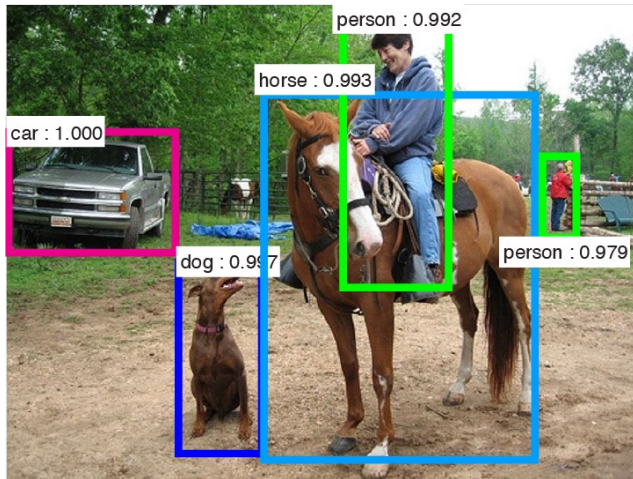


Recognition

Convolutional neural networks and sequence models for segmentation
object detection, and classification



CSC420

David Lindell

University of Toronto

cs.toronto.edu/~lindell/teaching/420

Slide credit: Babak Taati ← Ahmed Ashraf ← Sanja Fidler; Bill Freeman, Phillip Isola, Tianhong Li, Justin Johnson

Course Overview

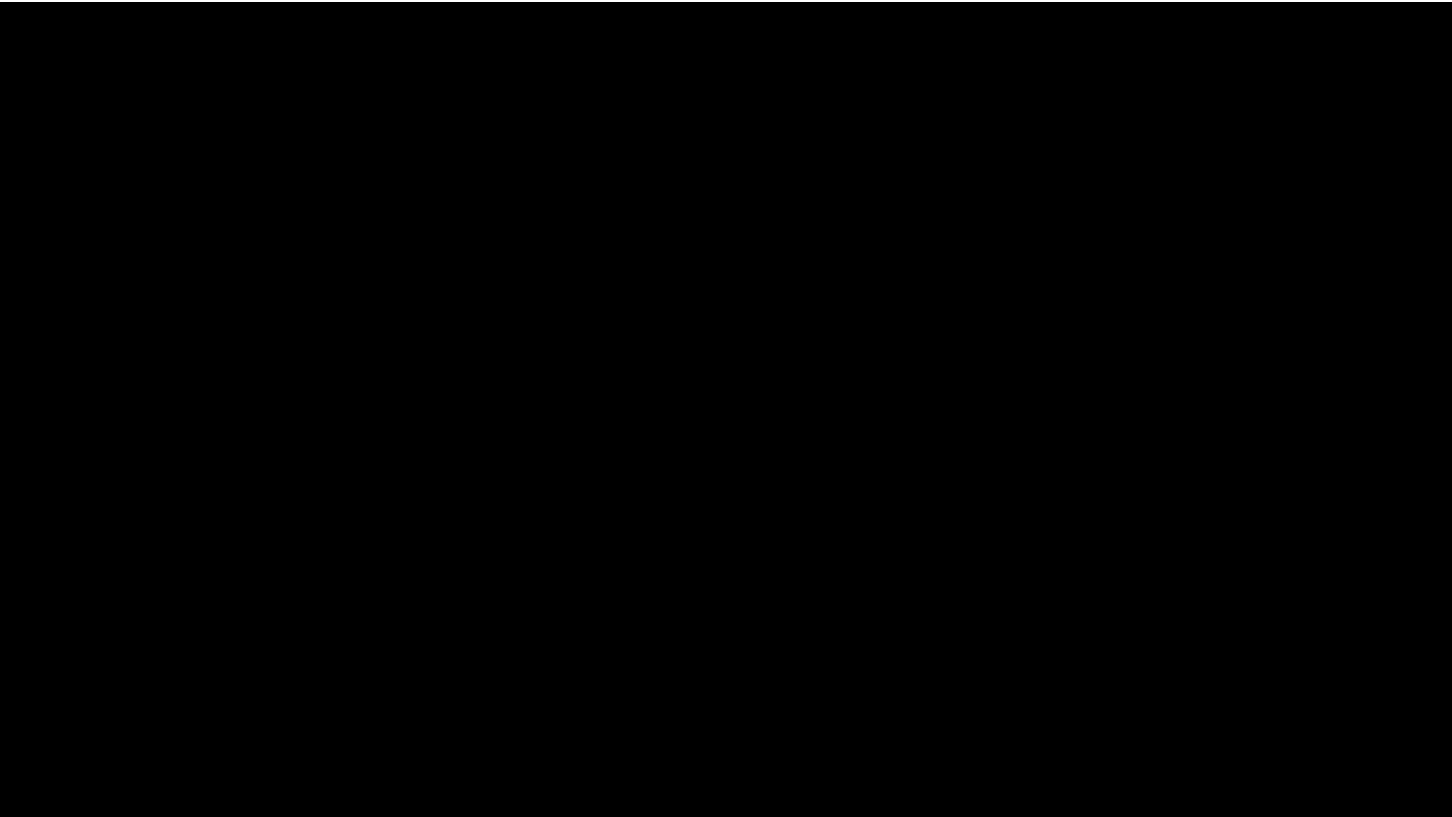


Outline

- motivation / overview of recognition tasks
- CNNs
 - Classification
 - Segmentation
 - Object Detection

Outline

- motivation / overview of recognition tasks
- CNNs
 - Classification
 - Segmentation
 - Object Detection



Visual Recognition Tasks

Classification: what
output: class label

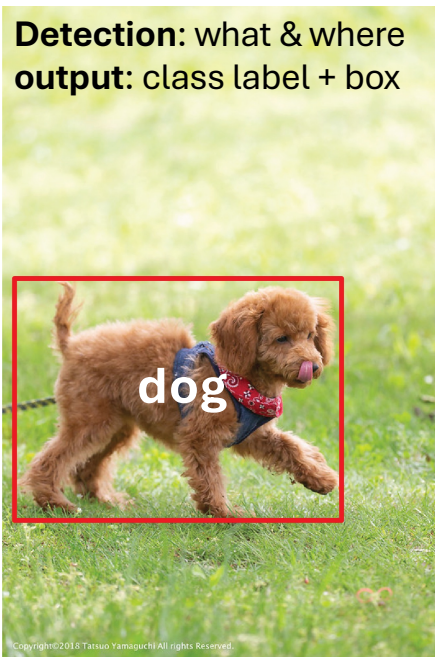


Visual Recognition Tasks

Classification: what
output: class label



Detection: what & where
output: class label + box

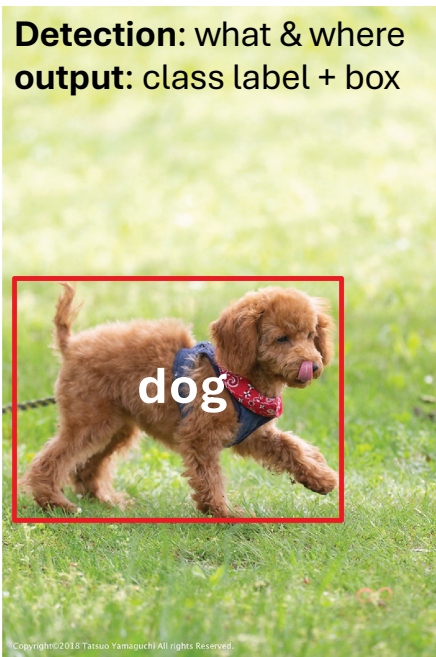


Visual Recognition Tasks

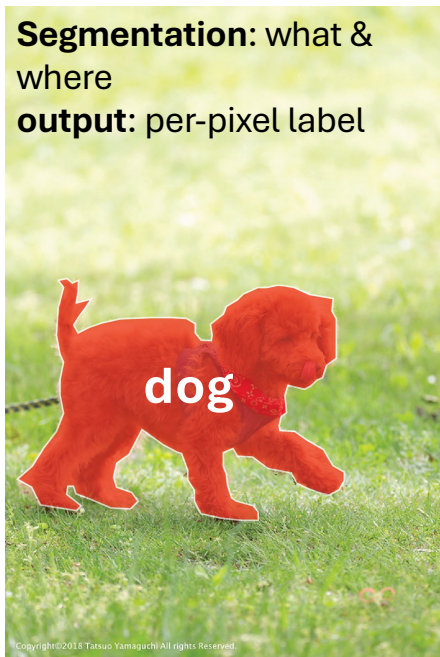
Classification: what
output: class label



Detection: what & where
output: class label + box



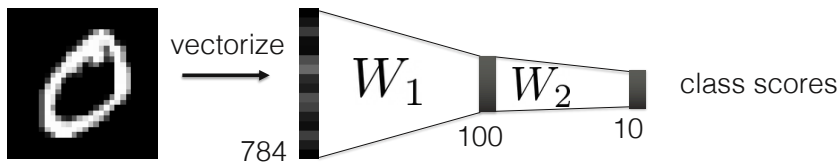
Segmentation: what & where
output: per-pixel label



Outline

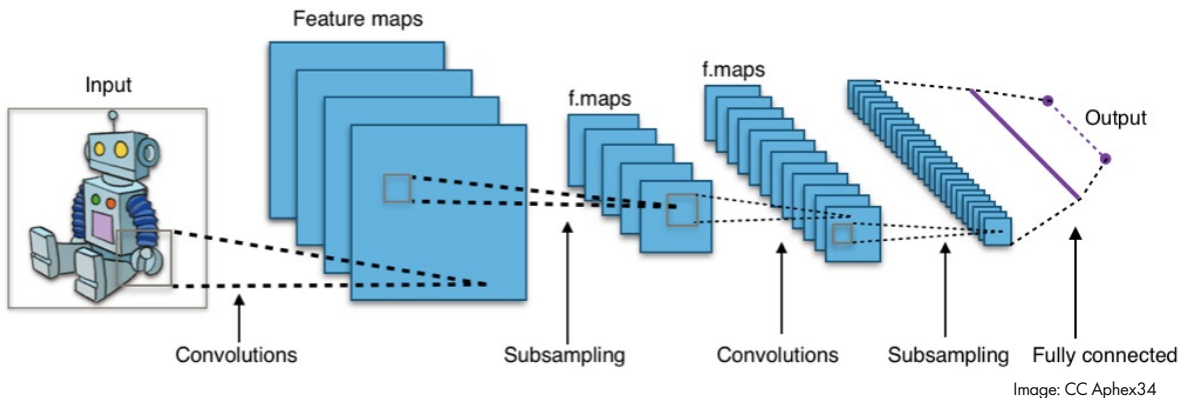
- motivation / overview of recognition tasks
- CNNs
 - Classification
 - Segmentation
 - Object Detection

Drawbacks of fully-connected networks



- spatial structure is destroyed
- fully-connected weights do not scale

Convolutional Neural Networks



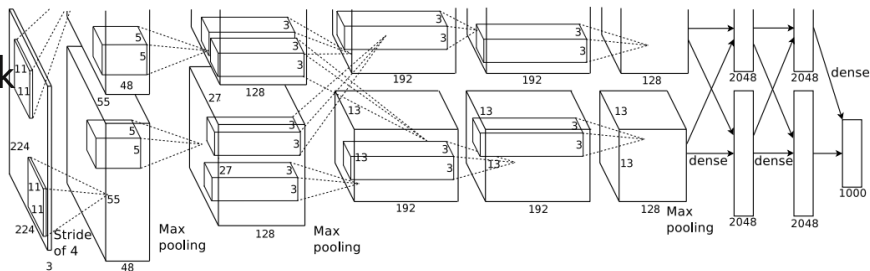
- Exploit spatial structure
- Scale to large inputs with fewer parameters
- Remarkable performance for processing visual data

AlexNet & surge in popularity

2010: ImageNet Large Scale Visual Recognition Challenge

- 14 million labeled images

First convolutional network
for image classification



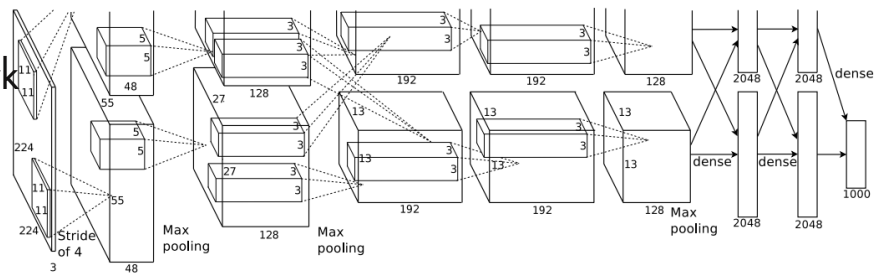
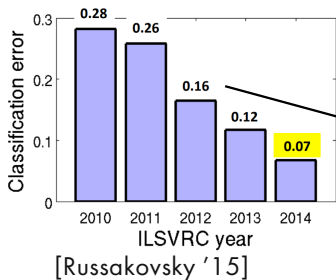
AlexNet [Krizhevsky '12]

AlexNet & surge in popularity

2010: ImageNet Large Scale Visual Recognition Challenge

- 14 million labeled images

First convolutional network
for image classification



AlexNet [Krizhevsky '12]

CNNs

AlexNet & surge in popularity

2010: ImageNet Large Scale Visual Recognition Challenge

- 14 million labeled images



Deep networks



VGG

[Simonyan '14]



GoogLeNet

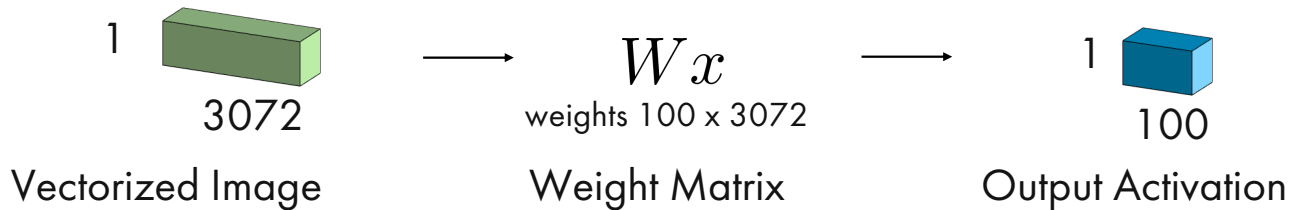
[Szegedy '14]



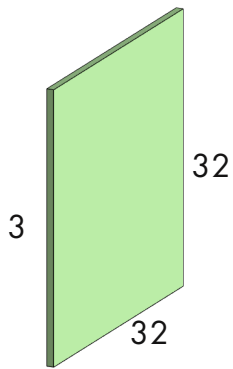
ResNet

[He '15]

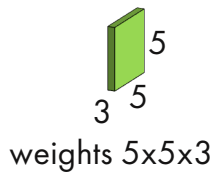
Fully-Connected Layer



Convolutional Layer

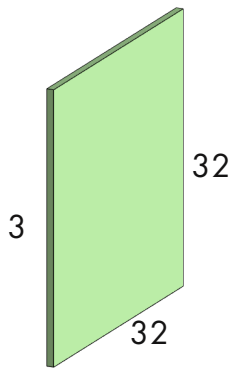


Input Image

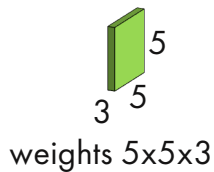


Filter

Convolutional Layer



Input Image

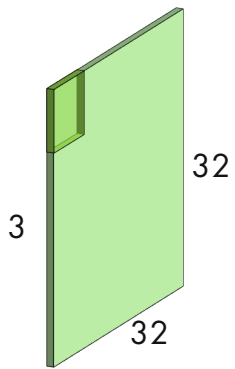


Filter

"Channel" dimension

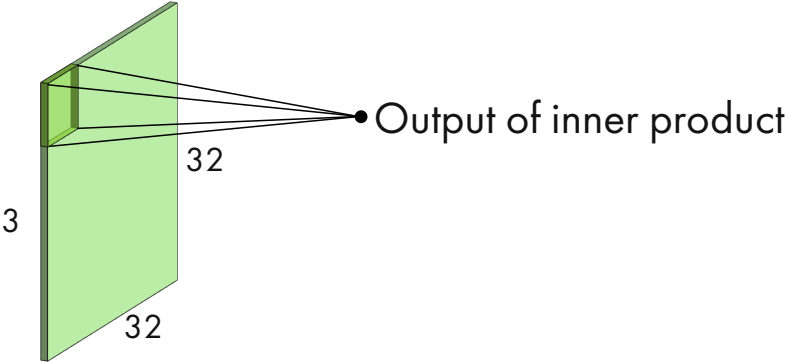


Convolutional Layer



Input Image

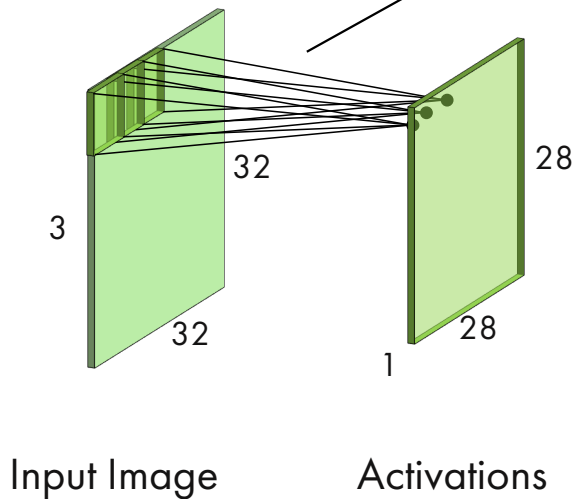
Convolutional Layer



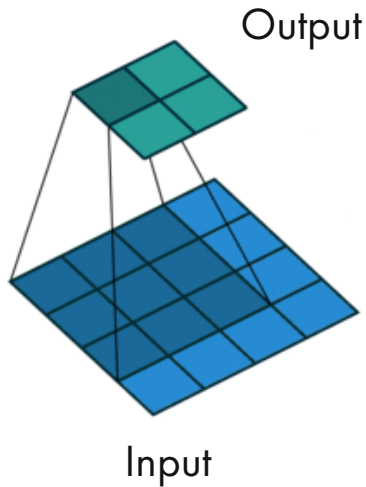
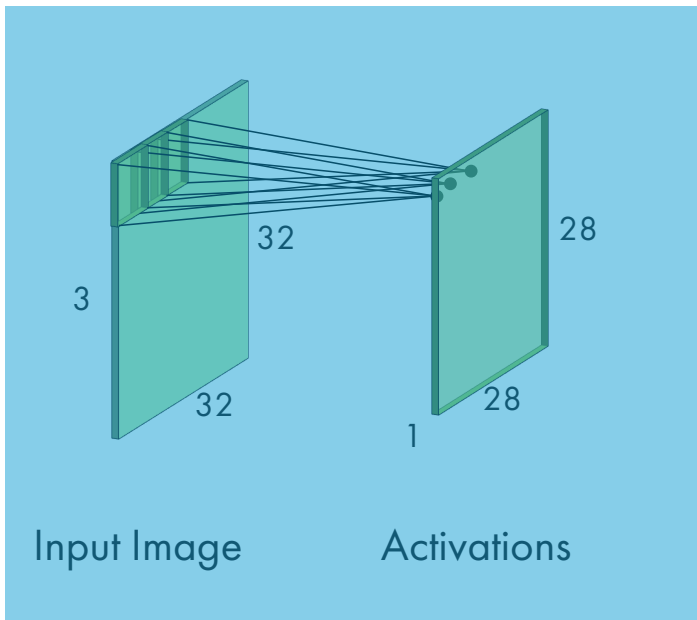
Input Image

Convolutional Layer

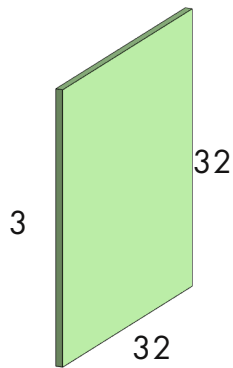
Convolution = sliding window + inner product



