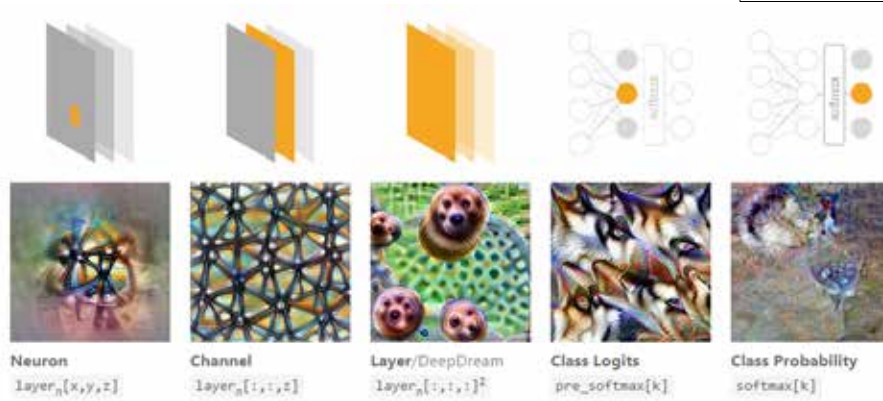
34

34

CNN Visualization

- Can modify the objective to get different types of insight to what the CNN is responding to.

n layer index
 x, y spatial position
 z channel index
 k class index



Olah, et al., "Feature Visualization", Distill, 2017.

R. Ptucha '20

35

Interrogate a network with different gradients

Start with any image, then when do back prop, use GT label of say a dog, then instead of updating the model weights, update the pixels in the image.



Inceptionism: Going Deeper into Neural Networks

A. Mordvintsev, C. Olah and M. Tyka (2015)

R. Ptucha '20

36

Interrogate a network with different gradients



Delta is the difference between logits and one-hot encoding of dog. After backprop, the next time you pass in the image, it will be closer to being classified as a dog!



Then iteratively keep updating pixels with dog GT label....eventually, the image will start to "look" like a dog

