

What can('t) we do with ConvNets?

Classification architectures + Semantic segmentation

Karel Zimmermann

Czech Technical University in Prague

Faculty of Electrical Engineering, Department of Cybernetics



Outline

- Architectures of classification networks
- Architectures of segmentation networks
- Architectures of regression networks
- Architectures of detection networks
- Architectures of feature matching networks



Classification results

<http://image-net.org/challenges/LSVRC/2017/index>

Label: **Steel drum**





Classification results

<http://image-net.org/challenges/LSVRC/2017/index>

Label: **Steel drum**



Output:
Scale
T-shirt
Steel drum
Drumstick
Mud turtle



Classification results

<http://image-net.org/challenges/LSVRC/2017/index>

Label: **Steel drum**



Output:
Scale
T-shirt
Steel drum
Drumstick
Mud turtle



Output:
Scale
T-shirt
Giant panda
Drumstick
Mud turtle



Classification results

<http://image-net.org/challenges/LSVRC/2017/index>

Label: **Steel drum**



Output:
Scale
T-shirt
Steel drum
Drumstick
Mud turtle



Output:
Scale
T-shirt
Giant panda
Drumstick
Mud turtle

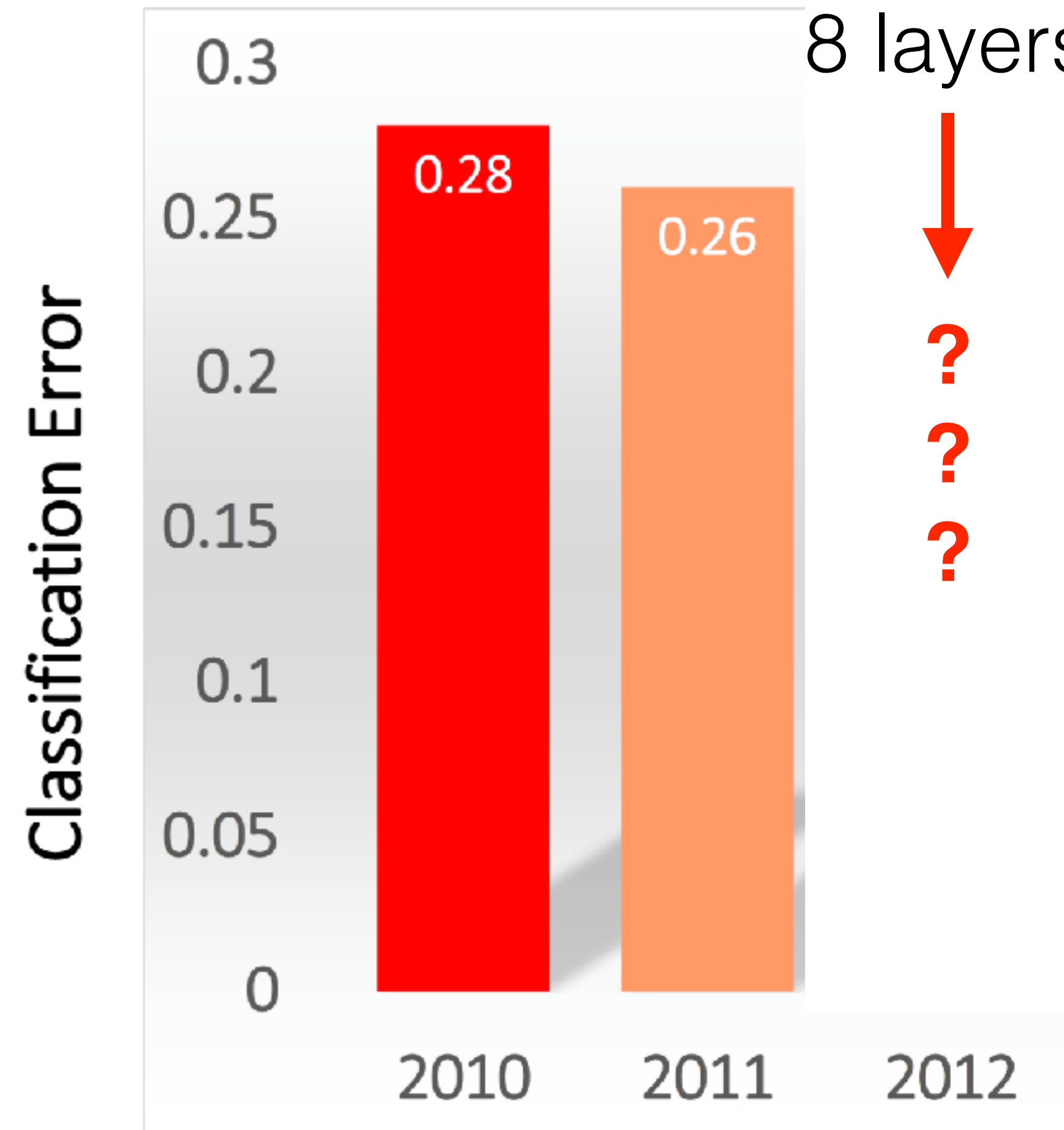


$$\text{Error} = \frac{1}{100,000} \sum_{100,000 \text{ images}} 1[\text{incorrect on image } i]$$