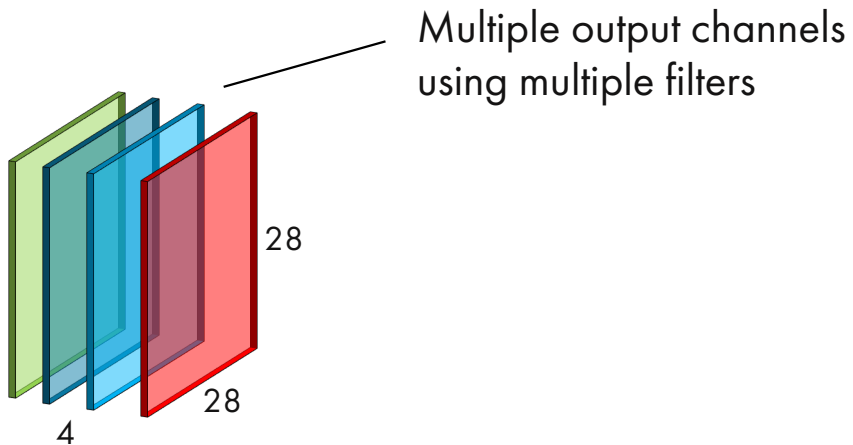
Convolutional Layer



Convolutional Layer

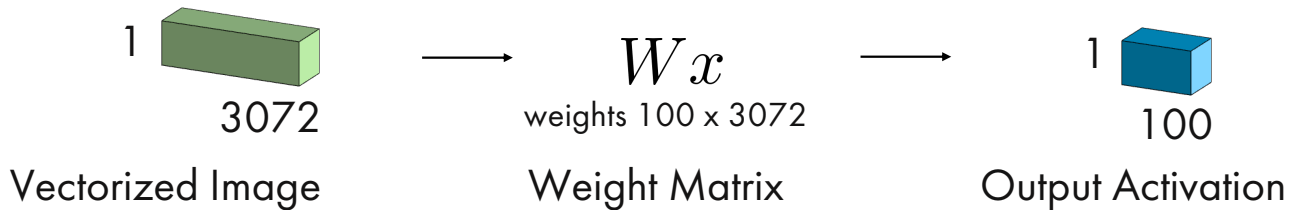


Input Image



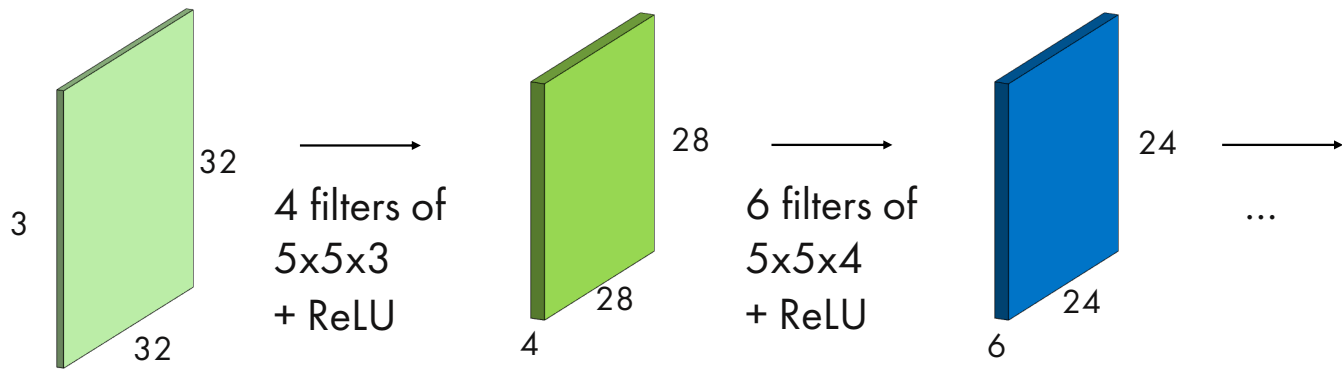
Activations

Fully-Connected Layer



Special case of convolutional layer when filter size = input size!

Convolutional Neural Network



Input Image

Layer 1
Activations

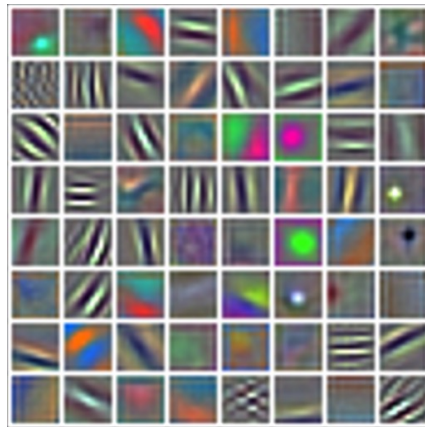
Layer 2
Activations

Case Study: AlexNet

Input Image



First-layer Filters

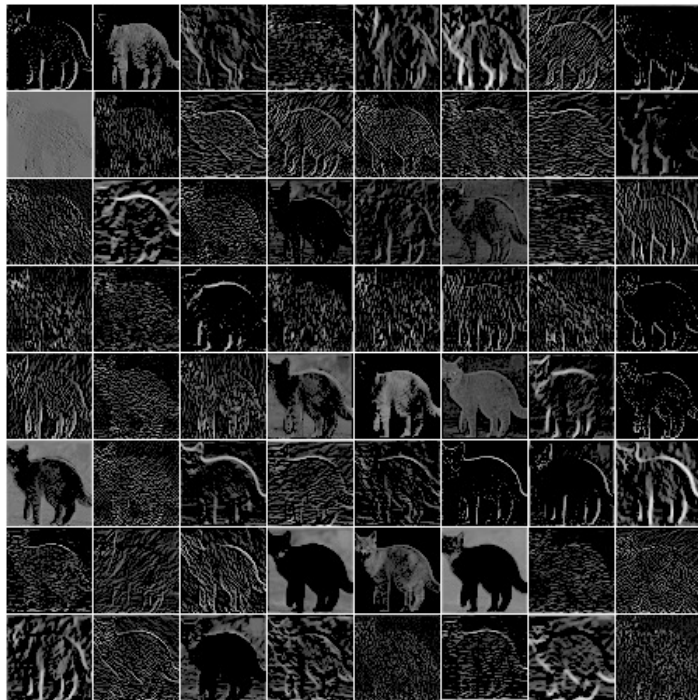
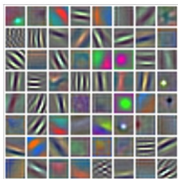


Case Study: AlexNet

Input Image

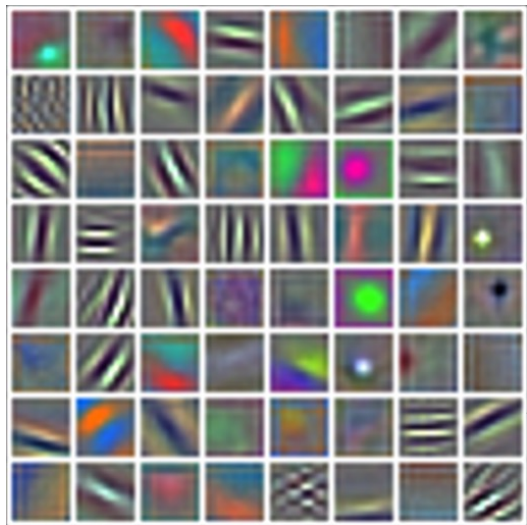


First-layer Filters



Activations

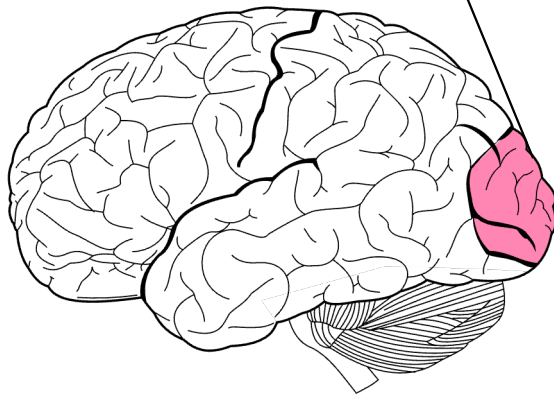
Case Study: AlexNet



First-layer Filters

Similar to simple cells in
visual cortex!

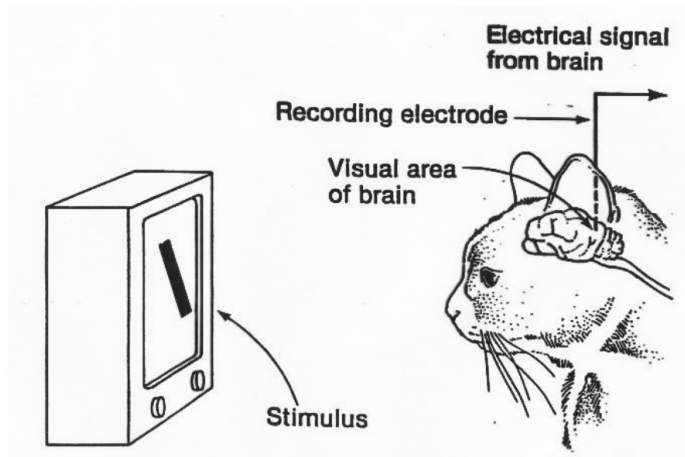
- Edge detectors



Case Study: AlexNet



[Hubel & Wiesel 1959]



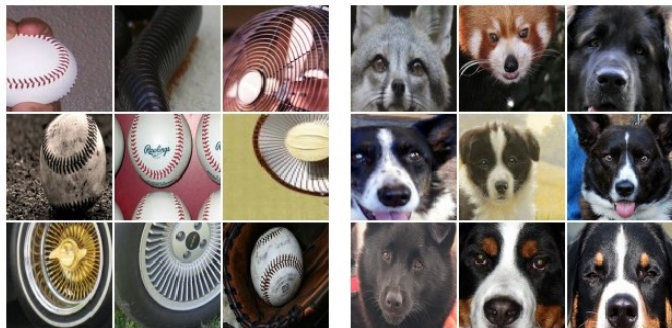
Simple cells in visual cortex
detect edges, complex cells
compose earlier responses

CNN higher layer filters



Dataset examples that maximize neuron outputs

CNN higher layer filters



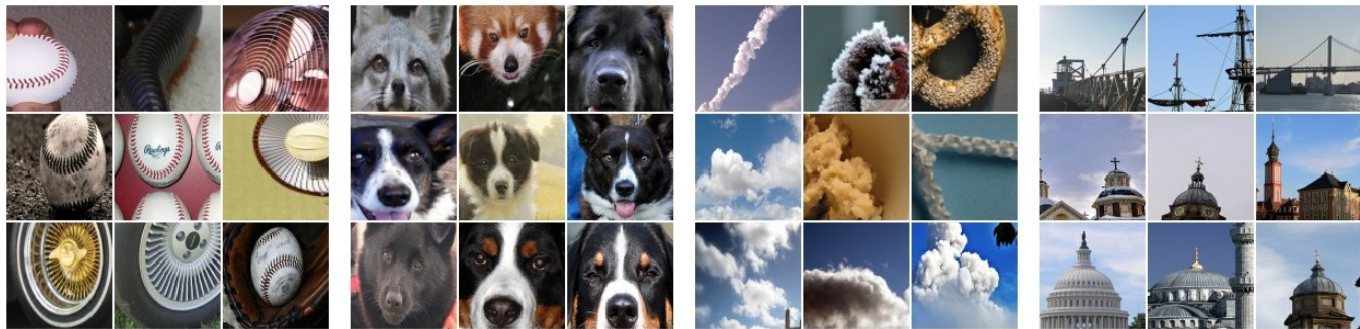
Dataset examples that maximize neuron outputs

CNN higher layer filters



Dataset examples that maximize neuron outputs

CNN higher layer filters



Dataset examples that maximize neuron outputs

CNN Building Blocks

Design choices:

- filter size
- number of filters
- padding
- stride

Layer types:

- pooling
- transpose convolutions
- upsampling layers*
- batch normalization*
- softmax layers*

* no time to cover all of these layers! Check out PyTorch docs for details...

e.g., <https://pytorch.org/docs/stable/generated/torch.nn.BatchNorm2d.html>

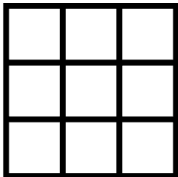
CNN Building Blocks

Filter size

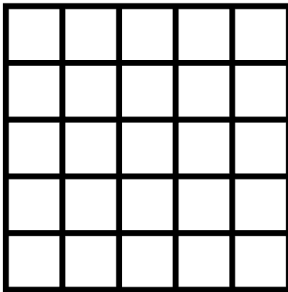
1x1



3x3



5x5

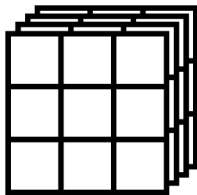


CNN Building Blocks

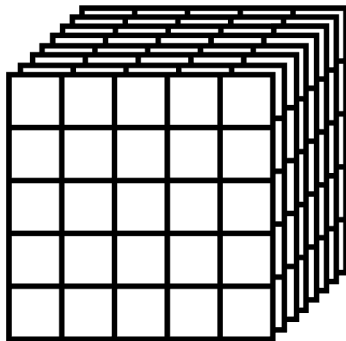
Number of channels

$N_{\text{out}} \times N_{\text{in}} \times 3 \times 3$

$N_{\text{out}} \times N_{\text{in}} \times 1 \times 1$

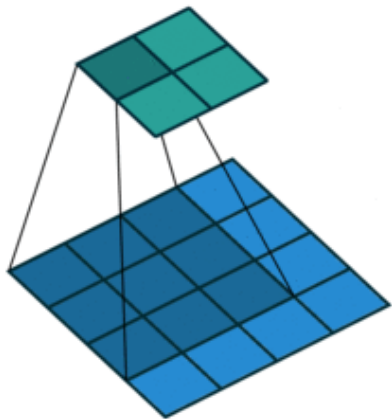


$N_{\text{out}} \times N_{\text{in}} \times 5 \times 5$

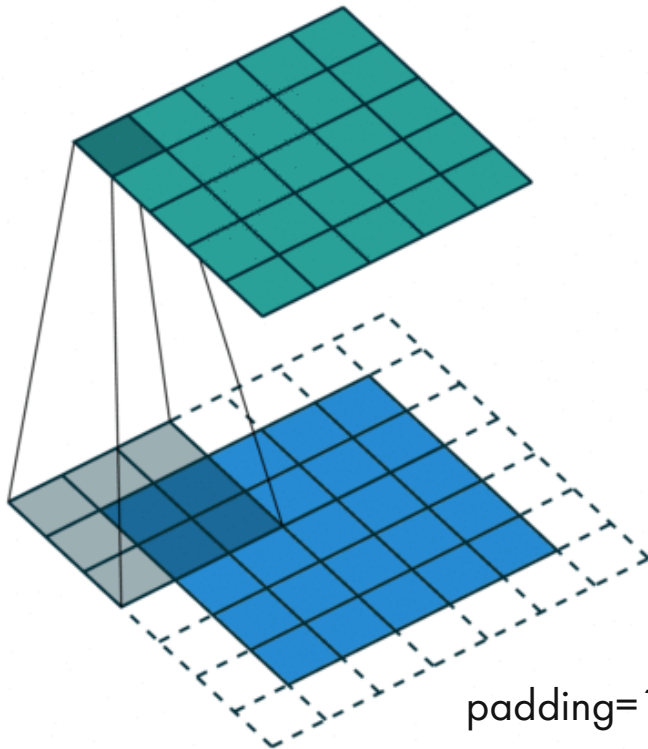


CNN Building Blocks

padding



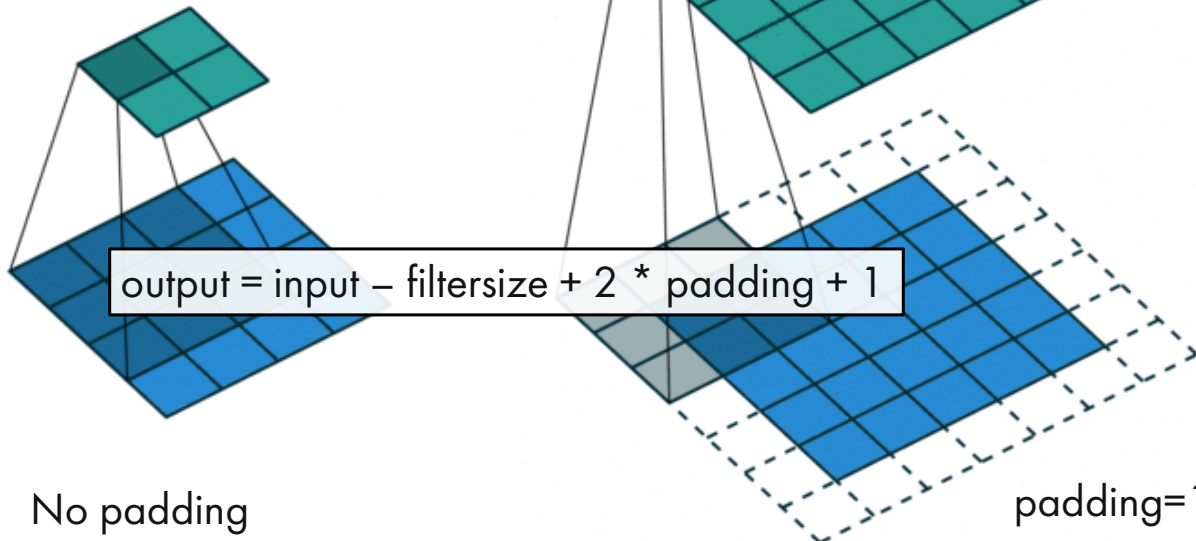
No padding



padding=1

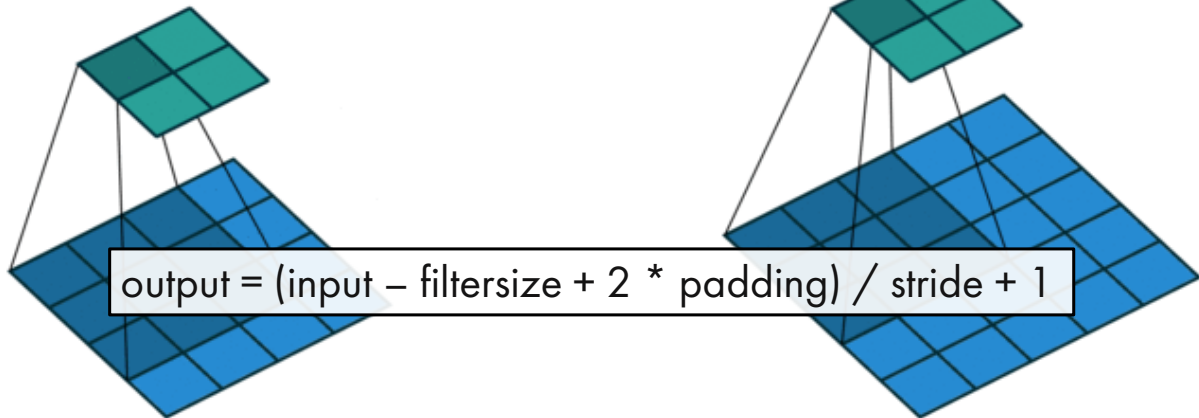
CNN Building Blocks

padding



CNN Building Blocks

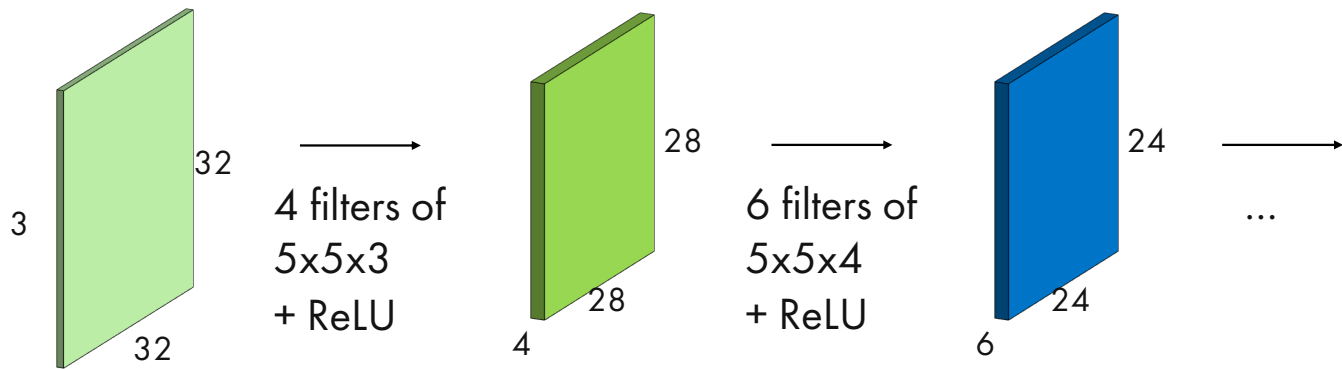
stride



stride = 1

stride = 2

Convolutional Neural Network

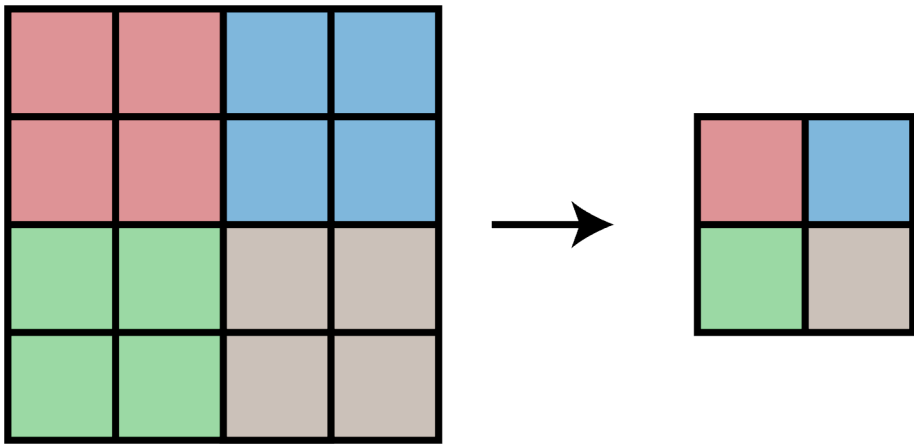


Input Image

Layer 1
Activations

Layer 2
Activations

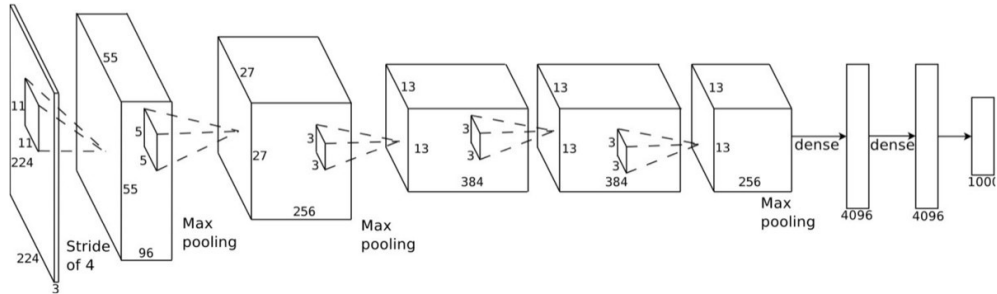
Layer types: Pooling



e.g., max pool size=2, stride=2

Convolutional Neural Networks (CNN)