Inceptionism: Going Deeper into Neural Networks

A. Mordvintsev, C. Olah and M. Tyka (2015)

R. Ptucha '20

37

37

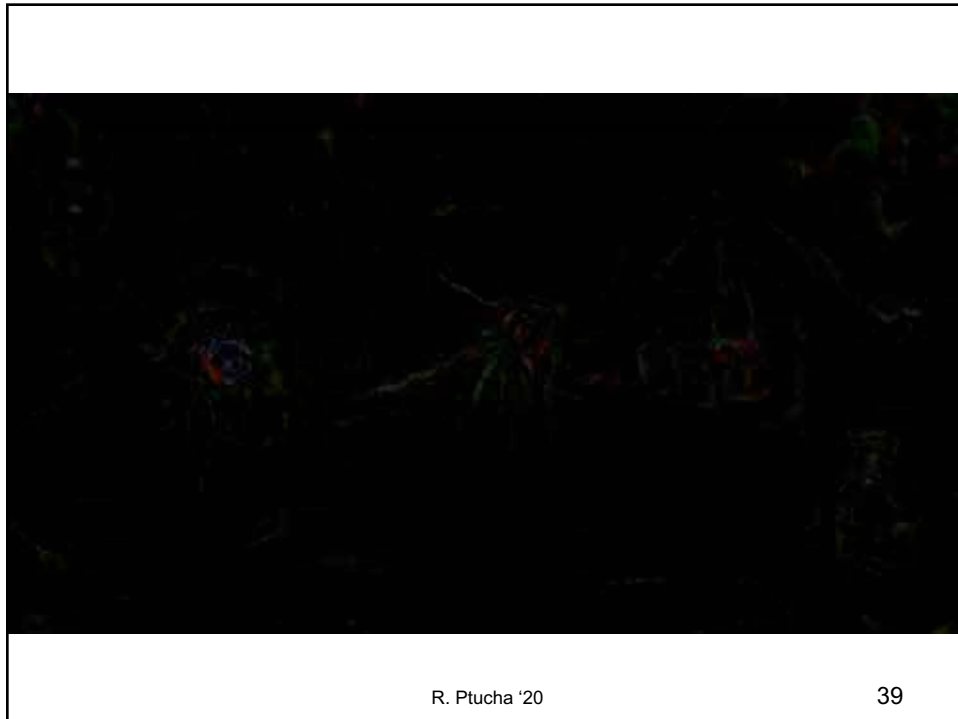


A. Mordvintsev, C. Olah and M. Tyka,
<http://googleresearch.blogspot.com/2015/06/inceptionism-going-deeper-into-neural.html>

R. Ptucha '20

38

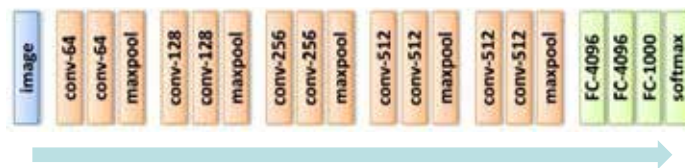
38



39

CNN as Vector Representation

Typical CNN Architecture



Input Image



2D Plot of fc8 Feature Vector

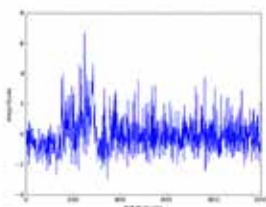
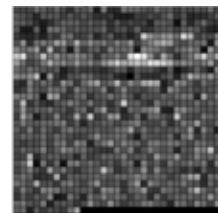


Image of fc8 Feature Vector



R. Ptucha '20

40

40

CNN as Vector Representation



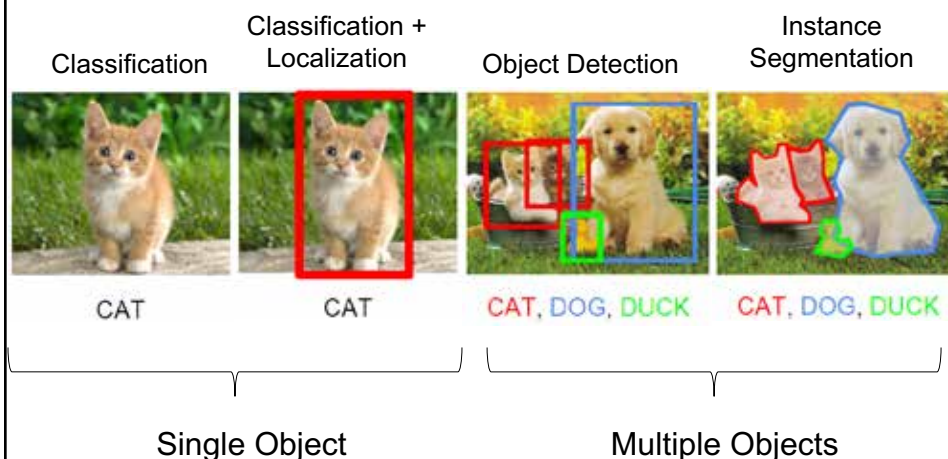
- As it turns out, these fully connected layers are excellent descriptors of the input image!
- For example, you can pass images through a pre-trained CNN, then take the output from a FC layer as input to a SVM classifier. (image2vec)
- Images in this vector space generally have the property that similar images are close in this latent representation.