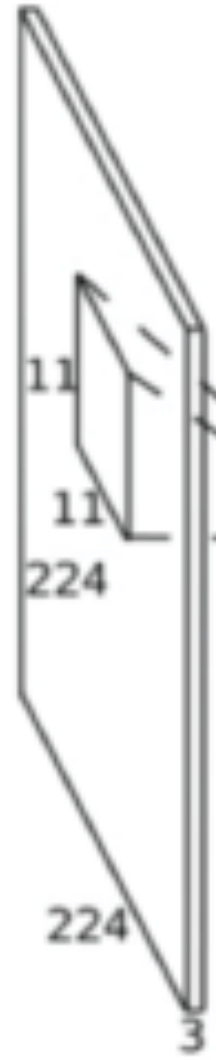
Classification results

AlexNet

8 layers

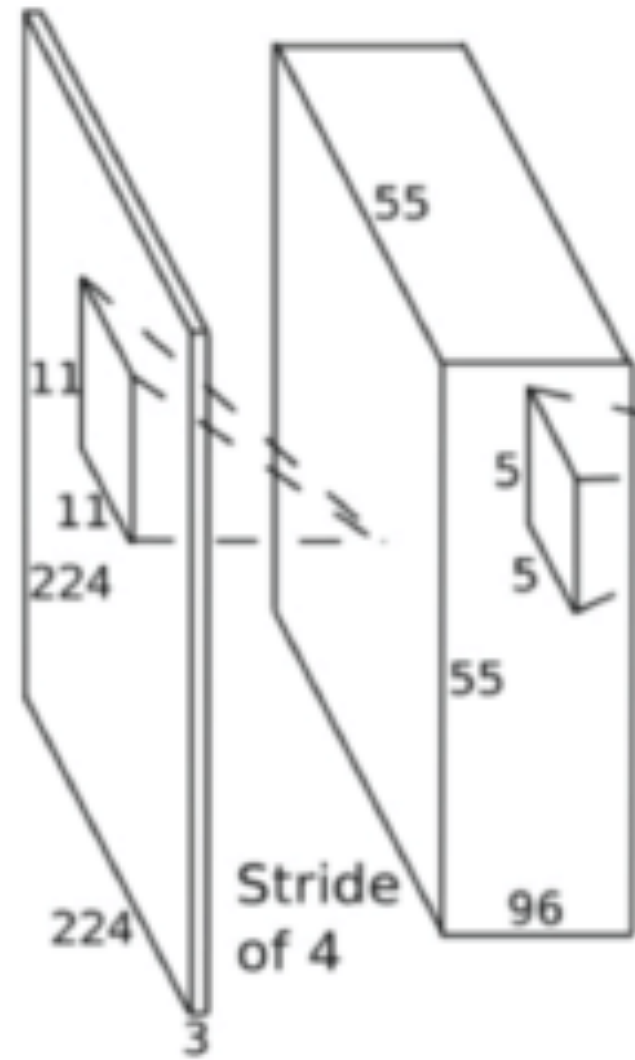


AlexNet on ImageNet 2012 (**over 27k citations !!!**)



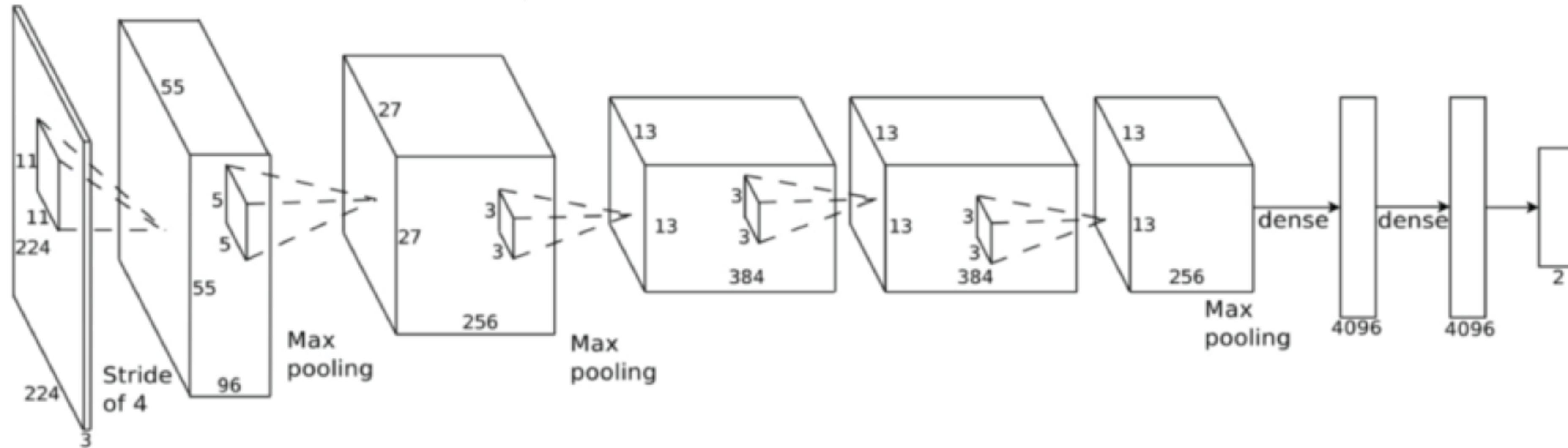
- Param in layer1 (conv, 96 11x11 filters, stride=4, pad=0)?

AlexNet on ImageNet 2012 (**over 27k citations !!!**)



- Param in layer1 (conv, 96 11x11 filters, stride=4, pad=0)?
- Param in layer2 (maxp, 3x3 filters, stride=2, pad=0)?

AlexNet on ImageNet 2012 (**over 27k citations !!!**)