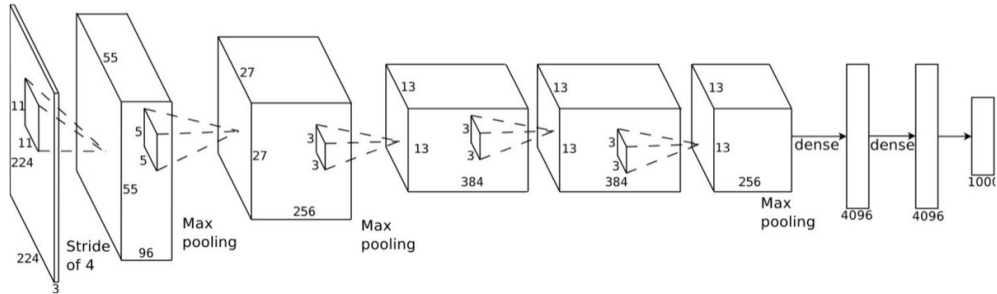
- AlexNet (from UofT!): A. Krizhevsky, I. Sutskever, G. E. Hinton, ImageNet Classification with Deep Convolutional Neural Networks, NeurIPS 2012. This network won the Imagenet Challenge of 2012, and revolutionized computer vision.



[Pic adopted from: A. Krizhevsky]

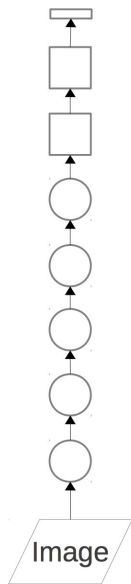
Convolutional Neural Networks (CNN)

- AlexNet (from UofT!): A. Krizhevsky, I. Sutskever, G. E. Hinton, ImageNet Classification with Deep Convolutional Neural Networks, NeurIPS 2012. This network won the Imagenet Challenge of 2012, and revolutionized computer vision.
- How many parameters (weights) does this network have?



[Pic adopted from: A. Krizhevsky]

Convolutional Neural Networks (CNN)



- Trained with stochastic gradient descent on two NVIDIA GPUs for about a week
- 650,000 neurons
- 60,000,000 parameters
- 630,000,000 connections
- **Final feature layer: 4096-dimensional**



Convolutional layer: convolves its input with a bank of 3D filters, then applies point-wise non-linearity



Fully-connected layer: applies linear filters to its input, then applies point-wise non-linearity

Figure: From <http://www.image-net.org/challenges/LSVRC/2012/supervision.pdf>

[Pic adopted from: A. Krizhevsky]

Outline

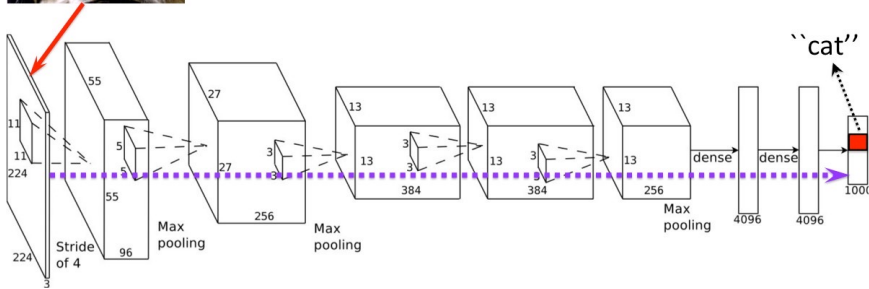
- motivation / overview of recognition tasks
- CNNs
 - Classification
 - Segmentation
 - Object Detection

Classification

- Once trained we can do classification. Just feed in an image or a crop of the image, run through the network, and read out the class with the highest probability in the last (classification) layer.



What's the class of this object?

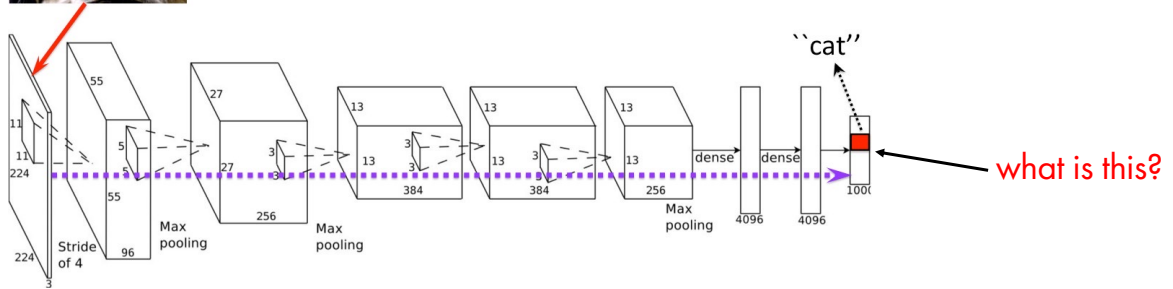


Classification

- Once trained we can do classification. Just feed in an image or a crop of the image, run through the network, and read out the class with the highest probability in the last (classification) layer.



What's the class of this object?

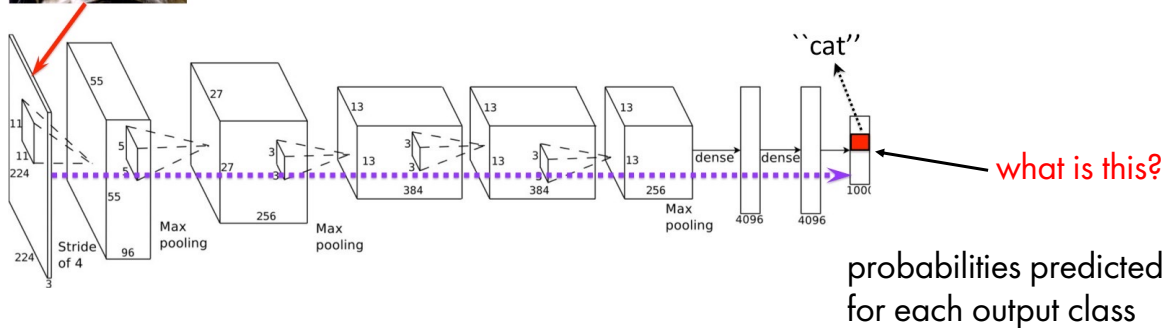


Classification

- Once trained we can do classification. Just feed in an image or a crop of the image, run through the network, and read out the class with the highest probability in the last (classification) layer.



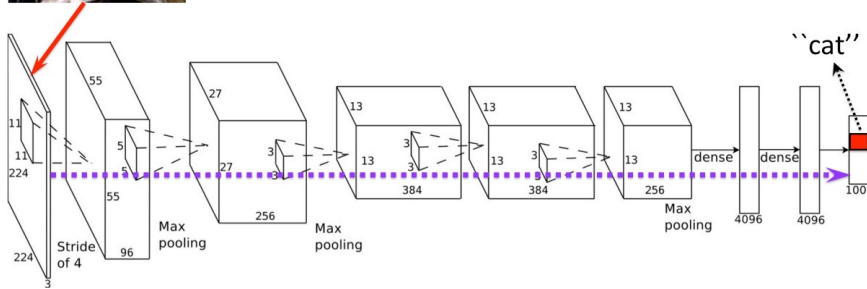
What's the class of this object?



Classification



What's the class of this object?



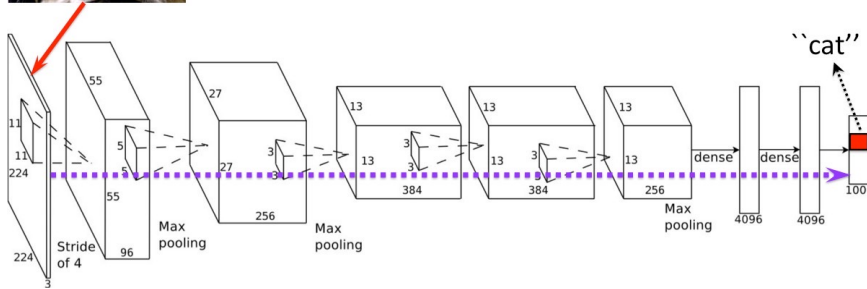
loss function

$$L_i = -\log\left(\frac{e^{f_{y_i}}}{\sum_j e^{f_j}}\right)$$

Classification



What's the class of this object?



loss function

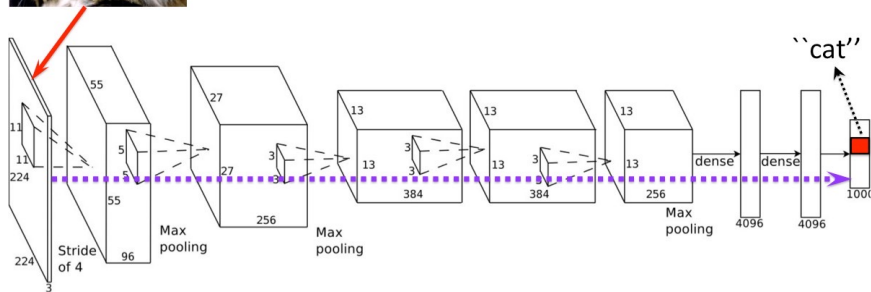
$$L_i = -\log \left(\frac{e^{f_{y_i}}}{\sum_j e^{f_j}} \right)$$

softmax function (squashes the values to be probability distribution)

Classification



What's the class of this object?



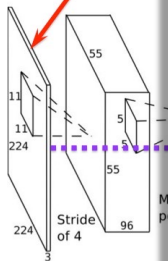
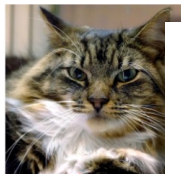
loss function

$$L_i = -\log \left(\frac{e^{f_{y_i}}}{\sum_j e^{f_j}} \right)$$

output for true class "slot"

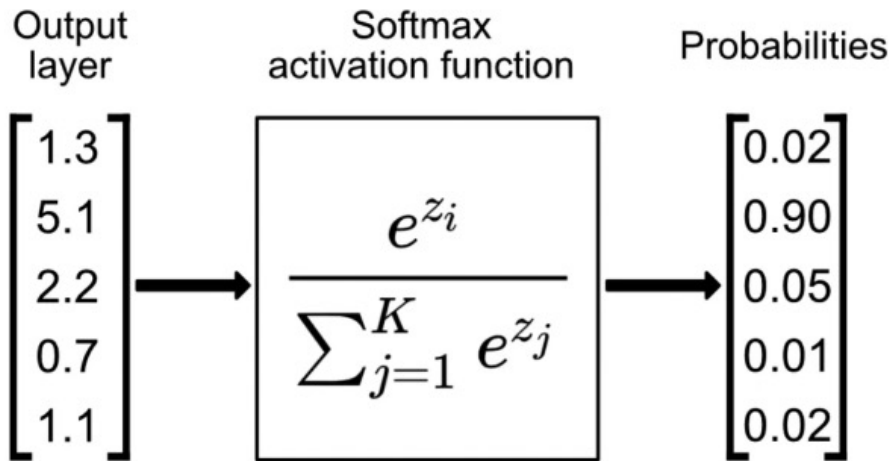
softmax function (squashes the values to be probability distribution)

Classification



loss

$$L_i = -1$$



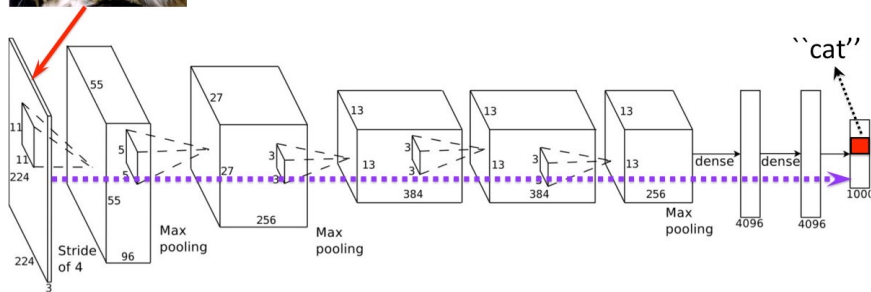
$$\left(\sum_j e^{z_j} \right)$$

softmax function (squashes the values to be probability distribution)

Classification



What's the class of this object?



loss function

$$L_i = -\log \left(\frac{e^{f_{y_i}}}{\sum_j e^{f_j}} \right)$$

output for true class "slot"

output for other classes

softmax function (squashes the values to be probability distribution)

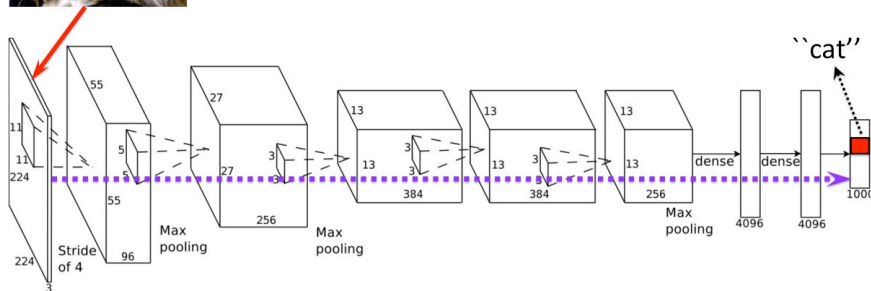
cross entropy!

$$H(p, q) = - \sum_x p(x) \log q(x)$$

Classification



What's the class of this object?



what is $p(x)$?

loss function

$$L_i = -\log \left(\frac{e^{f_{y_i}}}{\sum_j e^{f_j}} \right)$$

output for true class "slot"

output for other classes

softmax function (squashes the values to be probability distribution)

cross entropy!

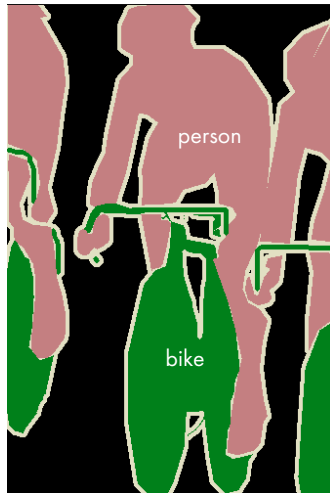
$$H(p, q) = - \sum_x p(x) \log q(x)$$

Outline

- motivation / overview of recognition tasks
- CNNs
 - Classification
 - **Segmentation**
 - Object Detection

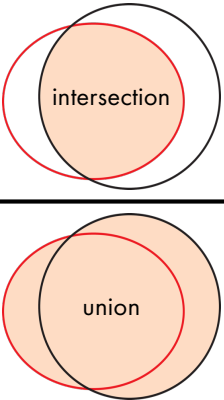
Semantic Segmentation: Problem Definition

- **Task:** Label every pixel
 - Don't differentiate instances



Semantic Segmentation: **Evaluation**

- Intersection over Union (IoU), averaged over classes

$$\text{IoU} = \frac{\text{intersection}}{\text{union}}$$




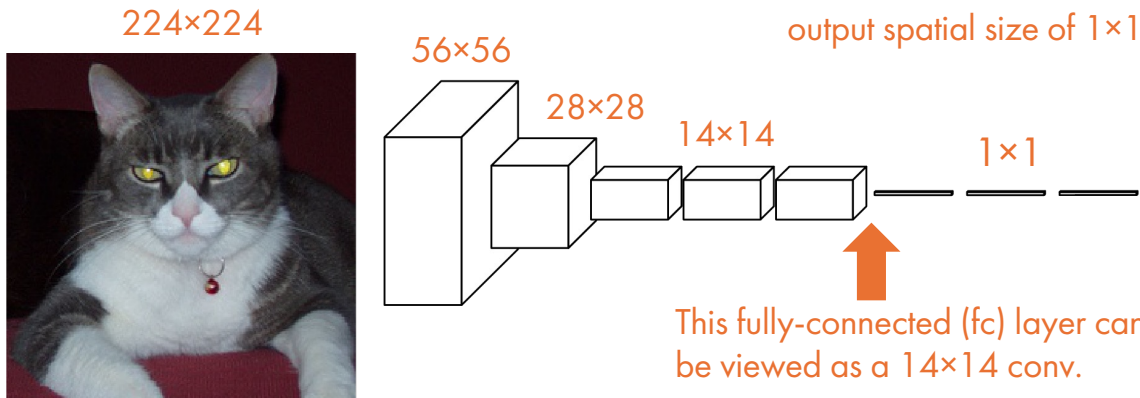
prediction



ground truth

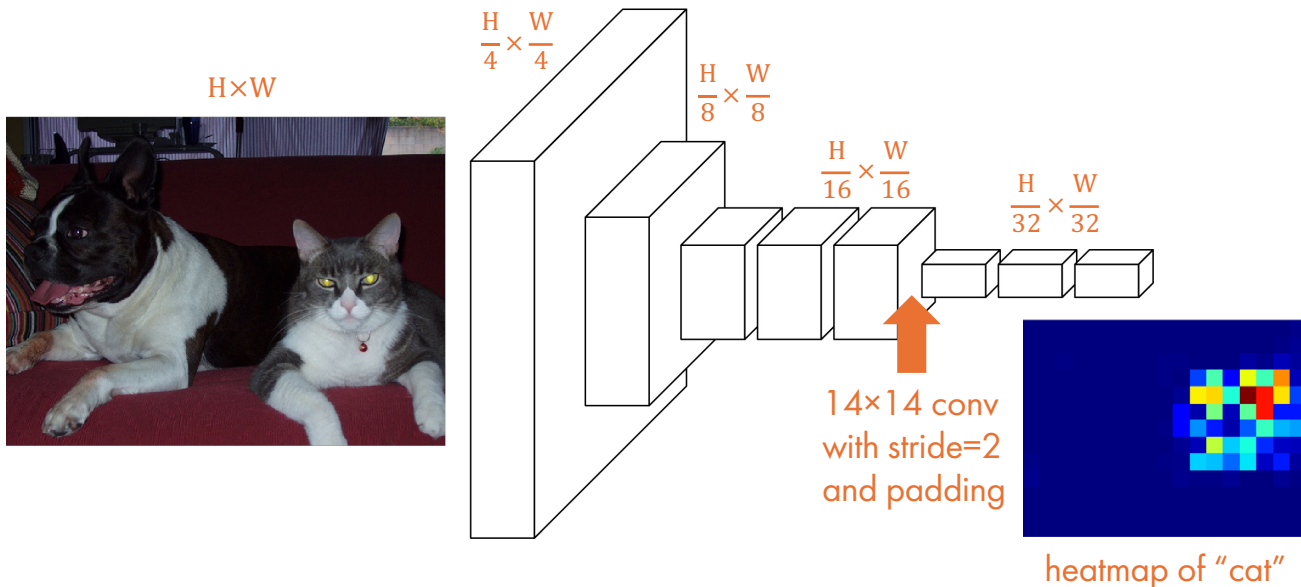
Fully Convolutional Networks (FCN)