R. Ptucha '20

41

41

Vision Tasks



R. Ptucha '20

53

53

Classification vs. Classification + Localization

Classification

Input: Image
Output: Class label
Evaluation metric: Accuracy



→ CAT

Classification + Localization

Input: Image
Output: Class label, Box coordinates
Evaluation metric: Intersection over Union (IoU)



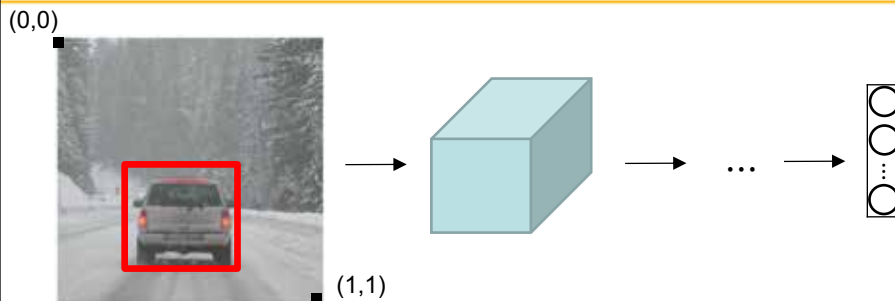
→ (CAT, x, y, w, h)

R. Ptucha '20

54

54

Classification with Localization



Lets allow a few classes:

1. Car
2. Truck
3. Pedestrian
4. Motorcycle

- For now, lets assume one object per image.
- Each object has $\{x, y, w, h\}$
- For this image, object location $\{x, y, w, h\} = \{0.3, 0.6, 0.4, 0.3\}$

Image from: deeplearning.ai, C4W3L01

R. Ptucha '20

55

55

Classification with Localization

Four classes:

1. Car
2. Truck
3. Pedestrian
4. Motorcycle

Localization $\{x, y, w, h\}$



Define y label: $y = \begin{bmatrix} P_c \\ b_x \\ b_y \\ b_w \\ b_h \\ C_1 \\ C_2 \\ C_3 \\ C_4 \end{bmatrix}$

Probability of an object

Bounding box location

0/1 for each class

Image from: deeplearning.ai, C4W3L01

R. Ptucha '20

56

56

Classification with Localization

Four classes:

1. Car
2. Truck
3. Pedestrian
4. Motorcycle

Localization $\{x, y, w, h\}$

Cost function (squared error):

$$Loss = \sum_{i=1}^9 (\hat{y}_i - y_i)^2 \quad \text{If } y_i=1$$

$$Loss = (\hat{y}_1 - y_1)^2 \quad \text{If } y_1=0$$



$y = \begin{bmatrix} 1 \\ 0.3 \\ 0.6 \\ 0.4 \\ 0.3 \\ 0 \\ 1 \\ 0 \\ 0 \\ 0 \end{bmatrix}$

$\begin{bmatrix} P_c \\ b_x \\ b_y \\ b_w \\ b_h \\ C_1 \\ C_2 \\ C_3 \\ C_4 \end{bmatrix}$

$y = \begin{bmatrix} 0 \\ ? \\ ? \\ ? \\ ? \\ ? \\ ? \\ ? \\ ? \end{bmatrix}$? = don't care

Image from: deeplearning.ai, C4W3L01

R. Ptucha '20

57

57

Classification with Localization

Four classes:

1. Car
2. Truck
3. Pedestrian
4. Motorcycle

Localization $\{x, y, w, h\}$

Alternate cost function:

- y_1 can be logistic loss
- $y_2 \rightarrow y_5$ can be squared error
- $y_6 \rightarrow y_9$ can be softmax cross entropy



$$y = \begin{bmatrix} 1 \\ 0.3 \\ 0.6 \\ 0.4 \\ 0.3 \\ 0 \\ 1 \\ 0 \\ 0 \end{bmatrix}$$

$$\begin{bmatrix} P_c \\ b_x \\ b_y \\ b_w \\ b_h \\ C_1 \\ C_2 \\ C_3 \\ C_4 \end{bmatrix}$$

$$y = \begin{bmatrix} 0 \\ ? \\ ? \\ ? \\ ? \\ ? \\ ? \\ ? \\ ? \end{bmatrix}$$

? = don't care

Image from: deeplearning.ai, C4W3L01

R. Ptucha '20

58

58

Snapchat Facewarp?