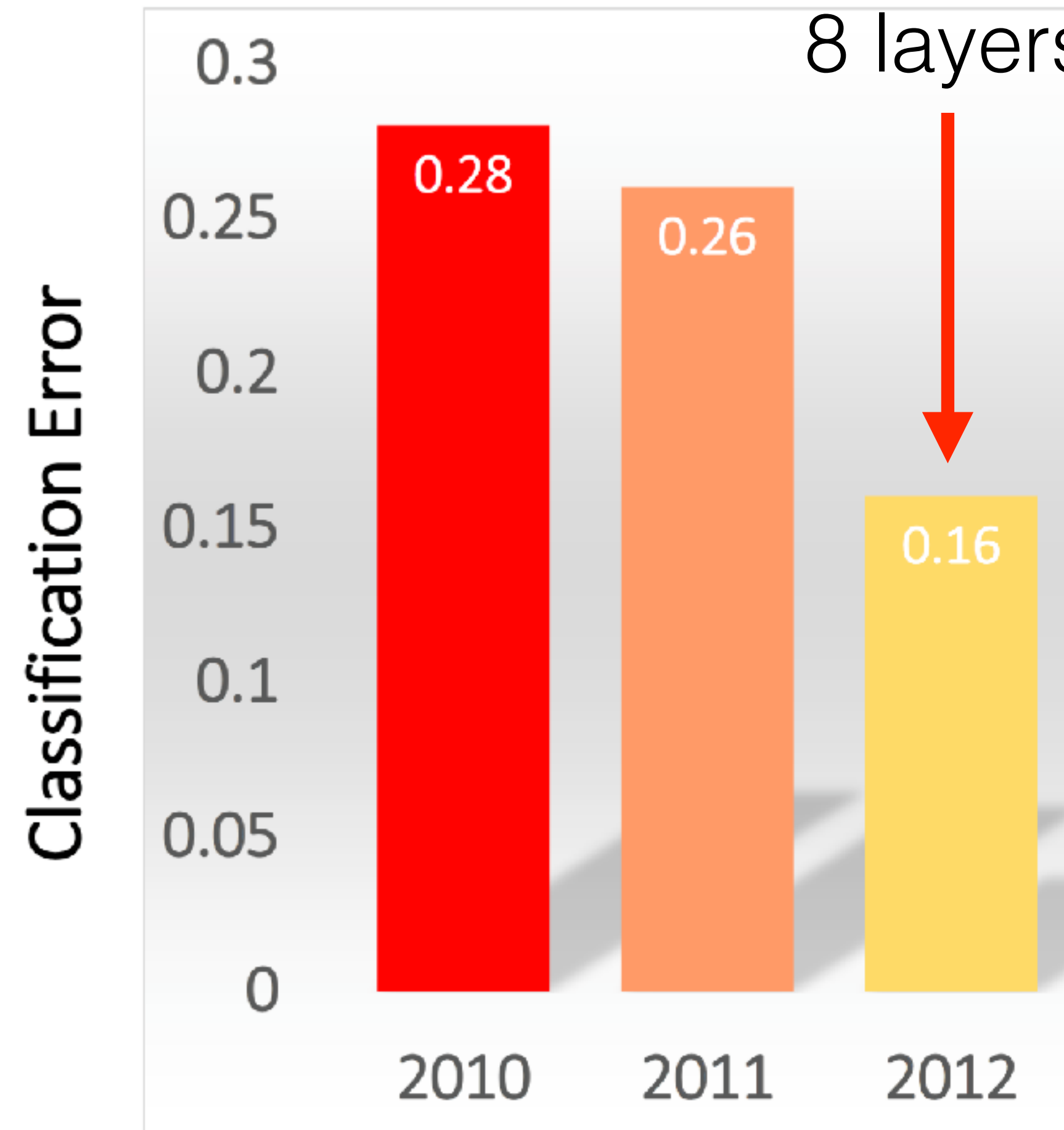


- Param in layer1 (conv, 96 11x11 filters, stride=4, pad=0)?
- Param in layer2 (maxp, 3x3 filters, stride=2, pad=0)?
- Param in layer3 (conv, 256 5x5 filters, stride=1, pad=2)?
- Parameters in total: 60M, Depth: 8 layers

Classification results

AlexNet

8 layers



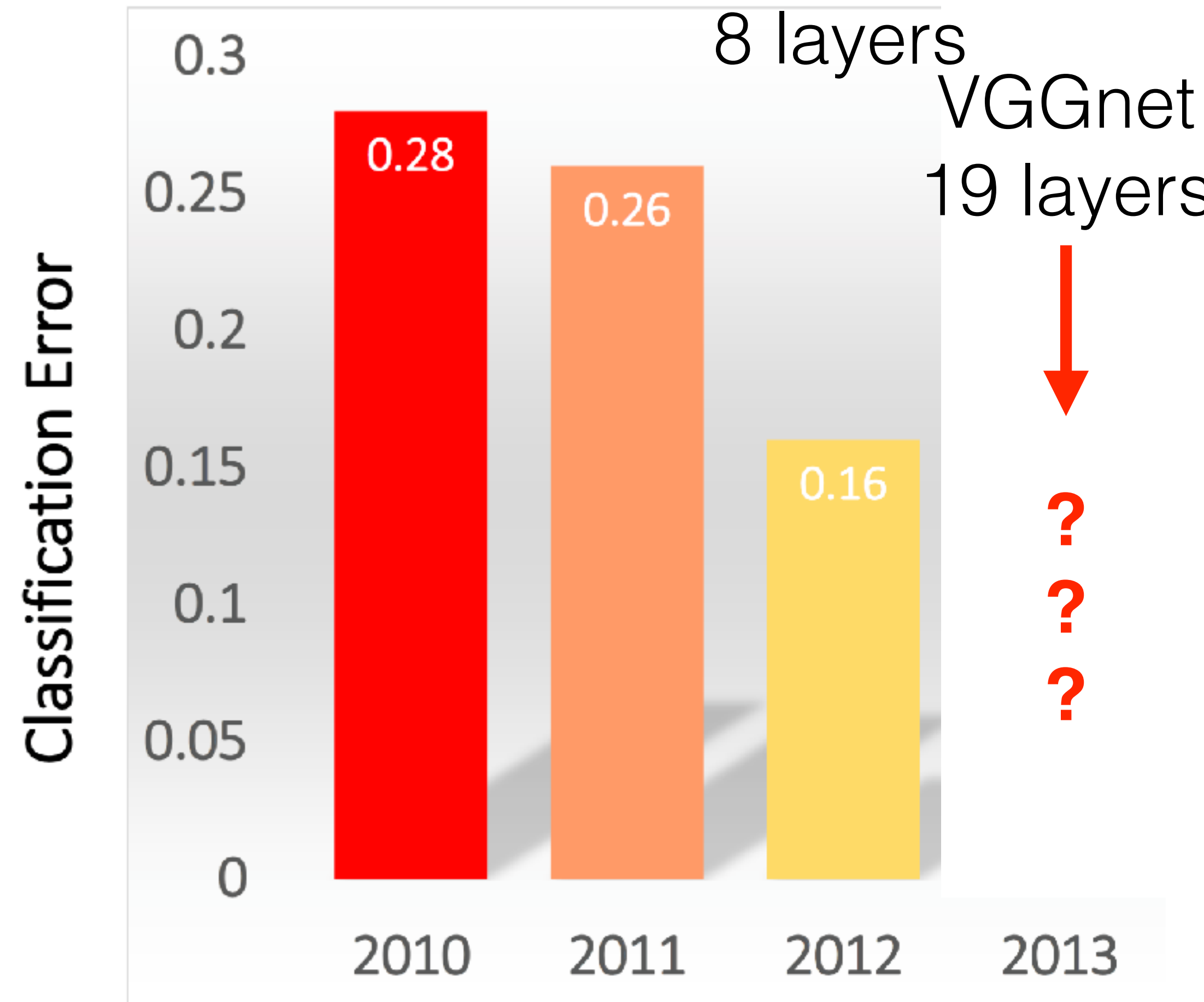
Classification results

AlexNet

8 layers

VGGnet

19 layers



VGGNet vs AlexNet



- large filters
- shallow (8 layers)