- Repurpose a classification ConvNet for semantic segmentation



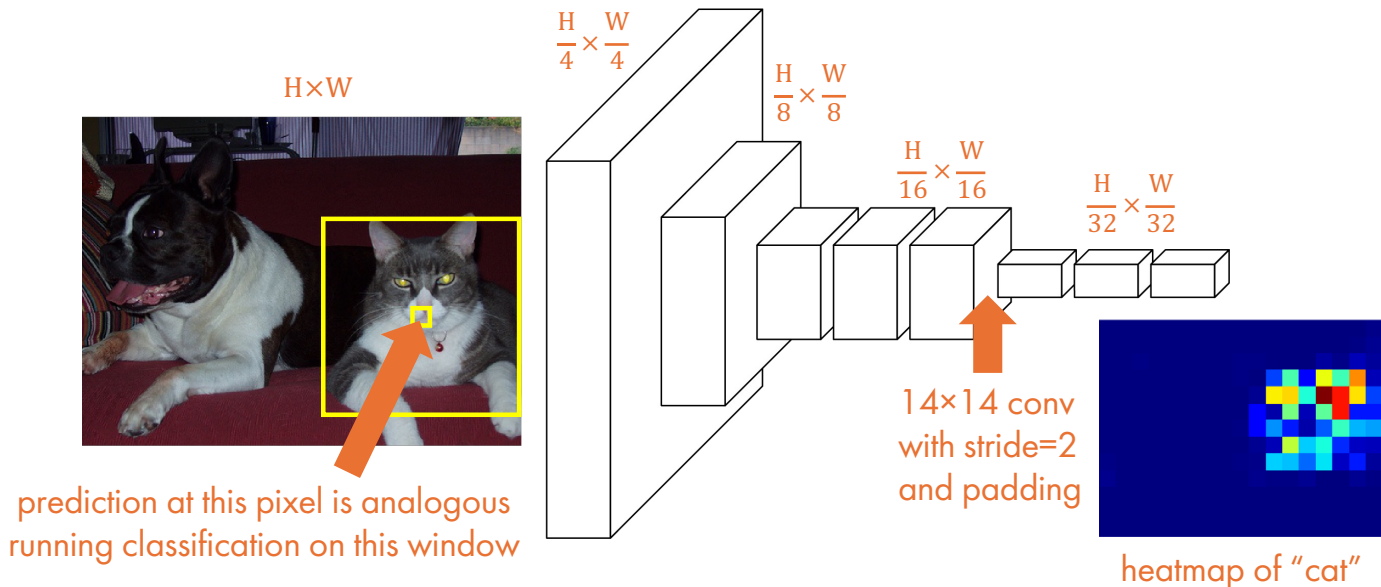
Fully Convolutional Networks (FCN)

- Repurpose a classification ConvNet for semantic segmentation



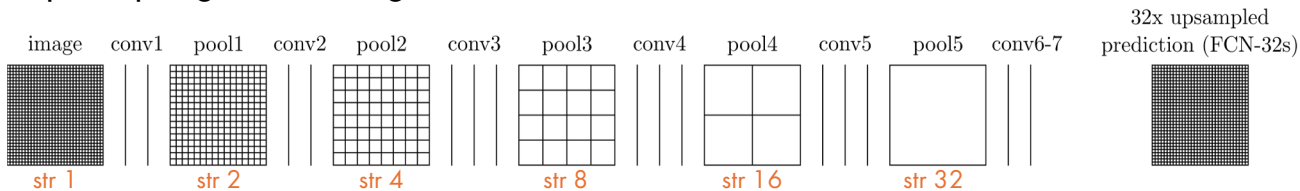
Fully Convolutional Networks (FCN)

- Repurpose a classification ConvNet for semantic segmentation



Fully Convolutional Networks (FCN)

- Upsampling: combining “where” and “what”

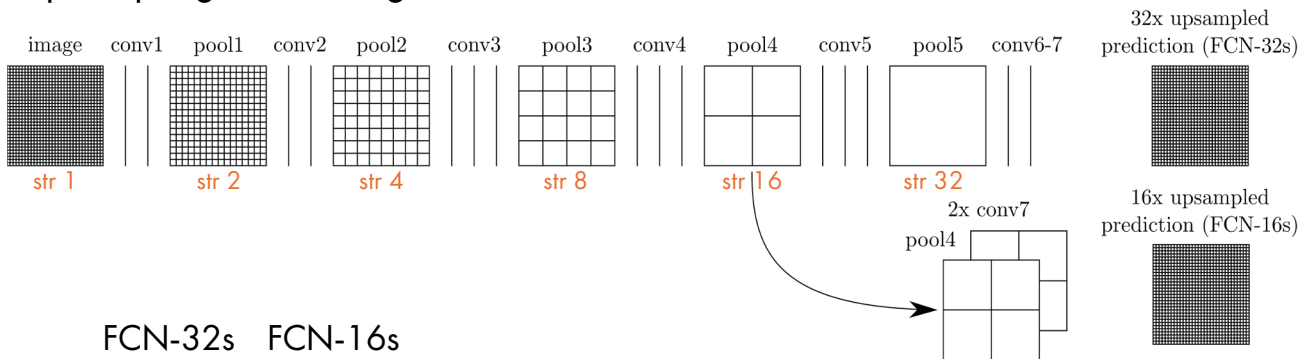


FCN-32s

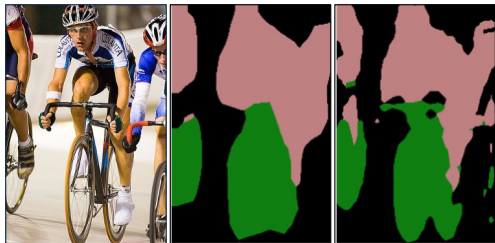


Fully Convolutional Networks (FCN)

- Upsampling: combining “where” and “what”

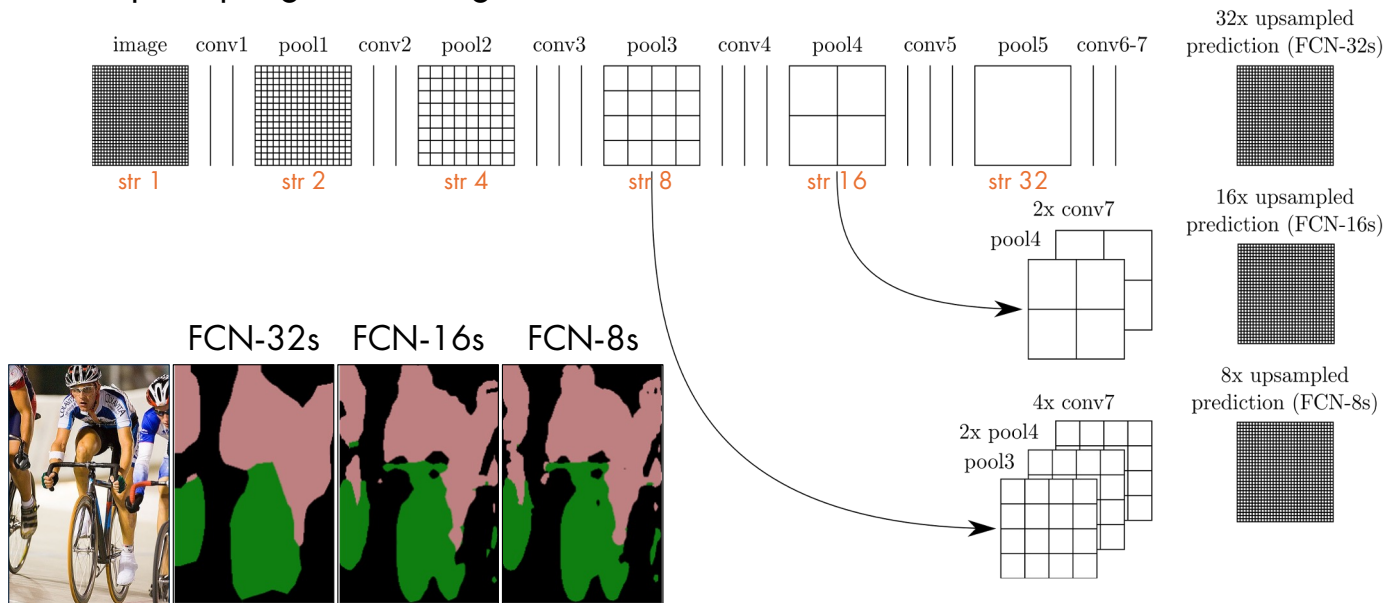


FCN-32s FCN-16s

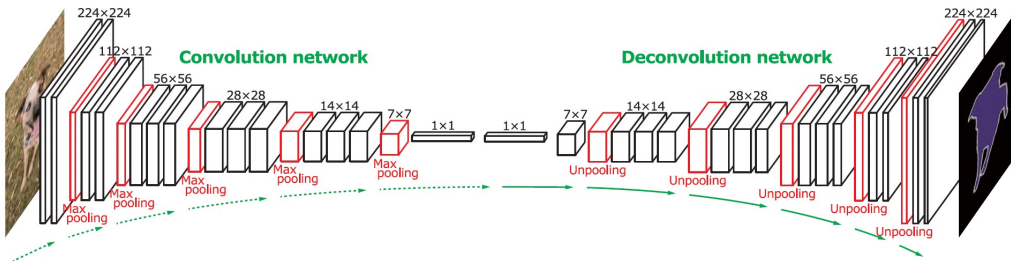
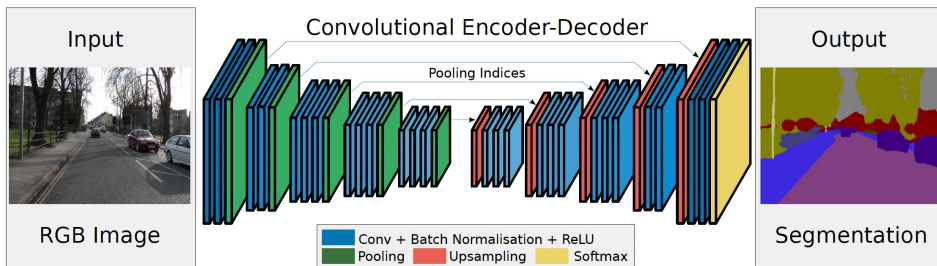


Fully Convolutional Networks (FCN)

- Upsampling: combining “where” and “what”



Again, an autoencoder!



Hyeonwoo Noh, Seunghoon Hong, Bohyung Han, "Learning Deconvolution Network for Semantic Segmentation", ICCV 2015

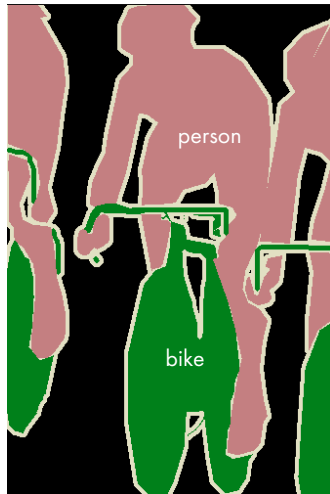
Vijay Badrinarayanan, Alex Kendall, Roberto Cipolla, "SegNet: A Deep Convolutional Encoder-Decoder Architecture for Image Segmentation", arXiv 2015, PAMI 2017

Semantic Segmentation

- Predict per-pixel class labels given an image
- Network structure: fully convolutional network (a CNN autoencoder)
- Repurpose a pre-trained image classification network for semantic segmentation
- Limitation: cannot differentiate instances

Semantic Segmentation Don't Differentiate Instances

- **Task:** Label every pixel
 - Don't differentiate instances

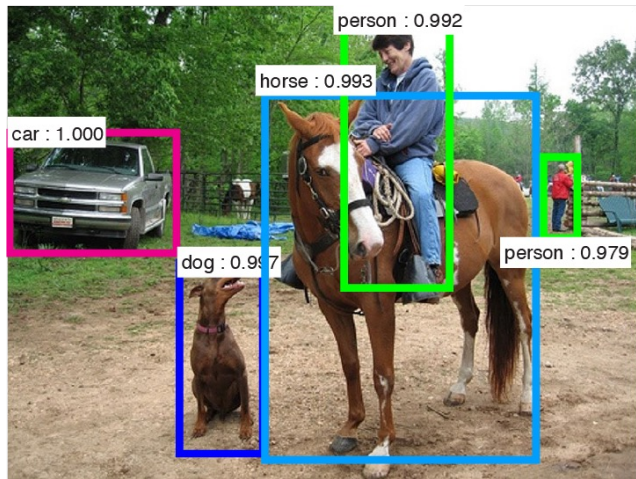


Outline

- motivation / overview of recognition tasks
- CNNs
 - Classification
 - Segmentation
 - Object Detection

Object Detection: Problem Definition

- **Task:** identify, locate, and recognize objects
- **Identify:**
 - Are there objects presented?
- **Locate:**
 - Where is each object?
 - (x, y, h, w)
- **Recognize:**
 - What is each object?
 - class label



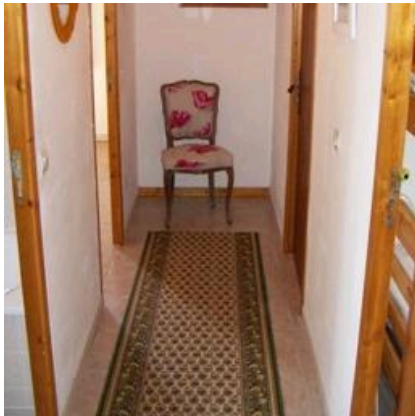
Object Recognition

- Object recognition is the task of **identifying** and **localizing** objects within images or video.

This is a chair



Find the chair in this image



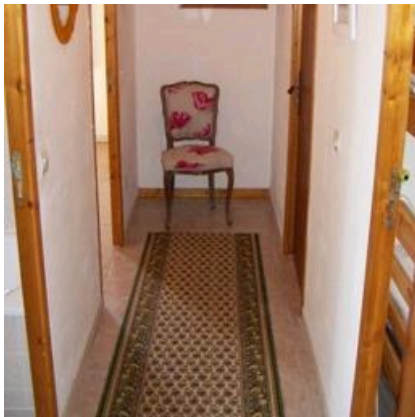
Object Recognition

- Object recognition is the task of **identifying** and **localizing** objects within images or video.

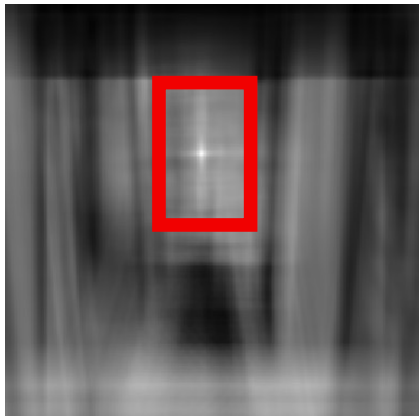
This is a chair



Find the chair in this image



Output of normalized correlation



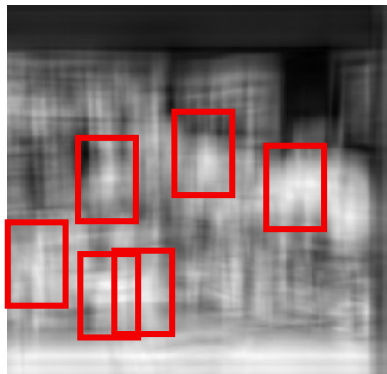
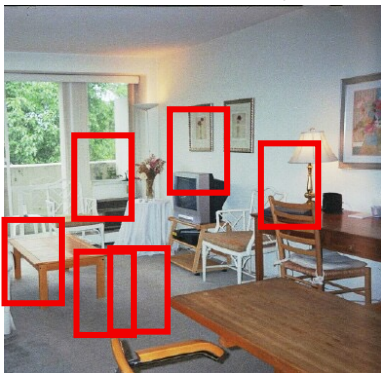
Object Recognition

- Object recognition is the task of **identifying** and **localizing** objects within images or video.

This is a chair



Find the chair in this image



Object Recognition

Find the chair in this image

This is a chair



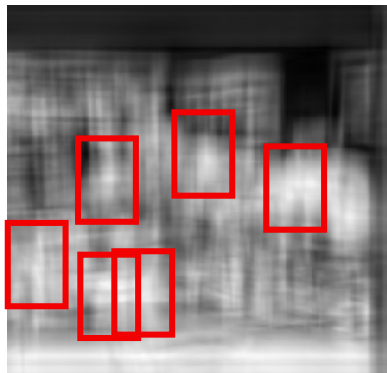
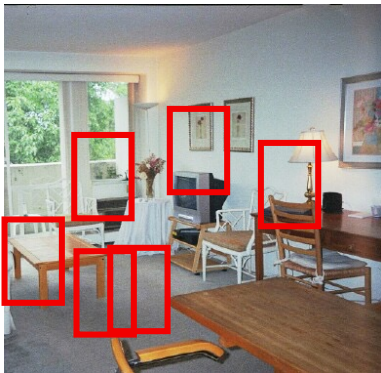
A “popular method is that of template matching, by point to point correlation of a model pattern with the image pattern. These techniques are inadequate for three-dimensional scene analysis for many reasons, such as occlusion, changes in viewing angle, and articulation of parts.” Nivatia & Binford, 1977.

Object Recognition

This is a chair



Find the chair in this image



A “popular method is that of template matching, by point to point correlation of a model pattern with the image pattern. These techniques are inadequate for three-dimensional scene analysis for many reasons, such as occlusion, changes in viewing angle, and articulation of parts.” Nivatia & Binford, 1977.