- Traditional approach:

Viola Jones
Face Detection

Search for actual point locations
using Mahalanobis distance



Repeat ~3-5x

Average eye and 82
facial feature points

Restrict based on
PCA statistics

R. Ptucha '20

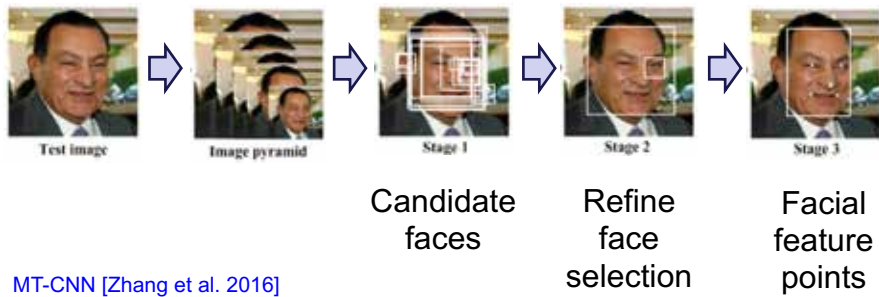
59

59

Snapchat Facewarp?



- Deep Learning approach:



R. Ptucha '20

60

60

Localization

- Facial feature



Each face has 68 points,
so CNN would output:

Face?

pt1X

pt1Y

pt2X

pt2Y

.

.

.

pt68X

pt68Y

137 outputs

Of course,
need GT for
thousands of
faces to train
model.

ptucha R. Ptucha '20

61

61

Can do same with Body Pose...



Pishchulin et al. CVPR'16

R. Ptucha '20

62

62

Object Detection

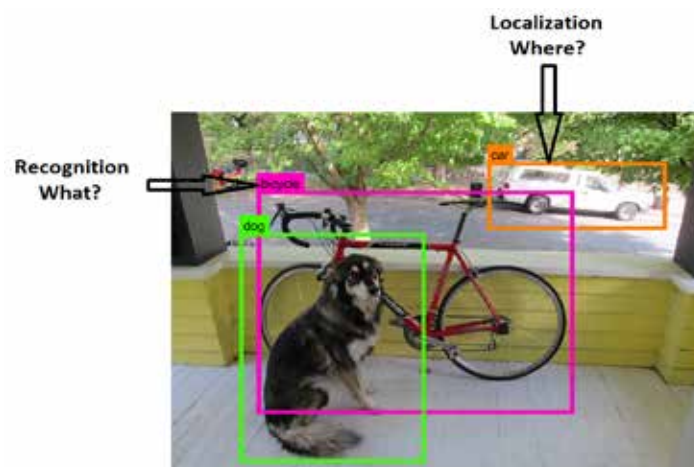


Image Credit: Redmon, Joseph, et al [4]

R. Ptucha '20

63

63

More than one object per image?

Training set:



Car detection example



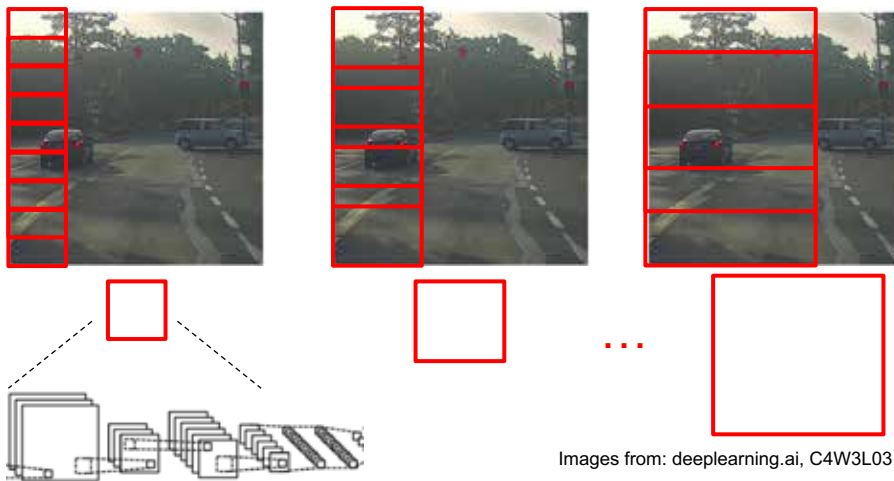
Images from: deeplearning.ai, C4W3L03

R. Ptucha '20

66

66

Sliding Window Detection



Images from: deeplearning.ai, C4W3L03

ptucha

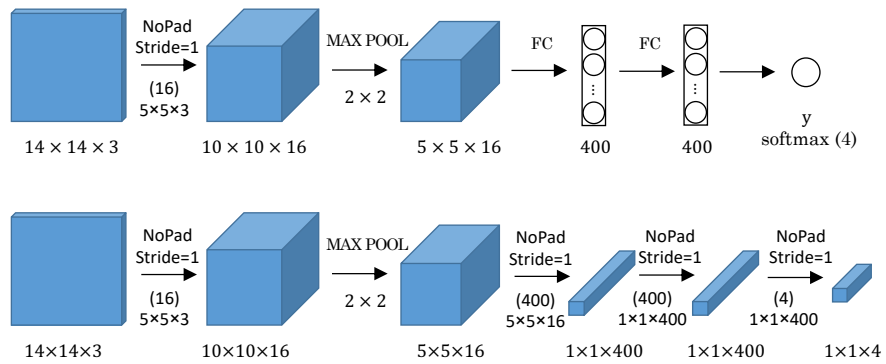
R. Ptucha '20

67

67

67

Computing FC layers with Convolution



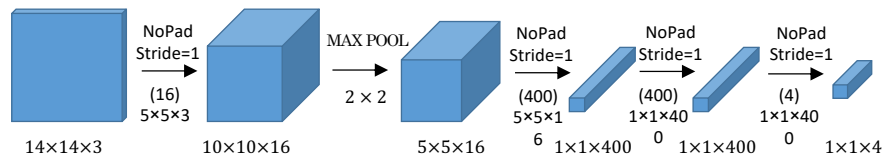
Images from: deeplearning.ai, C4W3L04

R. Ptucha '20

68

68

Replacing Sliding Windows w/Fully Convolutional CNNs

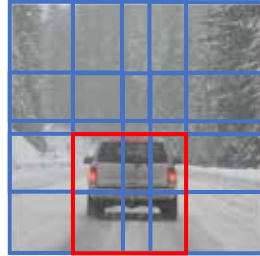


69

Replacing Sliding Windows w/Fully Convolutional CNNs

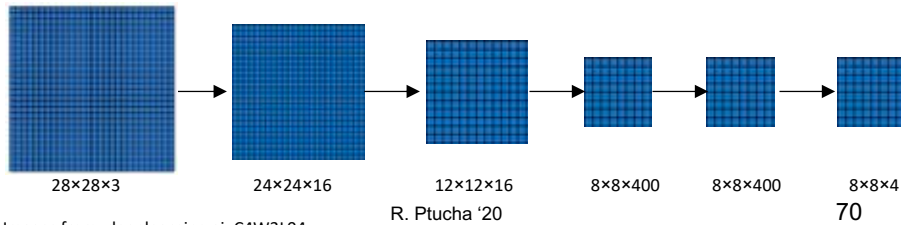
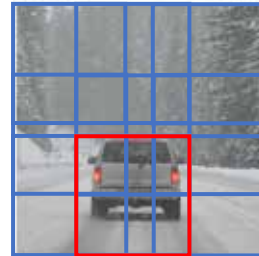
Sliding window approach:

Sequentially evaluate one window at a time



Fully convolutional approach:

Evaluate 64 regions at once



Images from: deeplearning.ai, C4W3L04

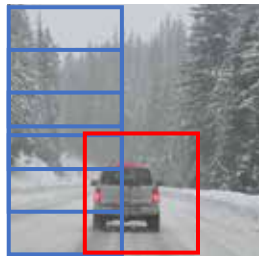
R. Ptucha '20

70

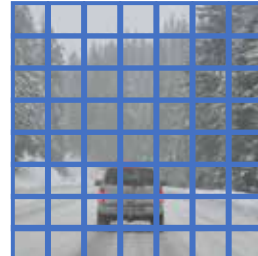
70

More Accurate Bounding Boxes

- Whether using sliding window or fully convolutional approach, bounding boxes often don't align well or may not match actual object.



$$y = \begin{bmatrix} P_c \\ b_x \\ b_y \\ b_w \\ b_h \\ C_{pers} \\ C_{car} \\ C_{truck} \\ C_{motorc} \end{bmatrix}$$



- At each location, solve for probability of object as well as actual coordinates, which don't have to be square.

Images from: deeplearning.ai, C4W3L05

R. Ptucha '20

71

71

Replacing Sliding Windows w/Fully Convolutional CNNs

- Can think of this as evaluating 8×8 grid, where each of the 64 cells is independently checked for an object:

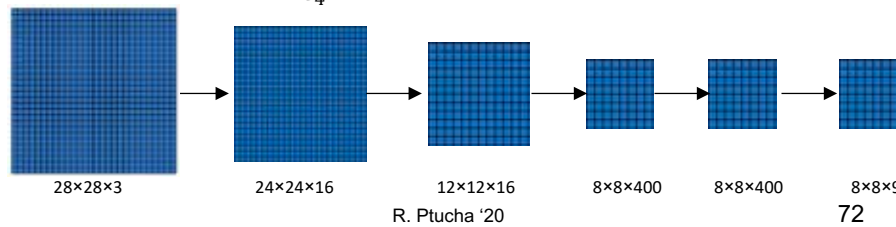
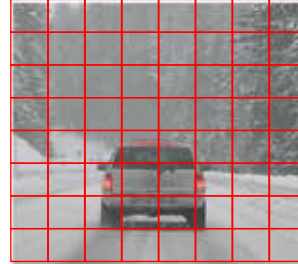
Each cell has a y label:

$$y = \begin{bmatrix} P_c \\ b_x \\ b_y \\ b_w \\ b_h \\ C_1 \\ C_2 \\ C_3 \\ C_4 \end{bmatrix}$$

Prob. of an object

Object location

0/1 for each class
(Four classes in this example)



72

Replacing Sliding Windows w/Fully Convolutional CNNs

- Overlay GT of object
- Cell where centroid lie is responsible.

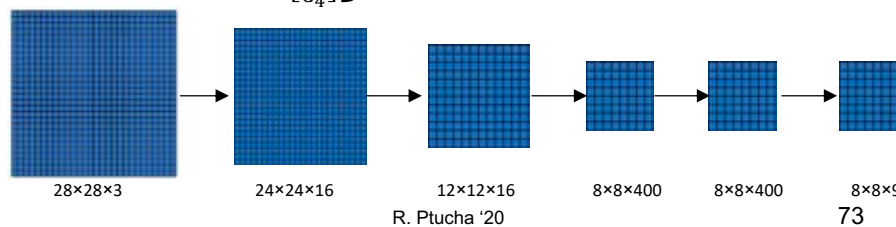
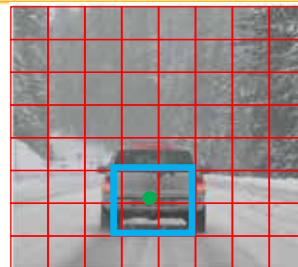
Each cell has a y label:

$$y = \begin{bmatrix} P_c \\ b_x \\ b_y \\ b_w \\ b_h \\ C_1 \\ C_2 \\ C_3 \\ C_4 \end{bmatrix}$$