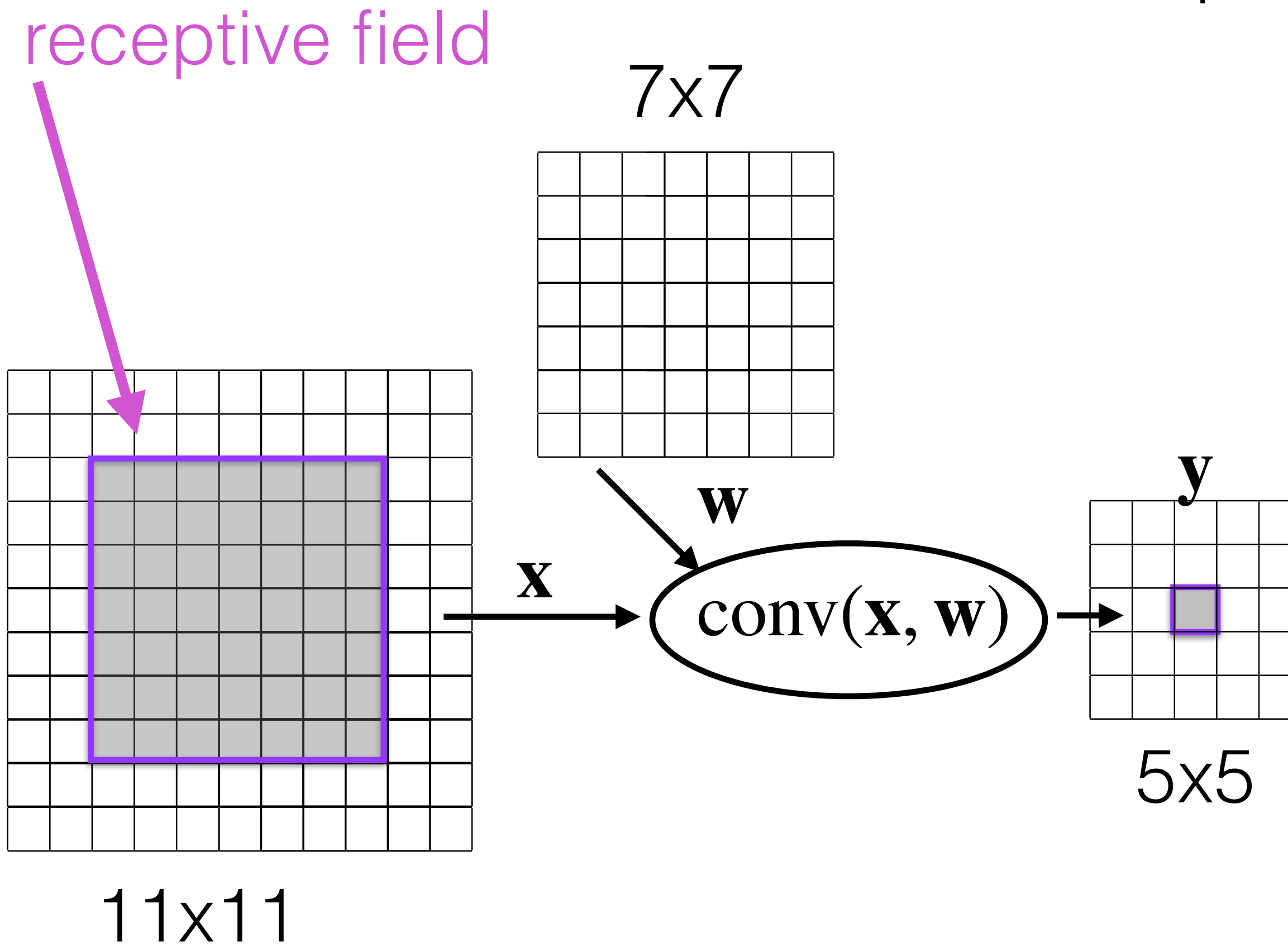


- small filters
- deeper (19 layers)

- Parameters in total: 138M, Depth: 19 layers

Simonyan and Zissermann, Very Deep Convolutional Networks for Large Scale Image Recognition, 2014
<https://arxiv.org/abs/1409.1556>

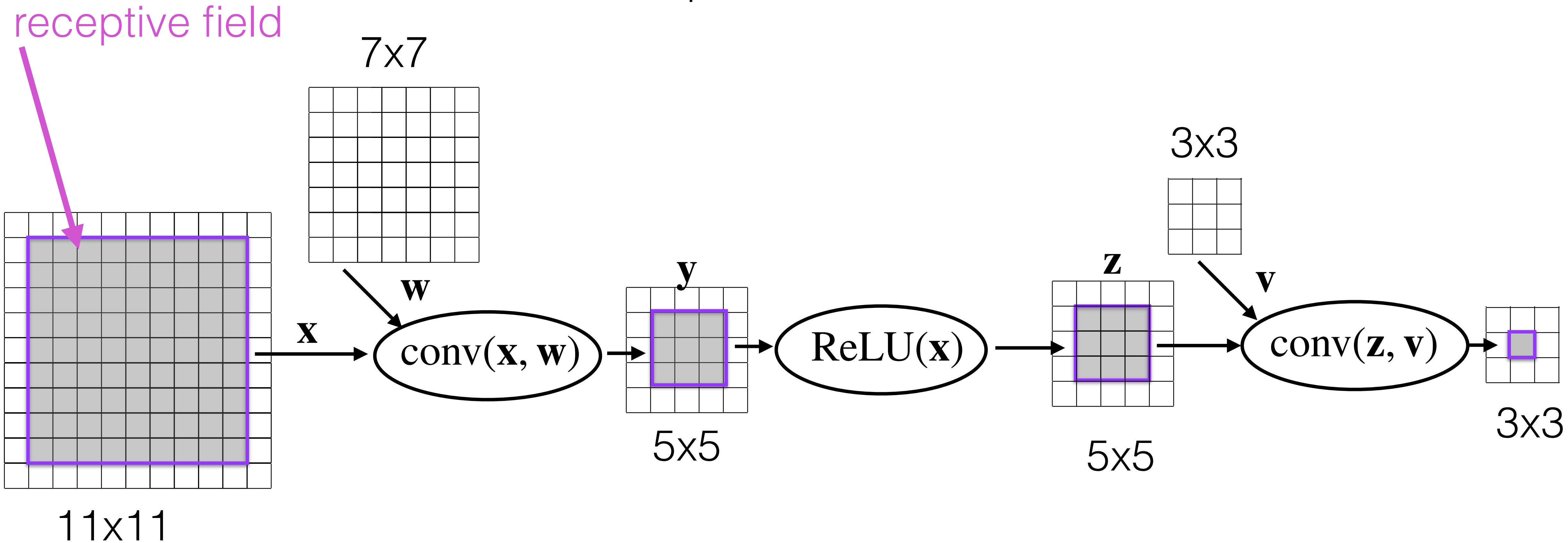
Receptive field



Receptive field: region in the input space that affect activation of a particular neuron.