Object Detection: Evaluation



correct
(true positive)

Object Detection: Evaluation



correct
(true positive)



missed
(false negative)

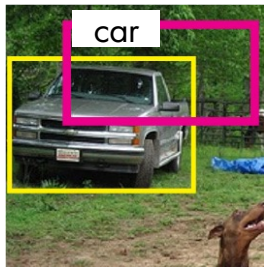
Object Detection: Evaluation



correct
(true positive)

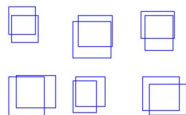


missed
(false negative)

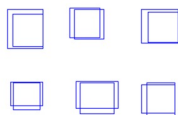


mislocalized
(false positive)

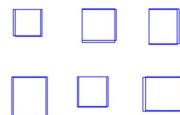
define by an IoU threshold



$\text{IoU} = 0.5$



$\text{IoU} = 0.7$



$\text{IoU} = 0.9$

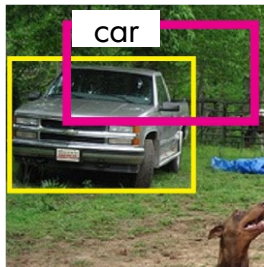
Object Detection: Evaluation



correct
(true positive)



missed
(false negative)

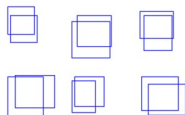


mislocalized
(false positive)

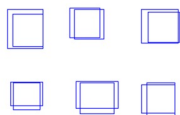


misrecognized
(false positive)

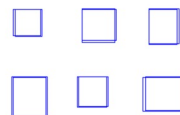
define by an IoU threshold



$\text{IoU} = 0.5$



$\text{IoU} = 0.7$



$\text{IoU} = 0.9$

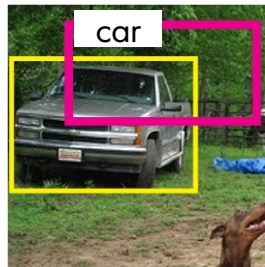
Object Detection: Evaluation



correct
(true positive)



missed
(false negative)



mislocalized
(false positive)

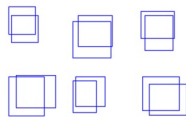


misrecognized
(false positive)

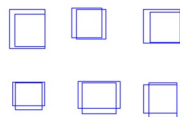


duplicated
(false positive)

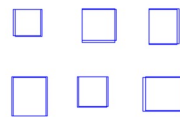
define by an IoU threshold



$\text{IoU} = 0.5$



$\text{IoU} = 0.7$



$\text{IoU} = 0.9$

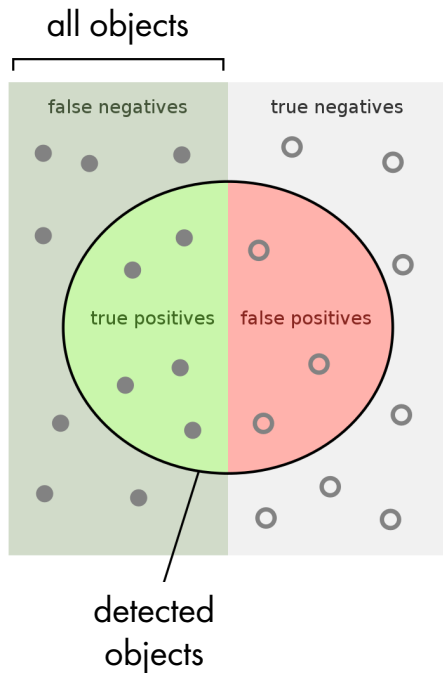
Object Detection: Evaluation

- Precision: What % of detected objects are correct?

Precision = $\frac{\text{true positive}}{\text{true positive} + \text{false positive}}$

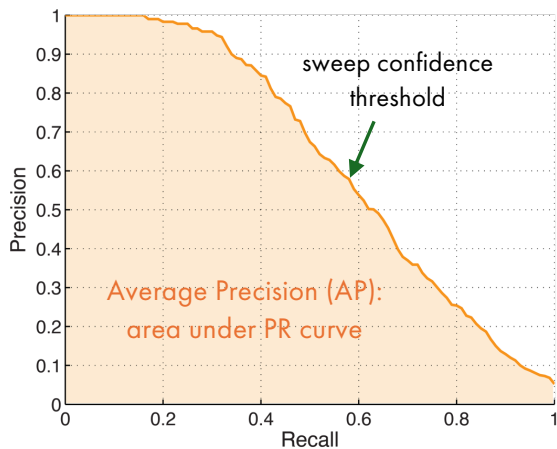
- Recall: What % of objects are detected?

$$\text{Recall} = \frac{\text{true positive}}{\text{true positive} + \text{false negative}}$$



Object Detection: Evaluation

- Precision-Recall Tradeoff

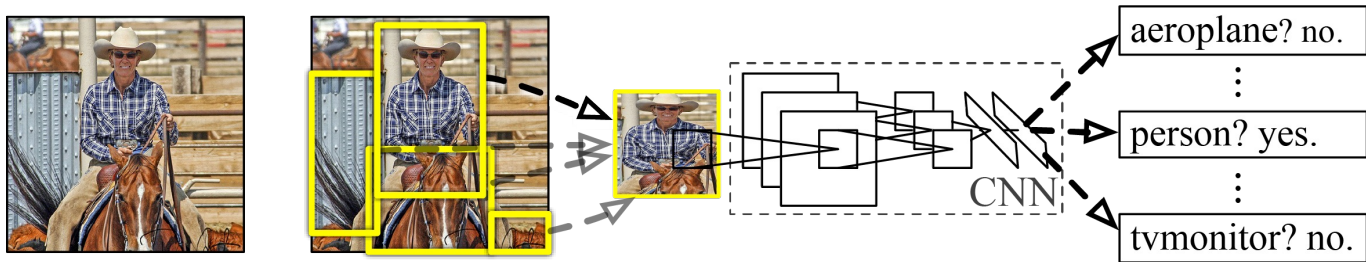


threshold = 0.4: Recall $\uparrow$, Precision $\downarrow$



threshold = 0.7: Recall $\downarrow$, Precision $\uparrow$

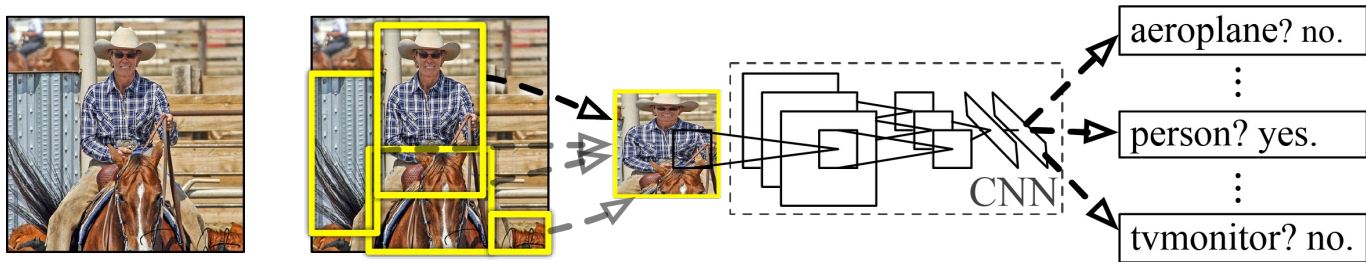
Region-Based CNN (R-CNN)



- THE deep learning revolution in the vision community.

the ILSVRC workshop at ECCV 2012.¹ The central issue of debate can be distilled as: **To what extent do the CNN *classification* results on ImageNet generalize to object *detection* results on the PASCAL VOC Challenge?** In this paper, we study object detection using features computed by a large convolutional neural network, answering this important scientific question.²

Region-Based CNN (R-CNN)



- THE deep learning revolution in the vision community.
- R-CNN: Regions with CNN features
- 40% relative improvement over classical methods
- Deep Learning made work and transferable beyond image classification

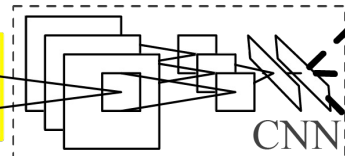
Region-Based CNN (R-CNN)



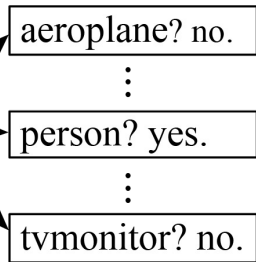
input



extract
"region proposals"

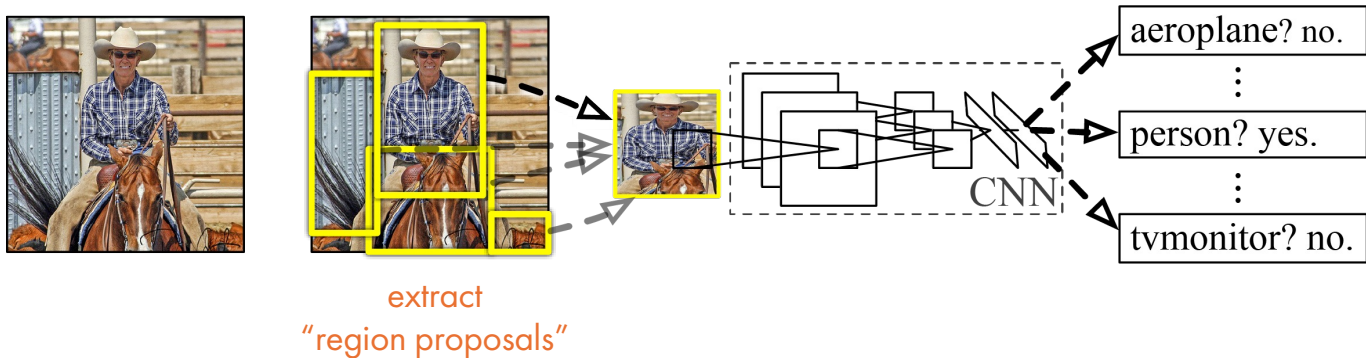


compute CNN features
for each region



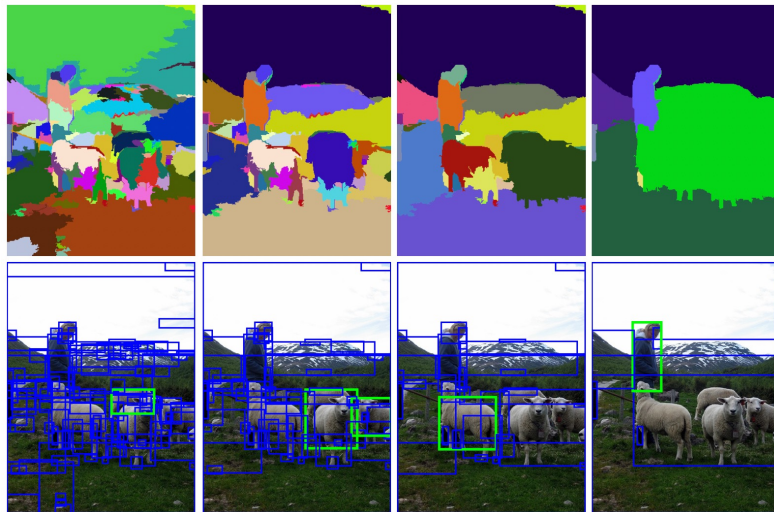
classify regions
& refine boxes

Region-Based CNN (R-CNN)

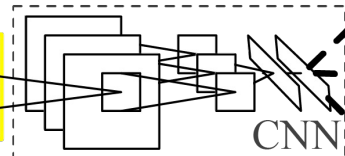
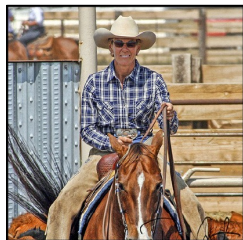


- Reduce the search space ("selective search")
- ~2000 regions / image
- Recall-driven
 - a proposal may not be an object
 - an object is highly likely in the proposals

e.g.: Selective Search for Region Proposal

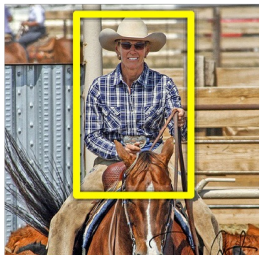


Region-Based CNN (R-CNN)



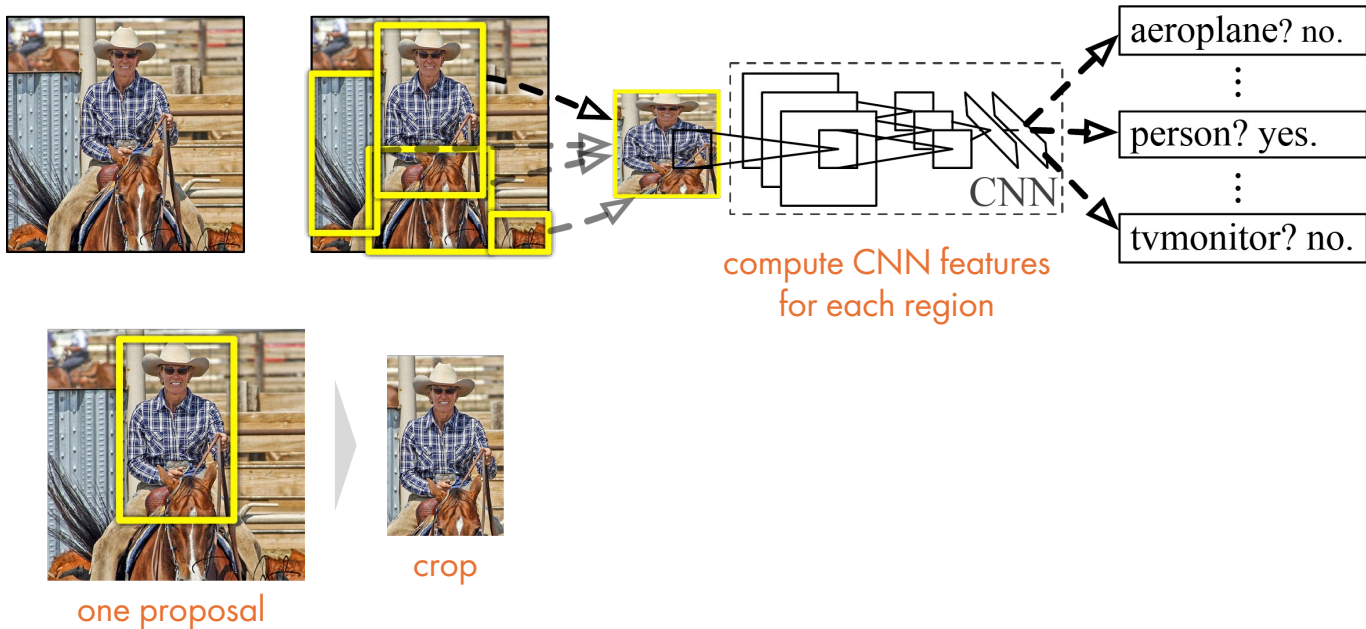
compute CNN features
for each region

aeroplane? no.
⋮
person? yes.
⋮
tvmonitor? no.

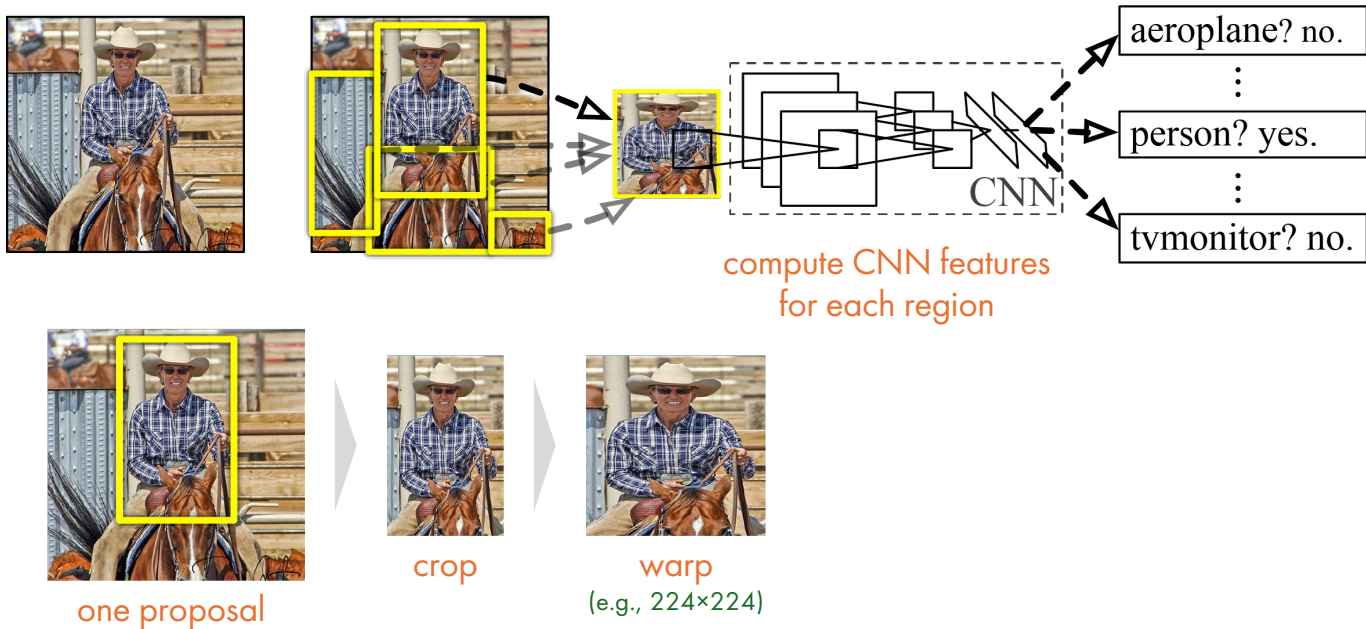


one proposal

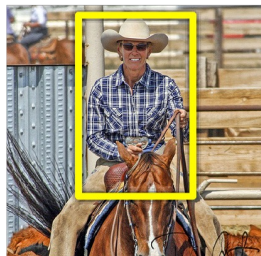
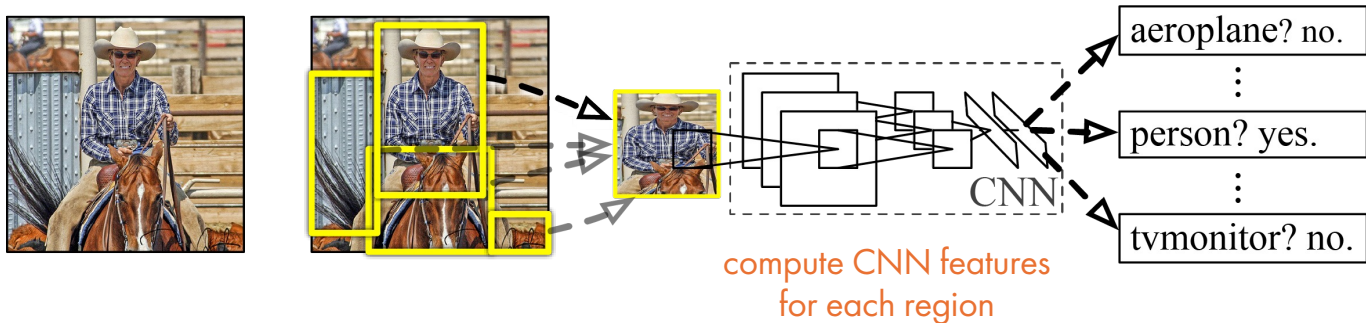
Region-Based CNN (R-CNN)



Region-Based CNN (R-CNN)



Region-Based CNN (R-CNN)



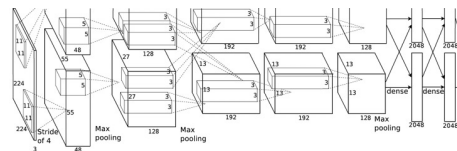
one proposal



crop

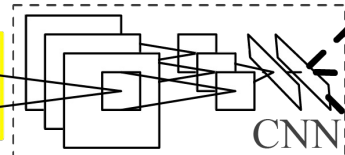
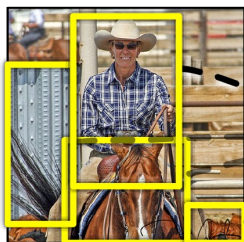
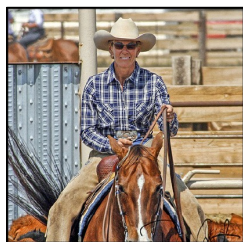


warp
(e.g., 224×224)



apply a classification ConvNet
(e.g., AlexNet)

Region-Based CNN (R-CNN)



aeroplane? no.

⋮

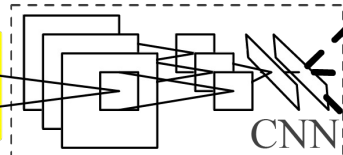
person? yes.

⋮

tvmonitor? no.