Prob. of an object

Object location

0/1 for each class
(Four classes in this example)



73

72

73

Mask R-CNN

<https://arxiv.org/pdf/1703.06870.pdf>



FAIR Mask R-CNN, COCO + Places Workshop, ICCV 2017

R. Ptucha '20

76

76

CNN Results

- Handwritten characters
 - MNIST: 0.17% error, Ciresan et al '11
 - Arabic & Chinese: Ciresan '12
- CIFAR-10 (60K images of 10 classes)
 - 9.3% error, Wan et al. '13
- Traffic Sign Recognition
 - 0.56% error vs 1.16% for humans, Ciresan '11



R. Ptucha '20

77

77

ImageNet

- Amazon Turk did bulk of labeling
- 14M labeled images
- 20K classes



Russakovsky et al., 2015

IMAGENET Large Scale Visual Recognition Challenge (ILSVRC)

- 1.2M images, 1000 categories
- Image classification, object localization, video detection

R. Ptucha '20

78

78

ImageNet: Examples of Hammer



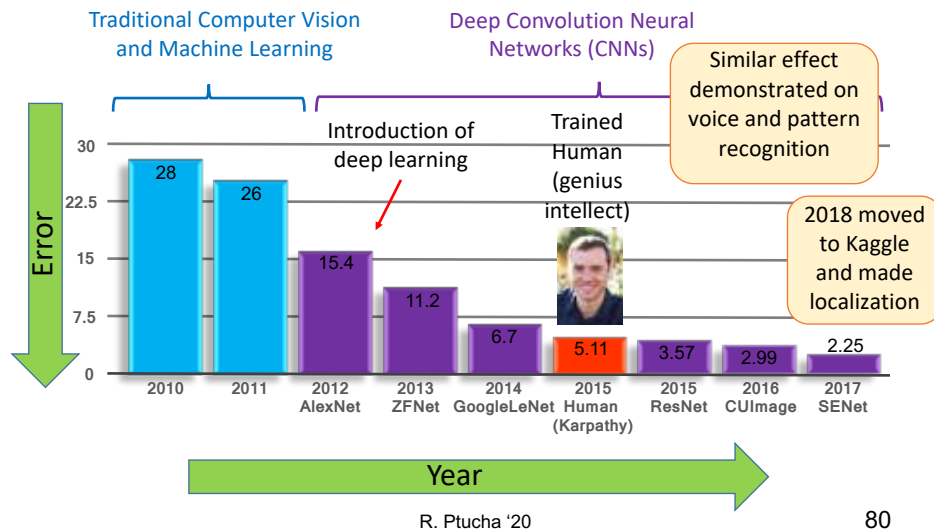
R. Ptucha '20

79

79

Deep Learning- Surpassing The Visual Cortex's Object Detection and Recognition Capability

Top-5 error on ImageNet



80



Andrew Ng, 2017

<https://www.deeplearning.ai/>

Deep Learning Specialization, Five courses:

1. Neural Networks and Deep Learning
2. Improving Deep Neural Networks
3. Structured Machine Learning Projects
4. Convolutional Neural Networks
5. Sequence Models

R. Ptucha '20

81

81



CS231n: Convolutional Neural Networks for Visual Recognition



Fei-Fei Li



Justin Johnson



Serena Young

Li, Johnson, Yeung 2017

<http://cs231n.stanford.edu/>

R. Ptucha '20

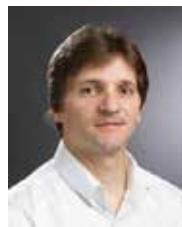
82

82



DEEP
LEARNING
INSTITUTE

For More Information: <http://www.rit.edu/mil>



Raymond W. Ptucha

Assistant Professor, Computer Engineering
Director, Machine Intelligence Laboratory
Rochester Institute of Technology

Email: rwpeec@rit.edu

Phone: +1 (585) 797-5561

Office: GLE (09) 3441



R. Ptucha '20

83

83