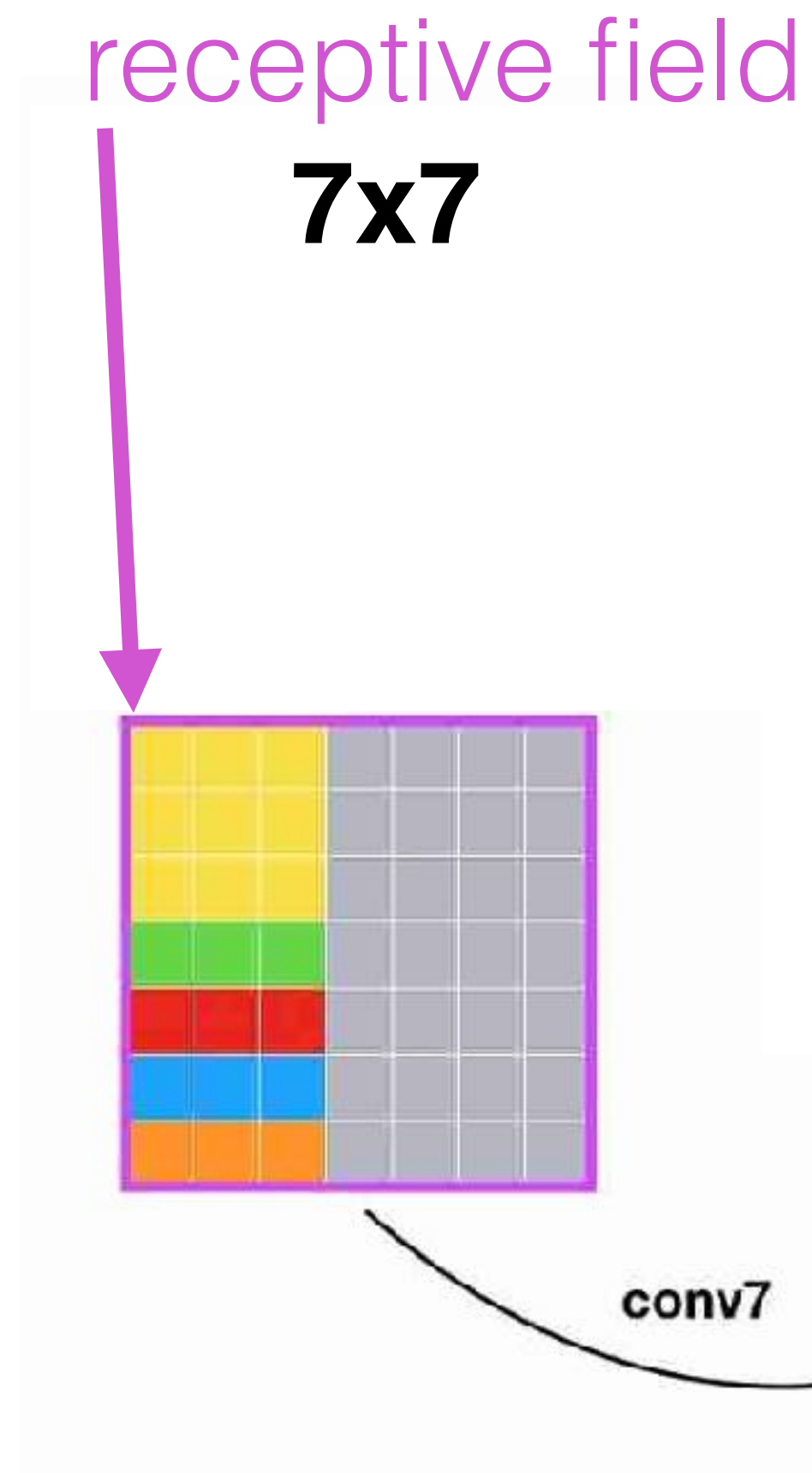
What is the size of the receptive field??? 7×7

Receptive field



Receptive field: region in the input space that affect activation of a particular neuron.
What is the size of the receptive field??? 9x9