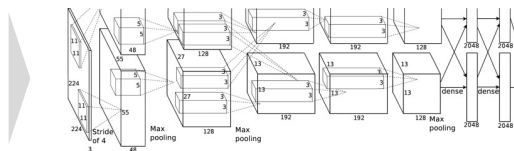
classify regions

Region-Based CNN (R-CNN)



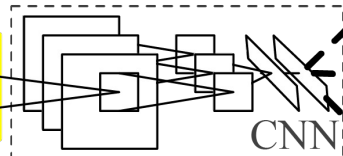
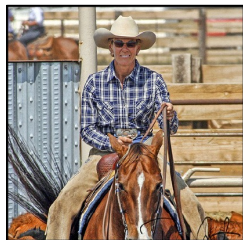
aeroplane? no.
:
person? yes.
:
tvmonitor? no.

classify regions



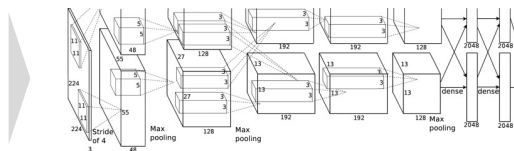
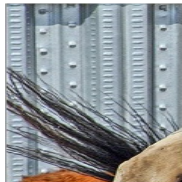
background: 0.0
:
person: 0.9
:
horse: 0.1
:

Region-Based CNN (R-CNN)



aeroplane? no.
:
person? yes.
:
tvmonitor? no.

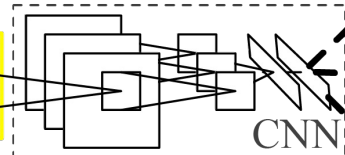
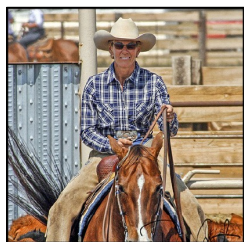
classify regions



Detection is first of all
a foreground/background
classification problem.

background: 0.9
:
person: 0.0
:
horse: 0.1
:

Region-Based CNN (R-CNN)

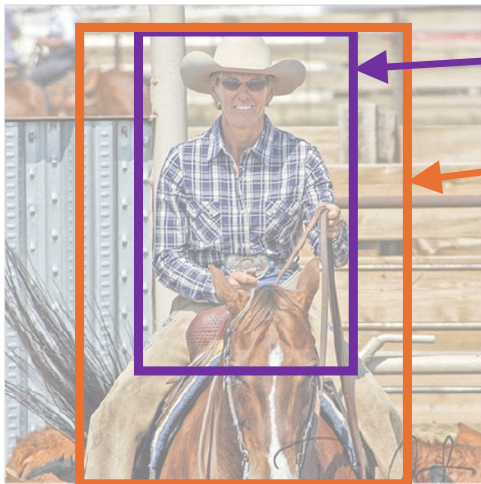


aeroplane? no.
⋮
person? yes.
⋮
tvmonitor? no.

classify regions
& refine boxes

Bounding-box Regression

- Geometric prediction from deep neural networks



reference box: (x, y, w, h)

target box: (x', y', w', h')

$$\Delta x = (x' - x)/w$$

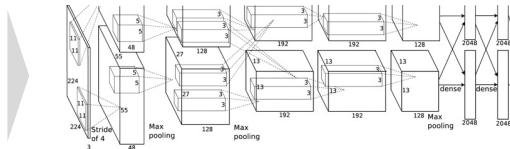
$$\Delta y = (y' - y)/h$$

$$\Delta \log w = \log w' - \log w$$

$$\Delta \log h = \log h' - \log h$$

Bounding-box Regression

- Geometric prediction from deep neural networks
- one network, multiple tasks



classification

bbox regression
 $\Delta x, \Delta y, \Delta \log w, \Delta \log h$

R-CNN



R-CNN

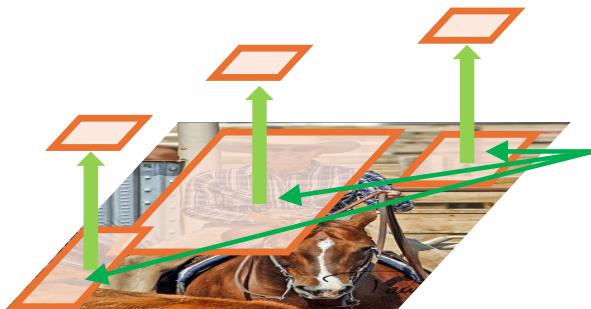


- Propose ~2000 Regions of Interest (RoI)

Slides adapted from: Ross Girshick, ICCV 2015

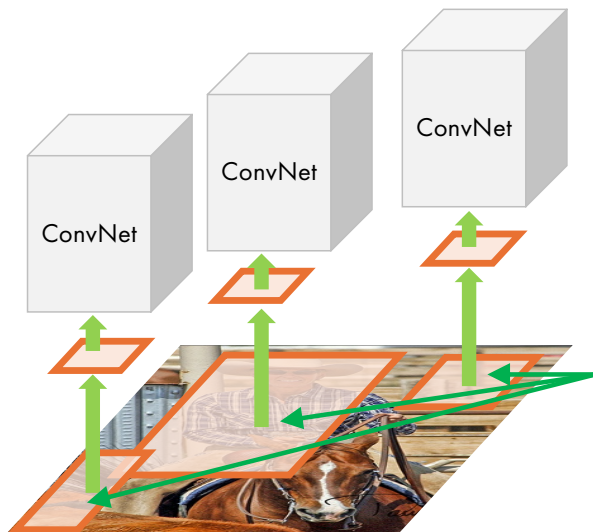
Ross Girshick, Jeff Donahue, Trevor Darrell, Jitendra Malik, "Rich Feature Hierarchies for Accurate Object Detection and Semantic Segmentation", CVPR 2014

R-CNN



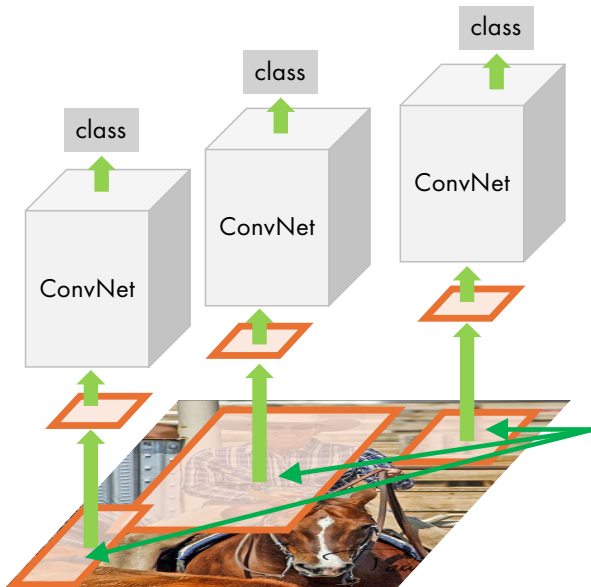
- Warp regions
- Propose ~2000 Regions of Interest (RoI)

R-CNN



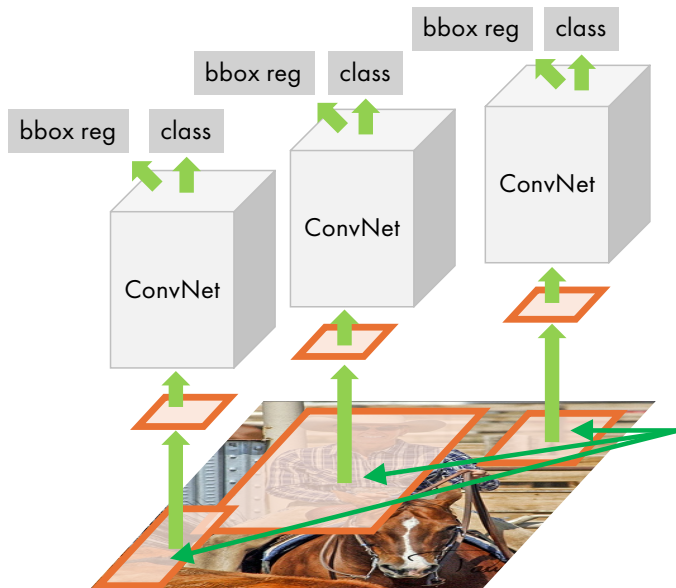
- Apply ConvNets to regions
- Warp regions
- Propose ~ 2000 Regions of Interest (RoI)

R-CNN



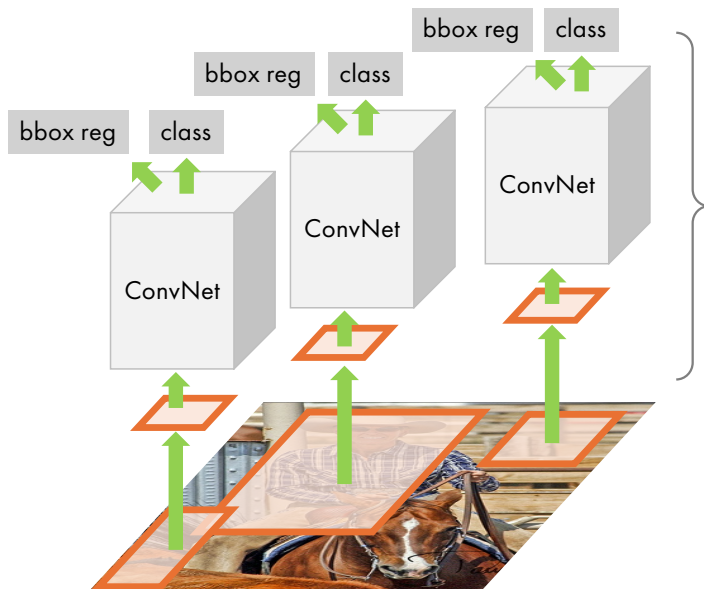
- Classify regions
- Apply ConvNets to regions
- Warp regions
- Propose ~2000 Regions of Interest (RoI)

R-CNN



- Regress bounding-boxes
- Classify regions
- Apply ConvNets to regions
- Warp regions
- Propose ~2000 Regions of Interest (RoI)

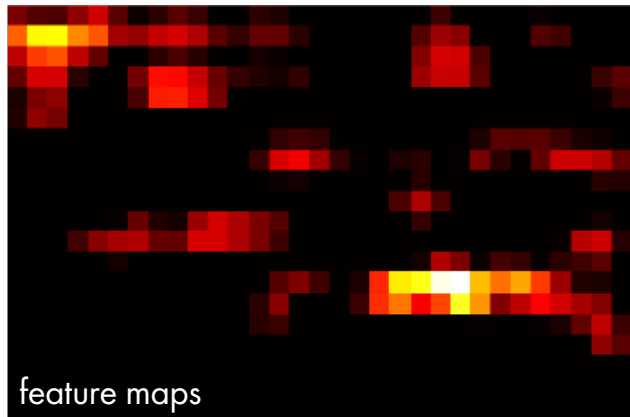
R-CNN is “Slow” R-CNN (quote [Ross Girshick, talk at ICCV 2015])



What's wrong w/ Slow R-CNN?
Many (~2000) ConvNets
applied on regions!

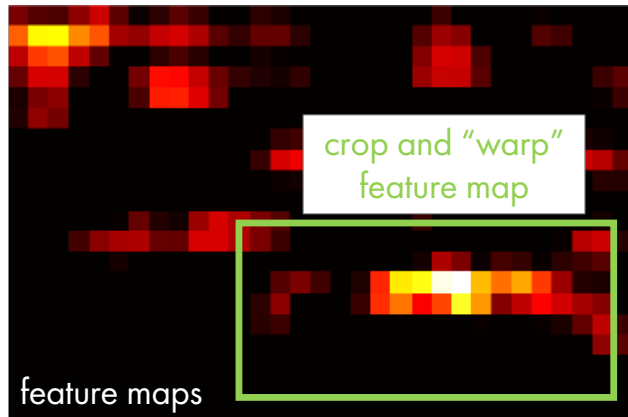
Towards Fast R-CNN (SPPnet)

- ConvNets are Fully Convolutional!



Towards Fast R-CNN (SPPnet)

- ConvNets are Fully Convolutional!



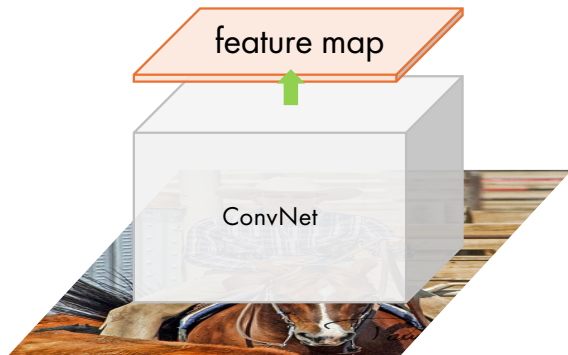
Towards Fast R-CNN (SPPnet)



Slides adapted from: Ross Girshick, ICCV 2015

Kaiming He, Xiangyu Zhang, Shaoqing Ren, Jian Sun, "Spatial Pyramid Pooling in Deep Convolutional Networks for Visual Recognition", ECCV 2014

Towards Fast R-CNN (SPPnet)

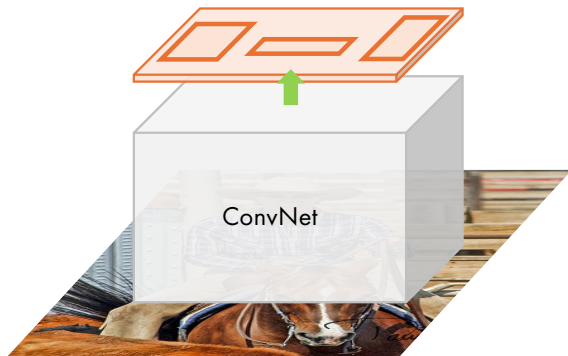


- Apply ConvNet to whole image

Slides adapted from: Ross Girshick, ICCV 2015

Kaiming He, Xiangyu Zhang, Shaoqing Ren, Jian Sun, "Spatial Pyramid Pooling in Deep Convolutional Networks for Visual Recognition", ECCV 2014

Towards Fast R-CNN (SPPnet)

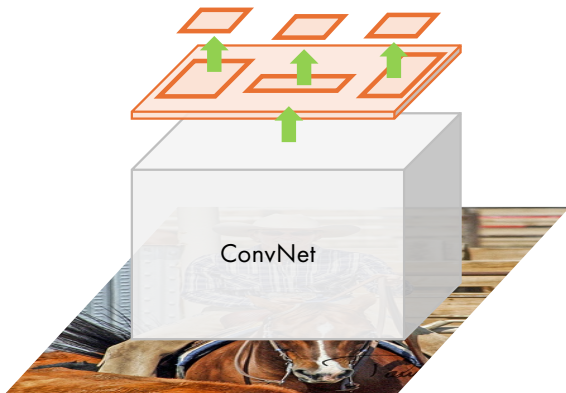


- Map proposed RoIs to feature maps
- Apply ConvNet to whole image

Slides adapted from: Ross Girshick, ICCV 2015

Kaiming He, Xiangyu Zhang, Shaoqing Ren, Jian Sun, "Spatial Pyramid Pooling in Deep Convolutional Networks for Visual Recognition", ECCV 2014

Towards Fast R-CNN (SPPnet)



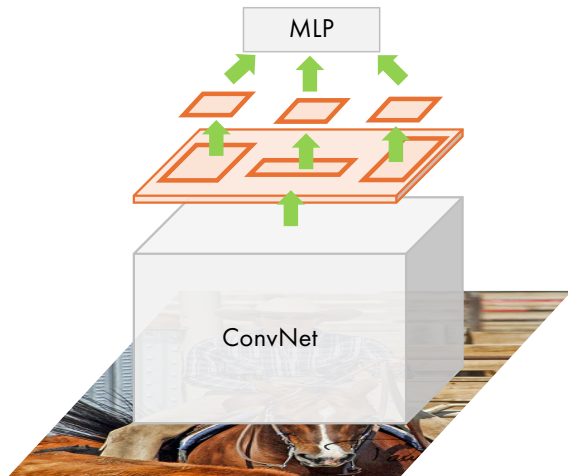
- Crop and pool features
- Map proposed RoIs to feature maps
- Apply ConvNet to whole image

* In original SPPnet, pooling is done with a fancier variant (spatial pyramid pooling)

Slides adapted from: Ross Girshick, ICCV 2015

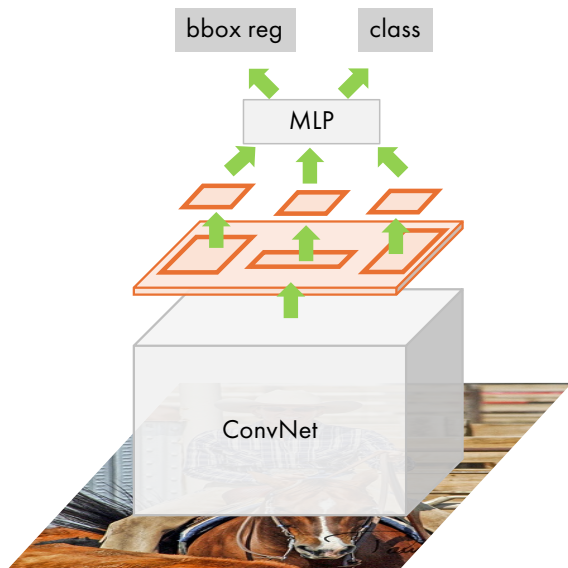
Kaiming He, Xiangyu Zhang, Shaoqing Ren, Jian Sun, "Spatial Pyramid Pooling in Deep Convolutional Networks for Visual Recognition", ECCV 2014

Towards Fast R-CNN (SPPnet)



- Apply MLP on regions
- Crop and pool features
- Map proposed RoIs to feature maps
- Apply ConvNet to whole image

Towards Fast R-CNN (SPPnet)



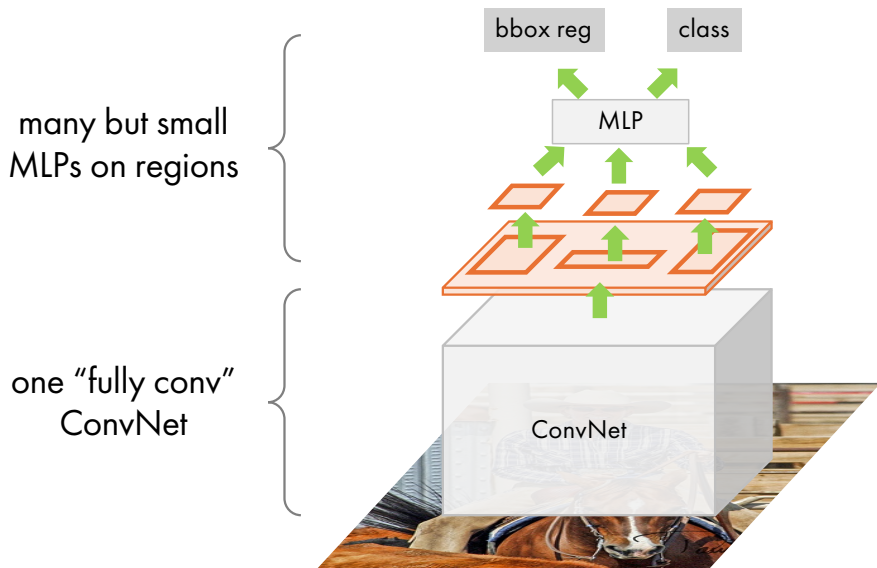
- Classify and regress boxes
- Apply MLP on regions
- Crop and pool features
- Map proposed Rols to feature maps
- Apply ConvNet to whole image

Slides adapted from: Ross Girshick, ICCV 2015

Kaiming He, Xiangyu Zhang, Shaoqing Ren, Jian Sun, "Spatial Pyramid Pooling in Deep Convolutional Networks for Visual Recognition", ECCV 2014

Towards Fast R-CNN (SPPnet)

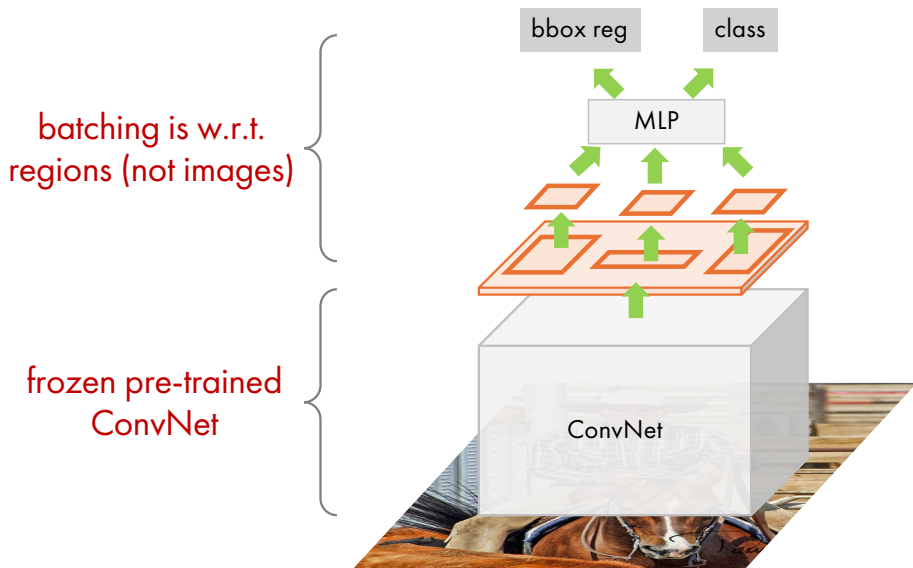
- What's good with "fully convolutional" ConvNet?



- Classify and regress boxes
- Apply MLP on regions
- Crop and pool features
- Map proposed Rols to feature maps
- Apply ConvNet to whole image

Towards Fast R-CNN (SPPnet)

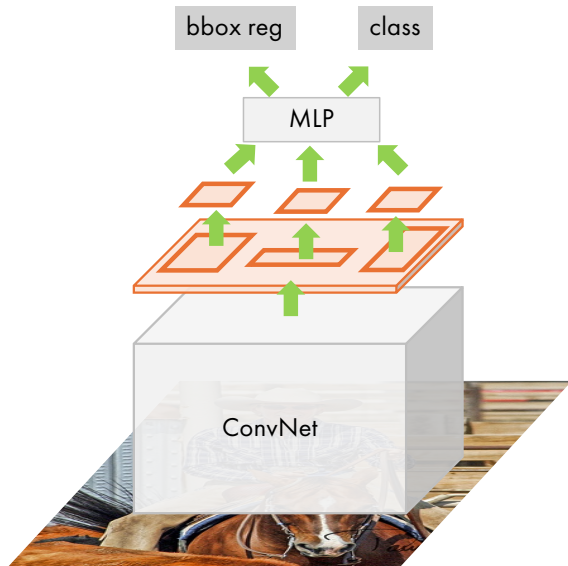
- What's still wrong? Not yet "Differentiable Programming"



- Classify and regress boxes
- Apply MLP on regions
- Crop and pool features
- Map proposed Rols to feature maps
- Apply ConvNet to whole image

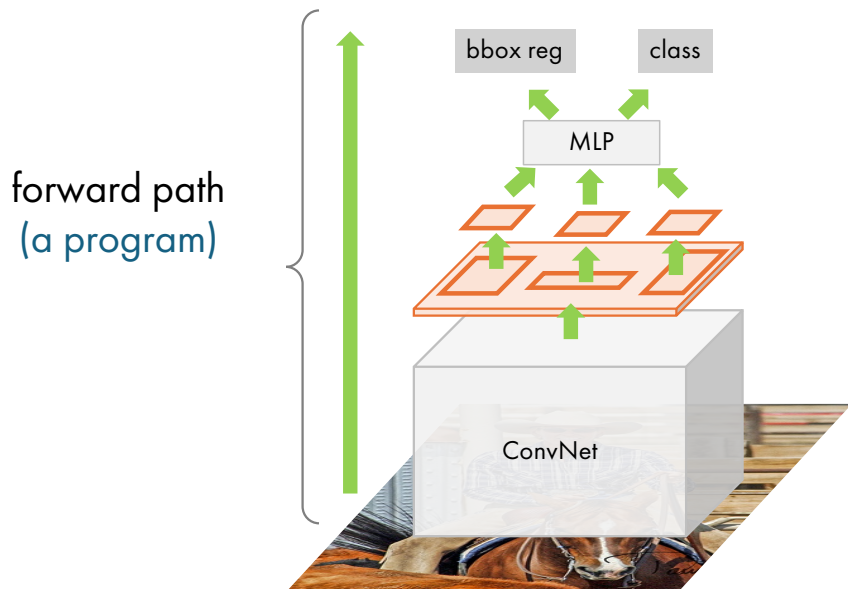
Fast R-CNN

- Differentiable Programming



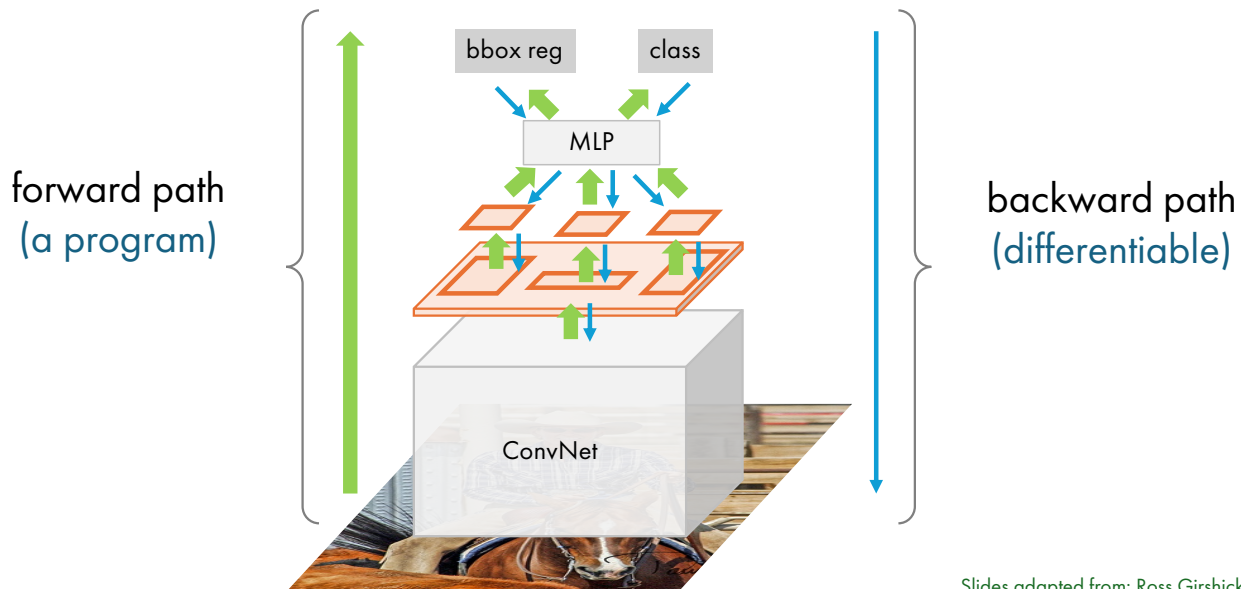
Fast R-CNN

- Differentiable Programming



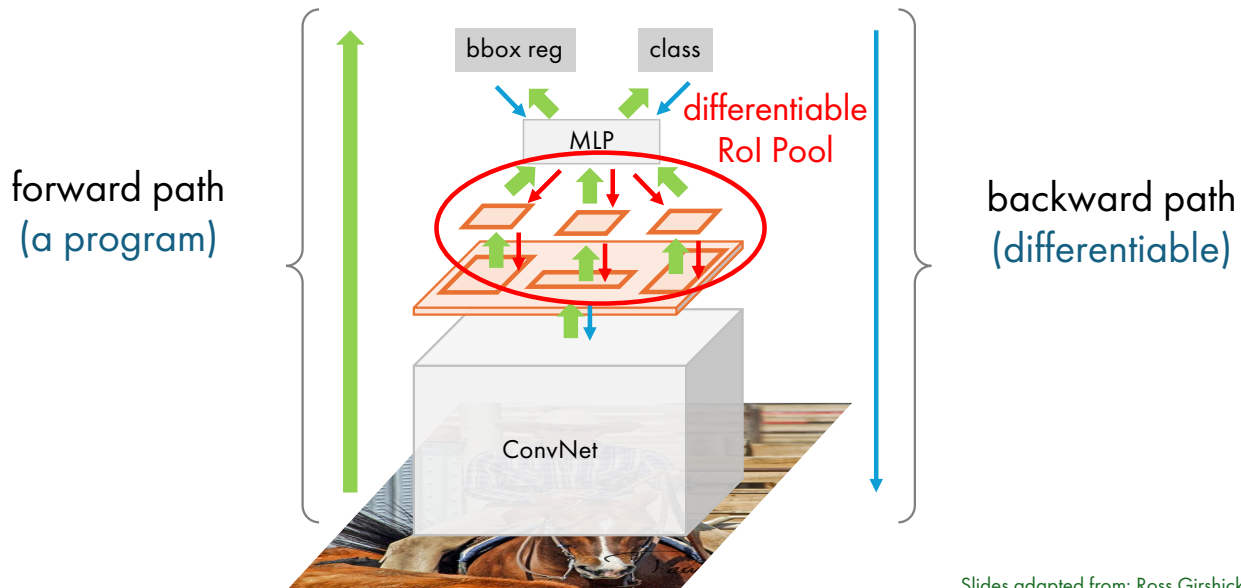
Fast R-CNN

- Differentiable Programing



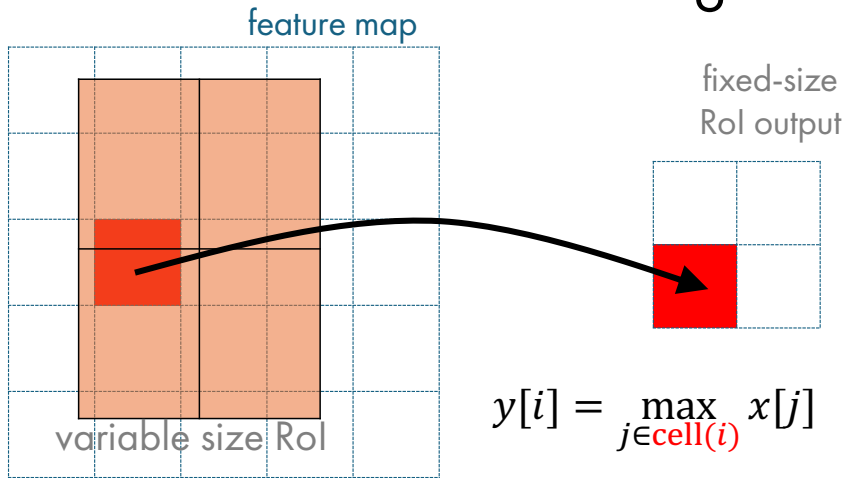
Fast R-CNN

- Differentiable Programing

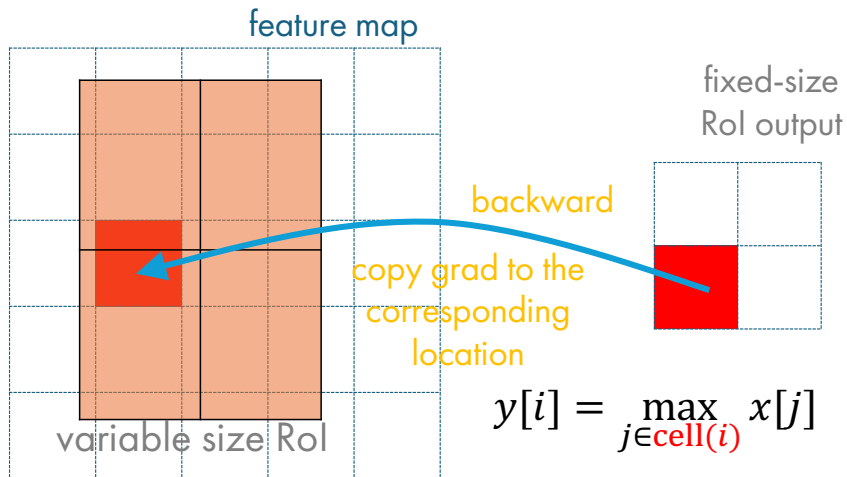


Fast R-CNN

Differentiable RoI Pooling

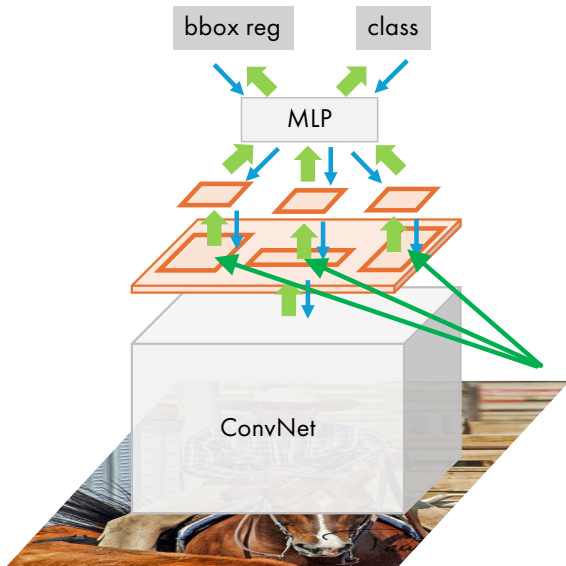


Differentiable RoI Pooling

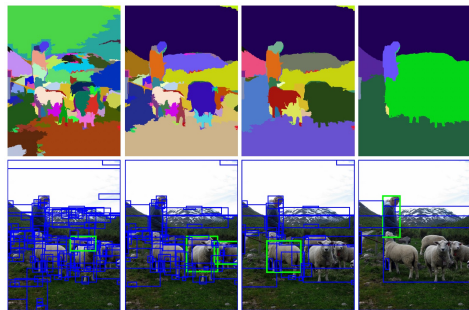


Fast R-CNN

- What's still wrong? Region proposals are not part of the program



recap: Selective Search



- Propose Regions of Interest (RoI)

Faster R-CNN

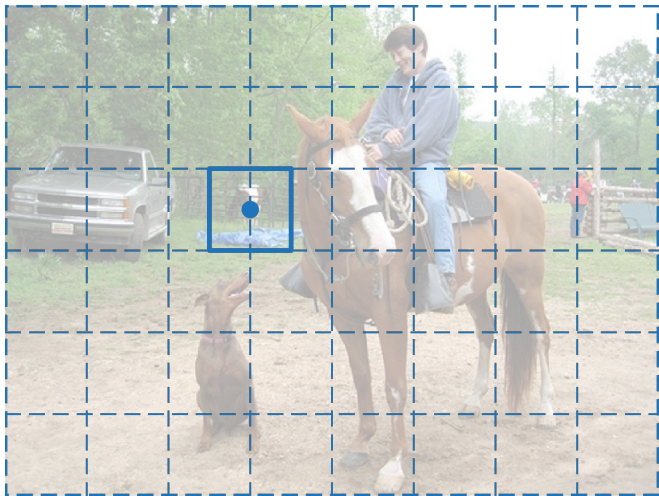
- How can neural nets produce regions?



- ConvNets are Fully Convolutional!
- Search objects w/ sliding windows (a.k.a., convolutions)

Faster R-CNN

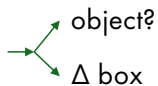
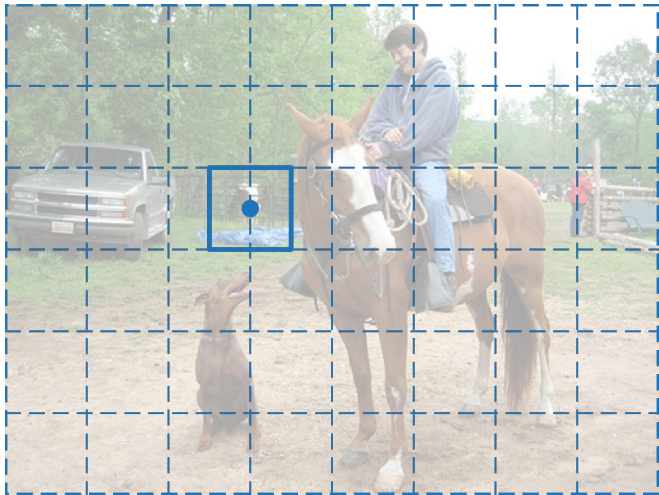
- How can neural nets produce regions?



- ConvNets are Fully Convolutional!
- Search objects w/ sliding windows (a.k.a., convolutions)

Faster R-CNN

- How can neural nets produce regions?

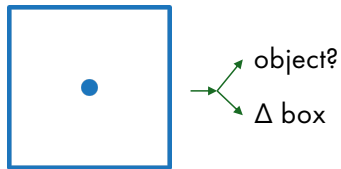
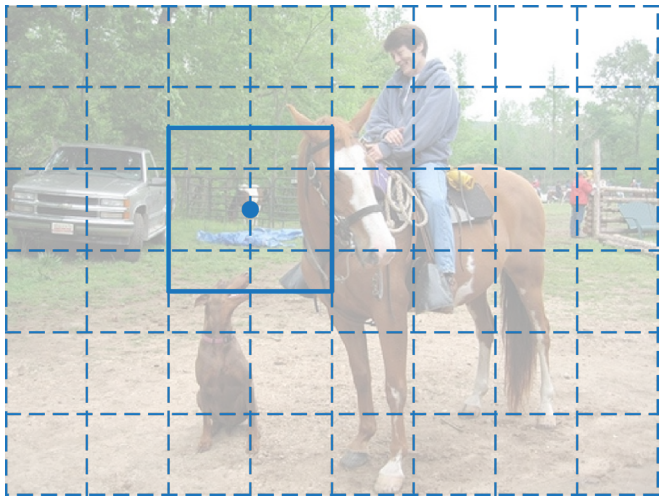


For each window, predict:

- Is it an object?
 - binary classification
- Object location w.r.t. window?
 - bbox regression

Faster R-CNN

- How can neural nets produce regions?

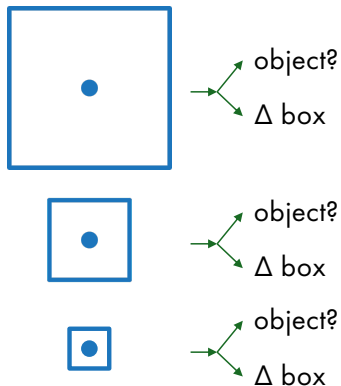
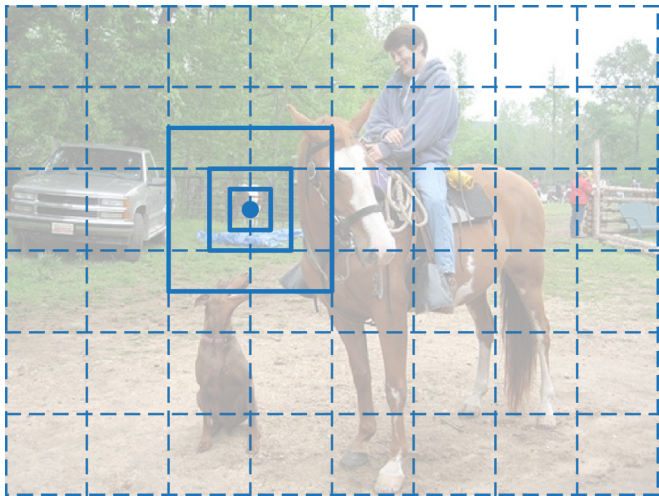


Anchor

- pred w.r.t. a reference window
- reference is called an "anchor"
- be different w/ sliding window

Faster R-CNN

- How can neural nets produce regions?

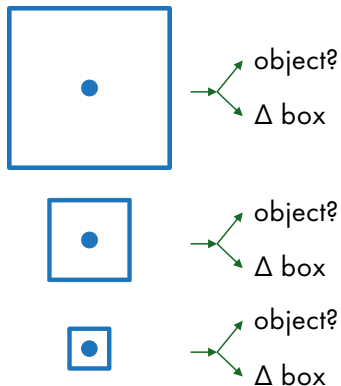
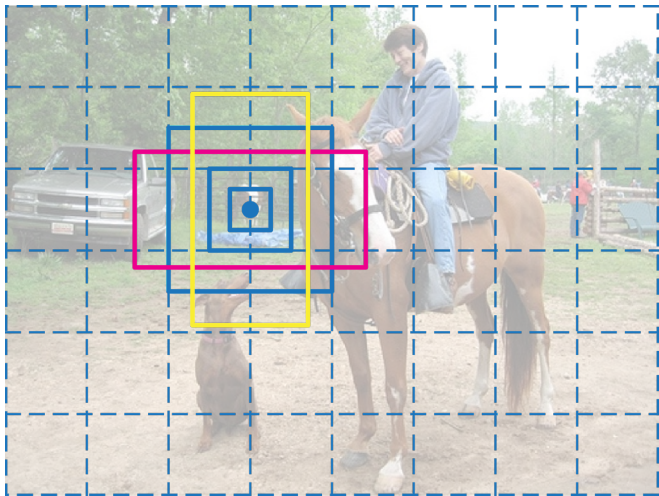


Anchors

- one sliding window, many anchors
- multiple scales

Faster R-CNN

- How can neural nets produce regions?