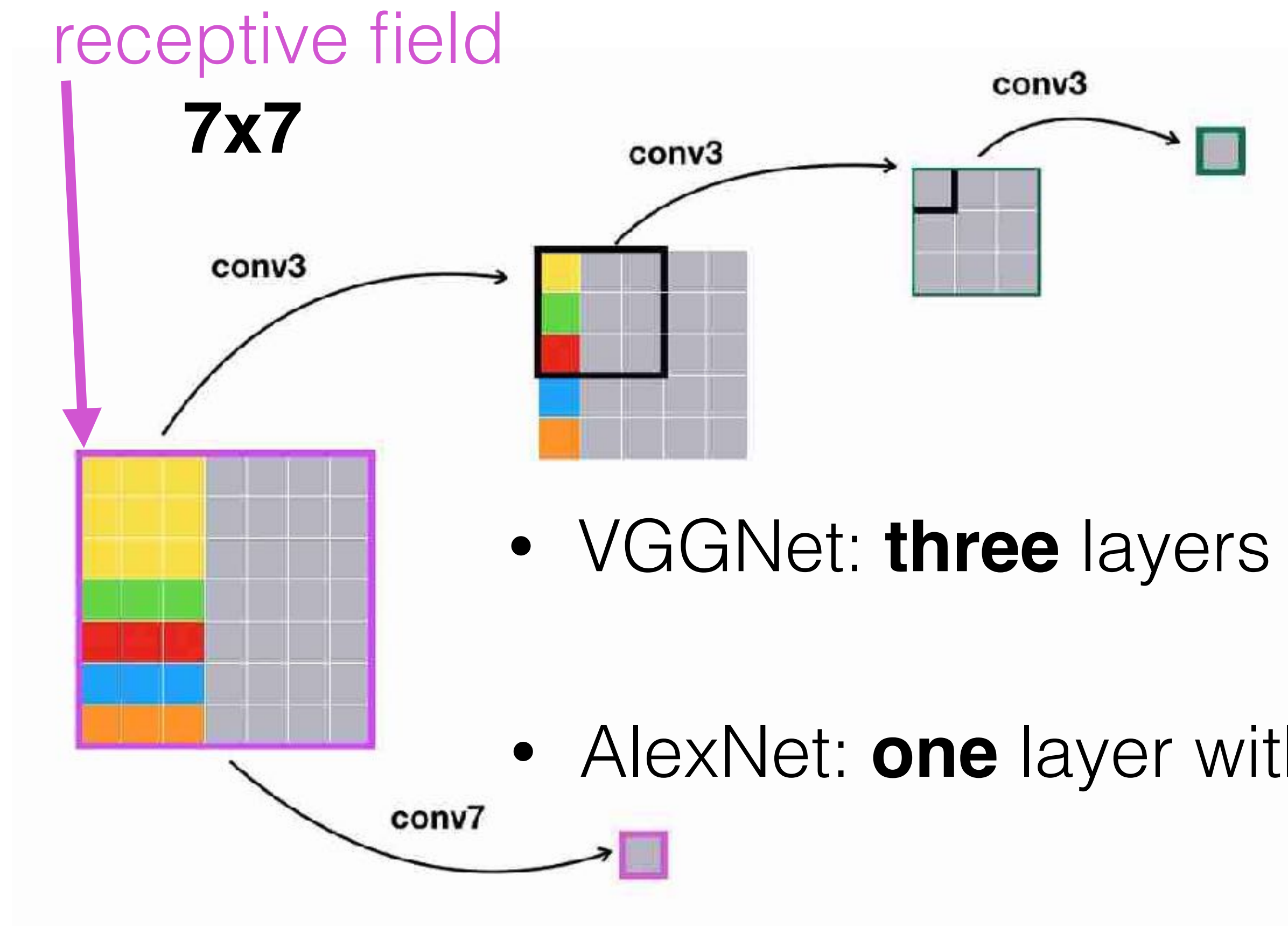
VGGNet vs AlexNet



- AlexNet: **one** layer with **7x7** filter (49+1 params)

Receptive field: region in the input space that affect activation of a particular neuron.

VGGNet vs AlexNet



- VGGNet: **three** layers with **3x3** filters (3x9+3 params)
- AlexNet: **one** layer with **7x7** filter (49+1 params)

Receptive field: region in the input space that affect activation of a particular neuron.

VGGNet has **the same receptive field** with **less parameters**

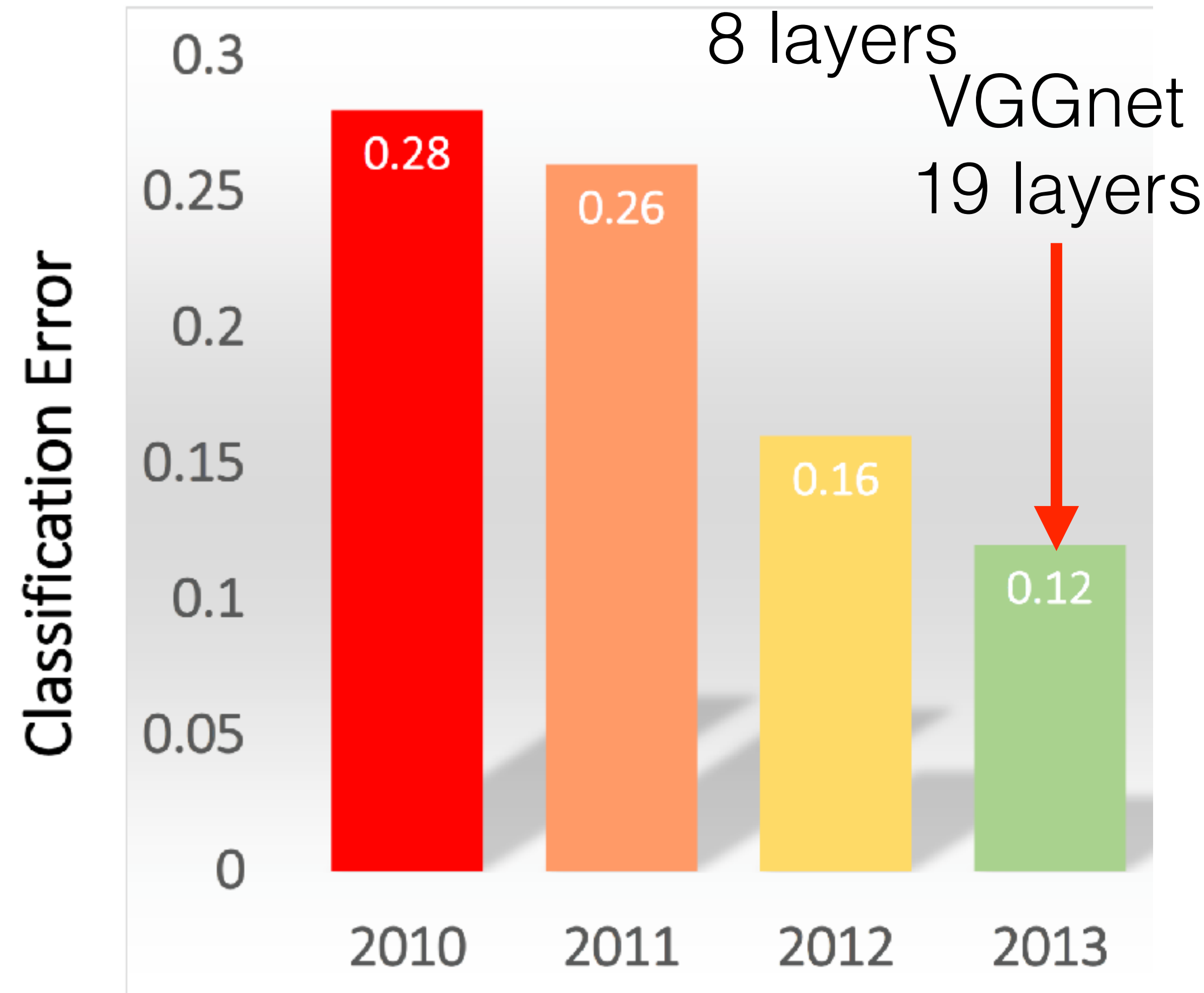
Classification results

AlexNet

8 layers

VGGnet

19 layers



Classification results

AlexNet

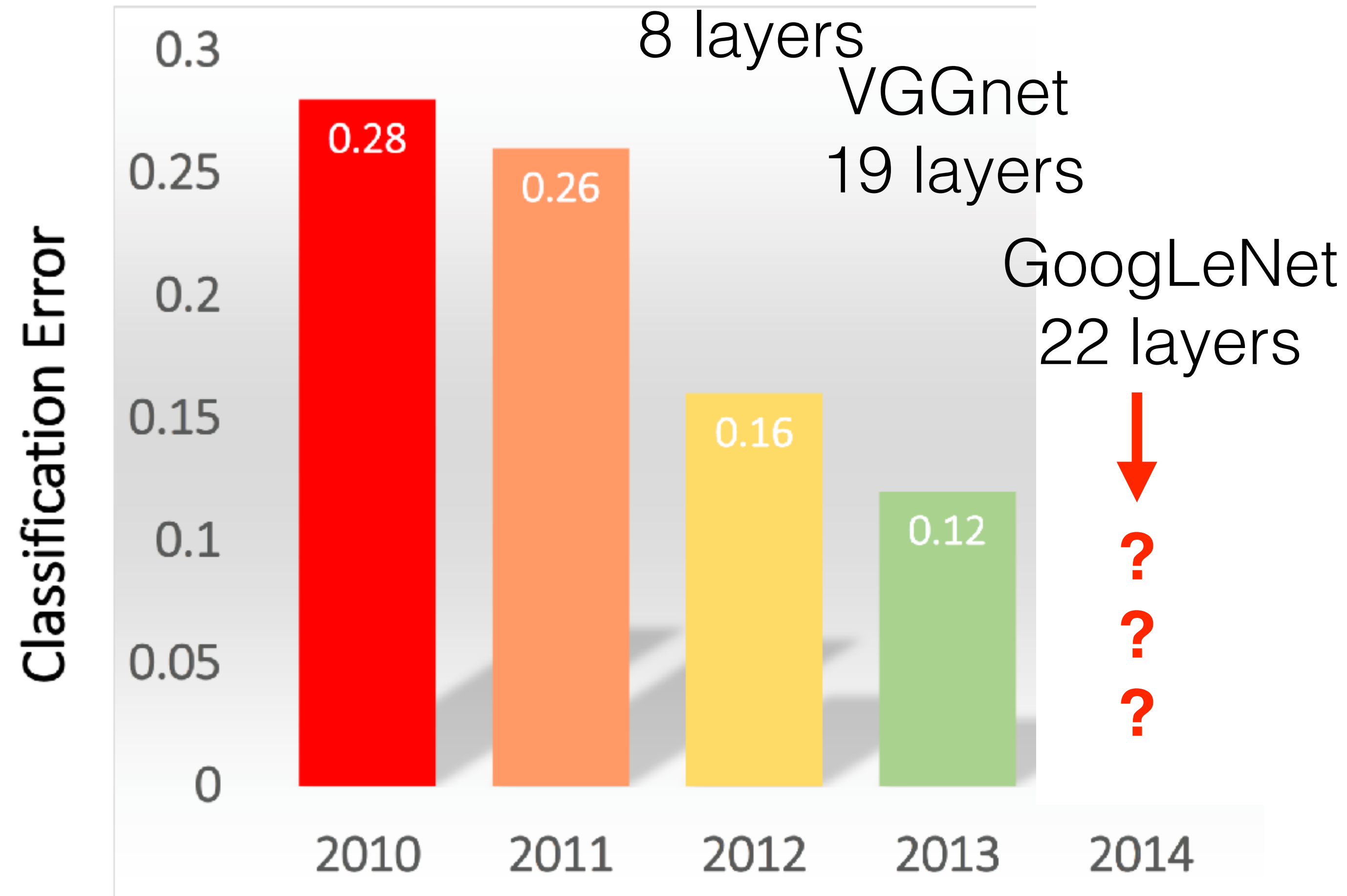
8 layers

VGGnet

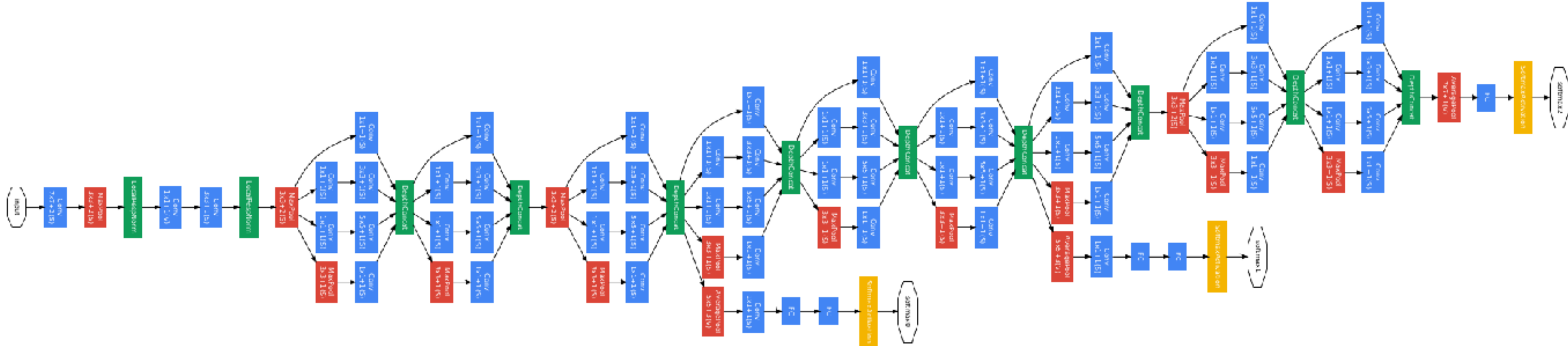
19 layers

GoogLeNet

22 layers

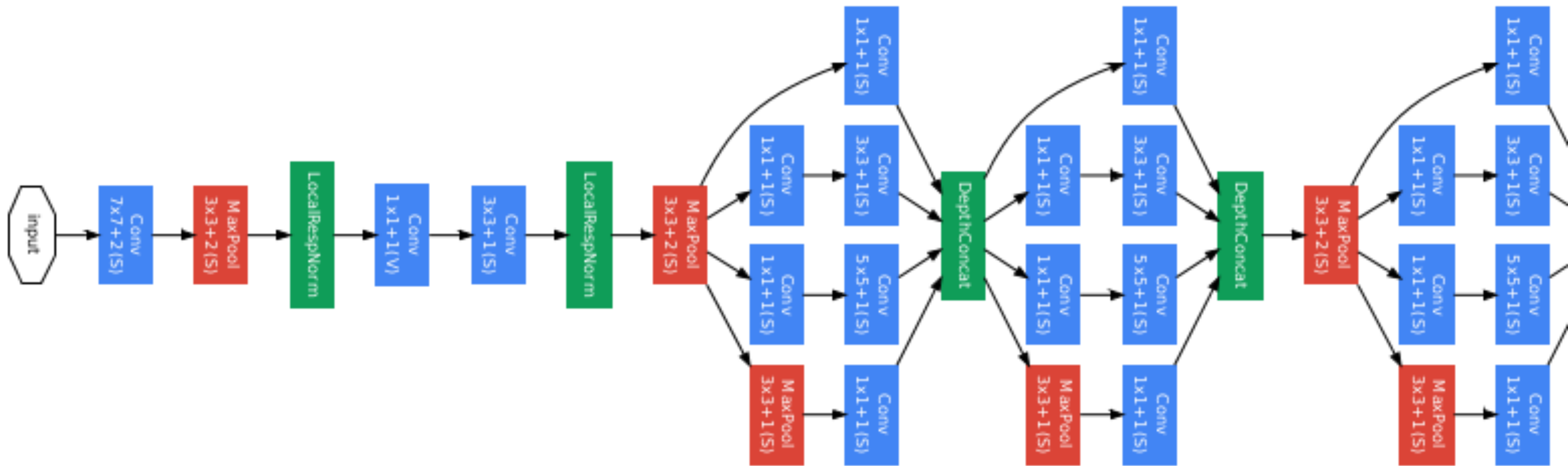


GoogLeNet