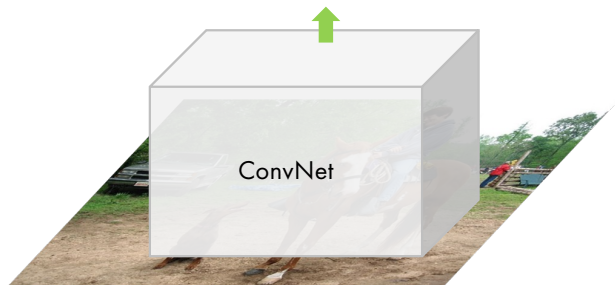


Anchors

- one sliding window, many anchors
- multiple scales
- multiple aspect ratios

Faster R-CNN: Region Proposal Network (RPN)

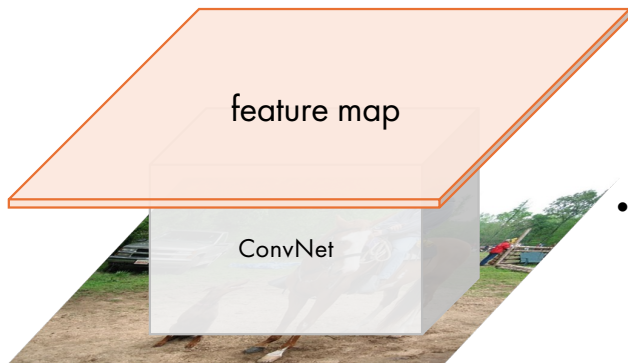
- How can neural nets produce regions?



- Apply ConvNet to whole image

Faster R-CNN: Region Proposal Network (RPN)

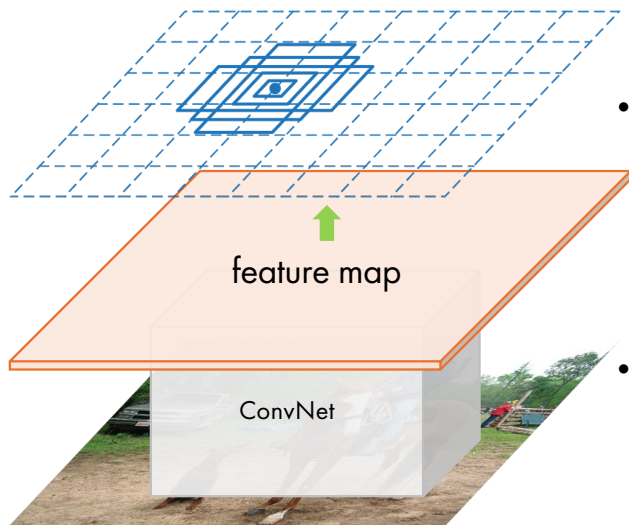
- How can neural nets produce regions?



- Apply ConvNet to whole image

Faster R-CNN: Region Proposal Network (RPN)

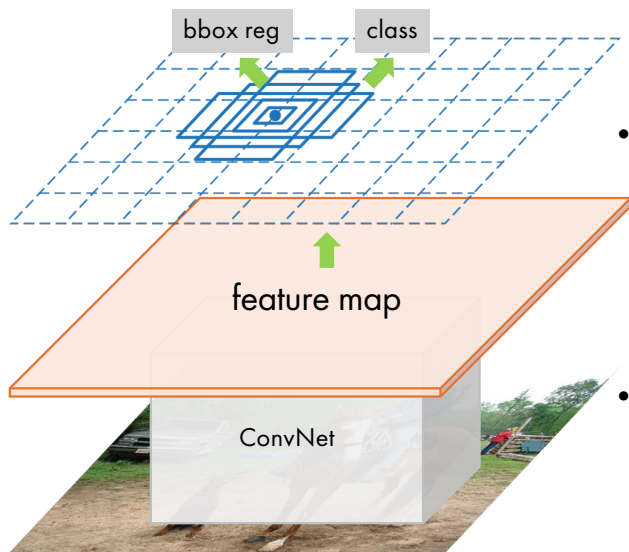
- How can neural nets produce regions?



- Convs for k anchors
 - k -d: binary classification
 - $4k$ -d: bbox regression
- Apply ConvNet to whole image

Faster R-CNN: Region Proposal Network (RPN)

- How can neural nets produce regions?



- Loss: reg + class
 - over all locations
 - over all anchors
- Convs for k anchors
 - k -d: binary classification
 - $4k$ -d: bbox regression
- Apply ConvNet to whole image

Faster R-CNN: End-to-End Object Detection

- End-to-End Differentiable Programming



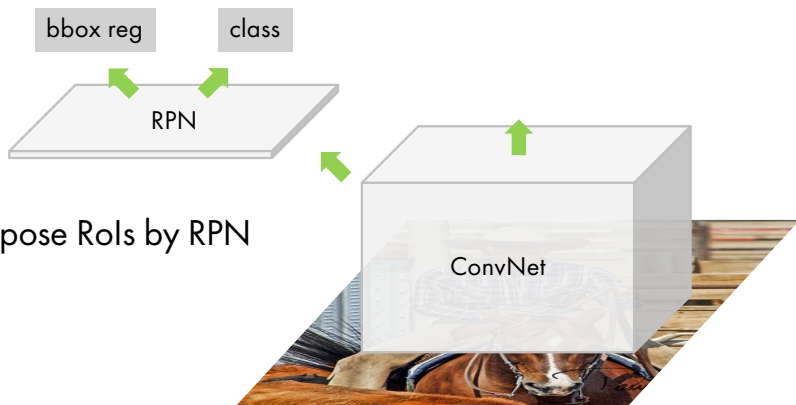
- Shared ConvNet backbone on whole image

Slides adapted from: Ross Girshick, ICCV 2015

Shaoqing Ren, Kaiming He, Ross Girshick, Jian Sun, "Faster R-CNN: Towards Real-Time Object Detection with Region Proposal Networks", NeurIPS 2015

Faster R-CNN: End-to-End Object Detection

- End-to-End Differentiable Programming



- Propose Rols by RPN

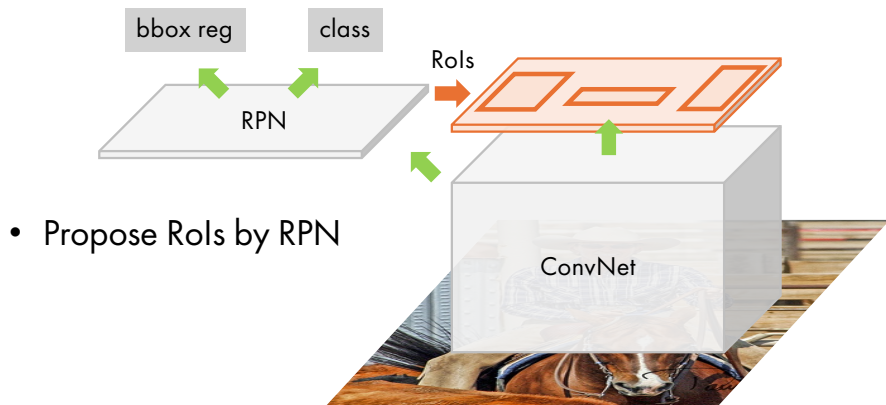
- Shared ConvNet backbone on whole image

Slides adapted from: Ross Girshick, ICCV 2015

Shaoqing Ren, Kaiming He, Ross Girshick, Jian Sun, "Faster R-CNN: Towards Real-Time Object Detection with Region Proposal Networks", NeurIPS 2015

Faster R-CNN: End-to-End Object Detection

- End-to-End Differentiable Programming



- Apply RPN's RoIs

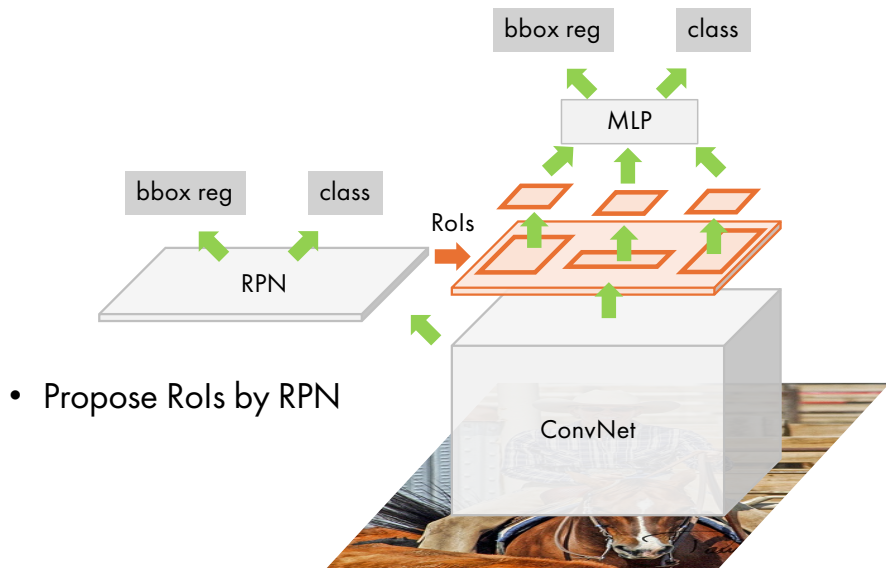
- Shared ConvNet backbone on whole image

Slides adapted from: Ross Girshick, ICCV 2015

Shaoqing Ren, Kaiming He, Ross Girshick, Jian Sun, "Faster R-CNN: Towards Real-Time Object Detection with Region Proposal Networks", NeurIPS 2015

Faster R-CNN: End-to-End Object Detection

- End-to-End Differentiable Programming



- Propose RoIs by RPN

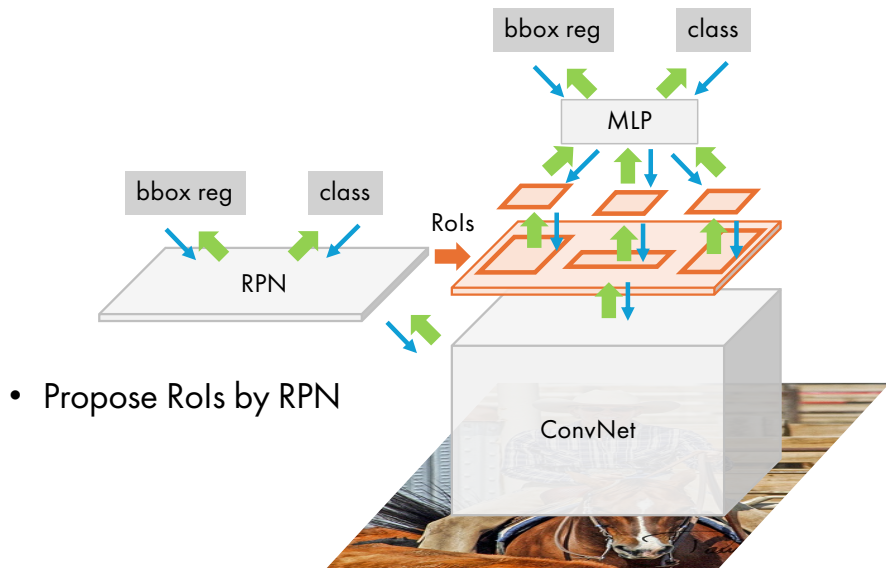
- Same as Fast R-CNN
- Apply RPN's RoIs
- Shared ConvNet backbone on whole image

Slides adapted from: Ross Girshick, ICCV 2015

Shaoqing Ren, Kaiming He, Ross Girshick, Jian Sun, "Faster R-CNN: Towards Real-Time Object Detection with Region Proposal Networks", NeurIPS 2015

Faster R-CNN: End-to-End Object Detection

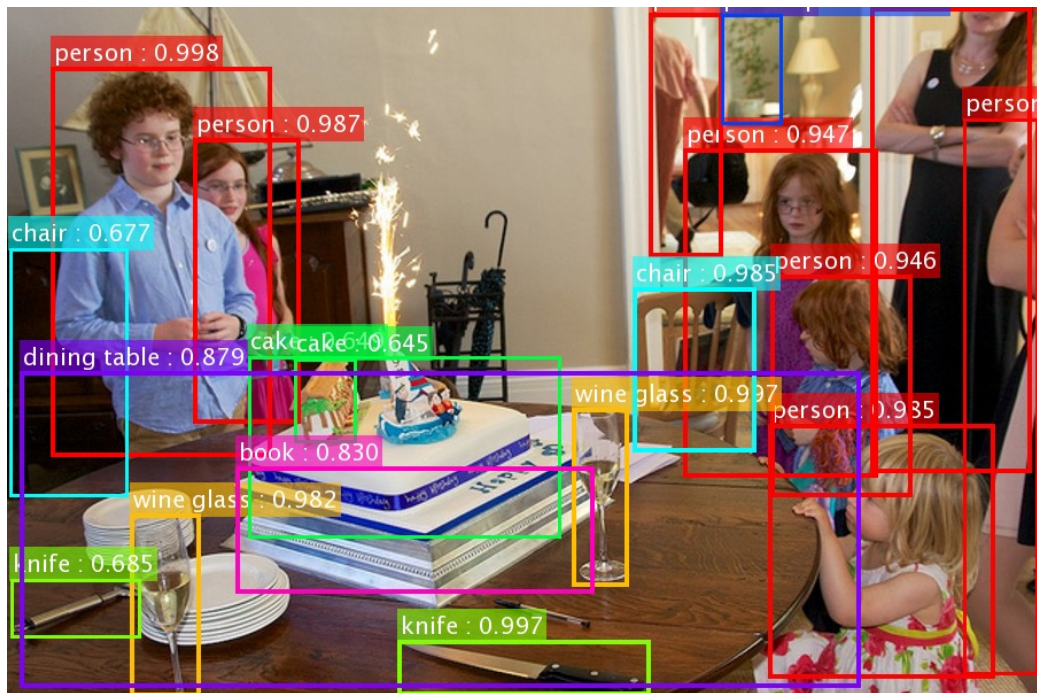
- End-to-End Differentiable Programming



- Propose RoIs by RPN

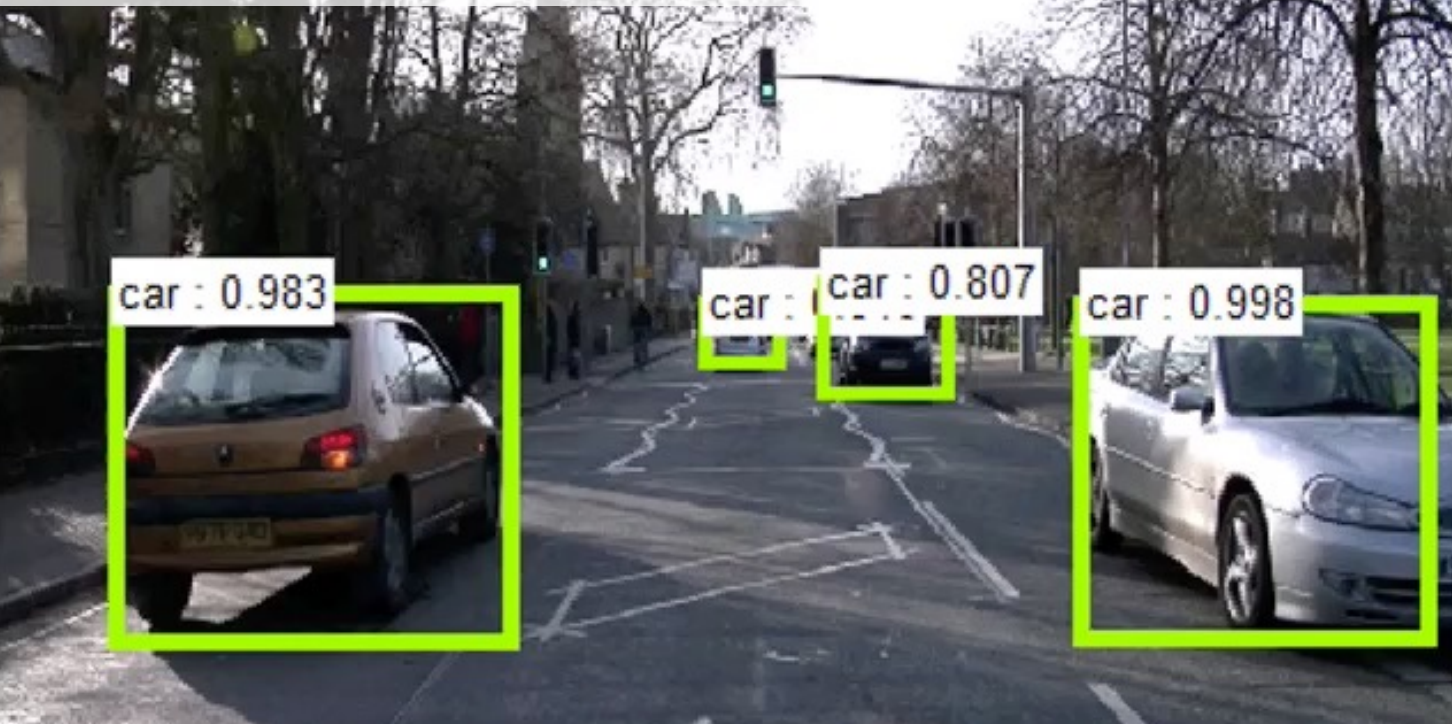
- Same as Fast R-CNN
- Apply RPN's RoIs
- Shared ConvNet backbone on whole image

Faster R-CNN result in 2015

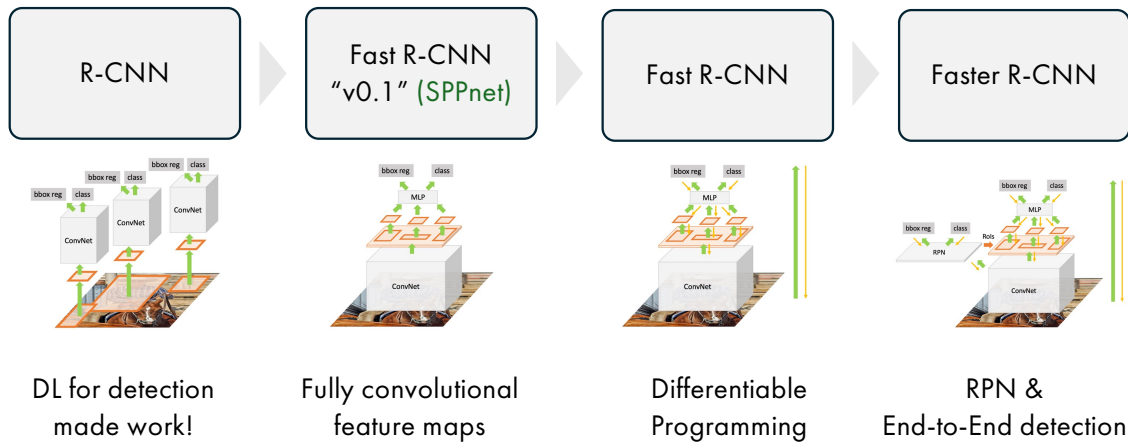


Shaoqing Ren, Kaiming He, Ross Girshick, Jian Sun, "Faster R-CNN: Towards Real-Time Object Detection with Region Proposal Networks", NeurIPS 2015

Faster R-CNN result in 2015



Evolution of R-CNN



Summary

- motivation / overview of recognition tasks
- CNNs
 - Classification
 - Segmentation
 - Object Detection

Takeaways

- what are the building blocks of CNNs?
- what are advantages of CNNs vs FCNs?
- what are the trainable parameters of CNNs?
- what are the inputs/outputs of recognition tasks?
- how do we measure success?
- what are the tradeoffs between different methods for object detection?

Next Time

- Transformers!