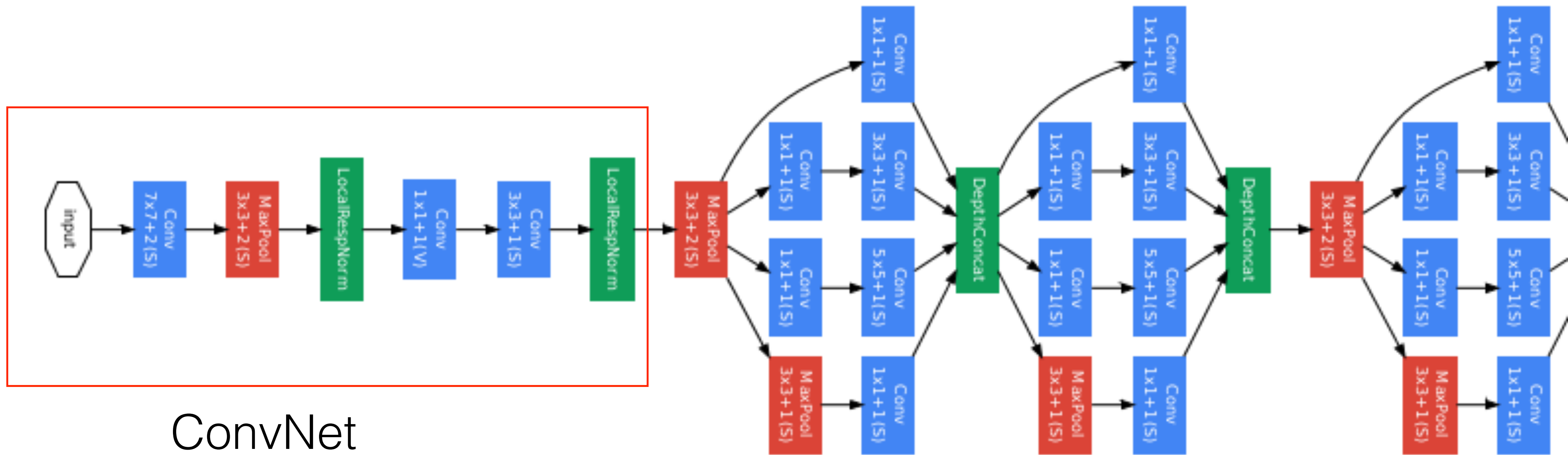
Szegedy et al. Going Deeper with Convolutions, CVPR, 2014
<https://arxiv.org/abs/1409.4842>

GoogLeNet



Szegedy et al. Going Deeper with Convolutions, CVPR, 2014
<https://arxiv.org/abs/1409.4842>

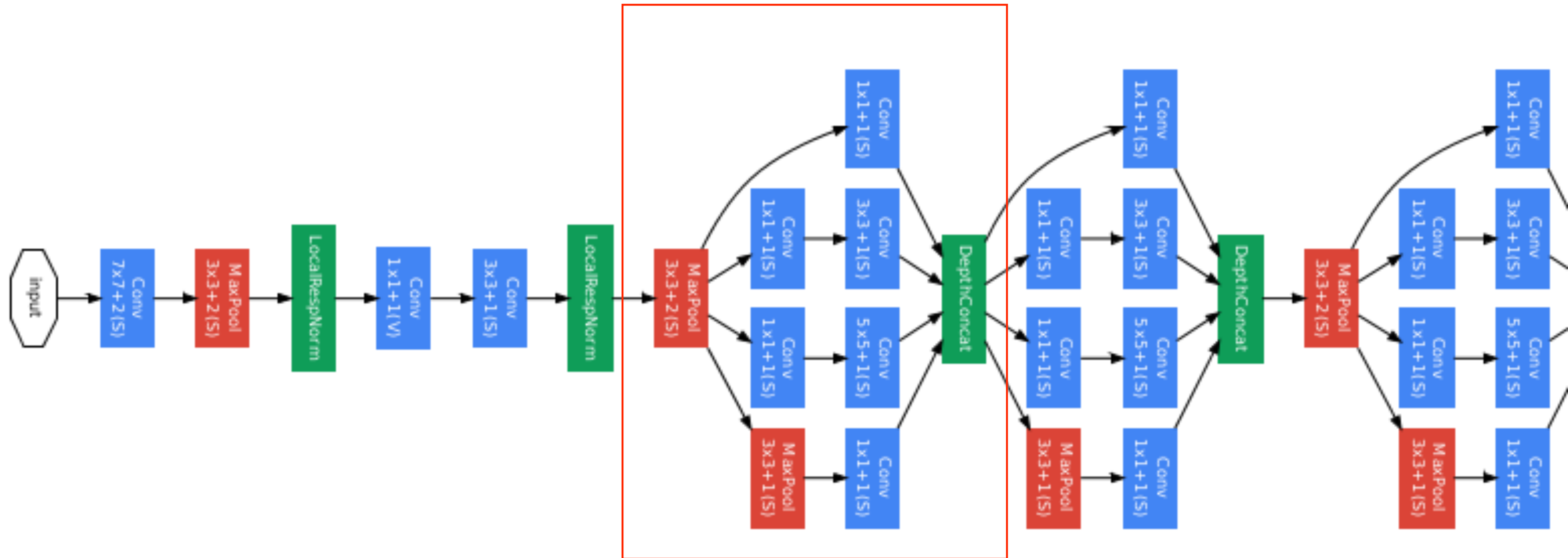
GoogLeNet



ConvNet

Szegedy et al. Going Deeper with Convolutions, CVPR, 2014
<https://arxiv.org/abs/1409.4842>

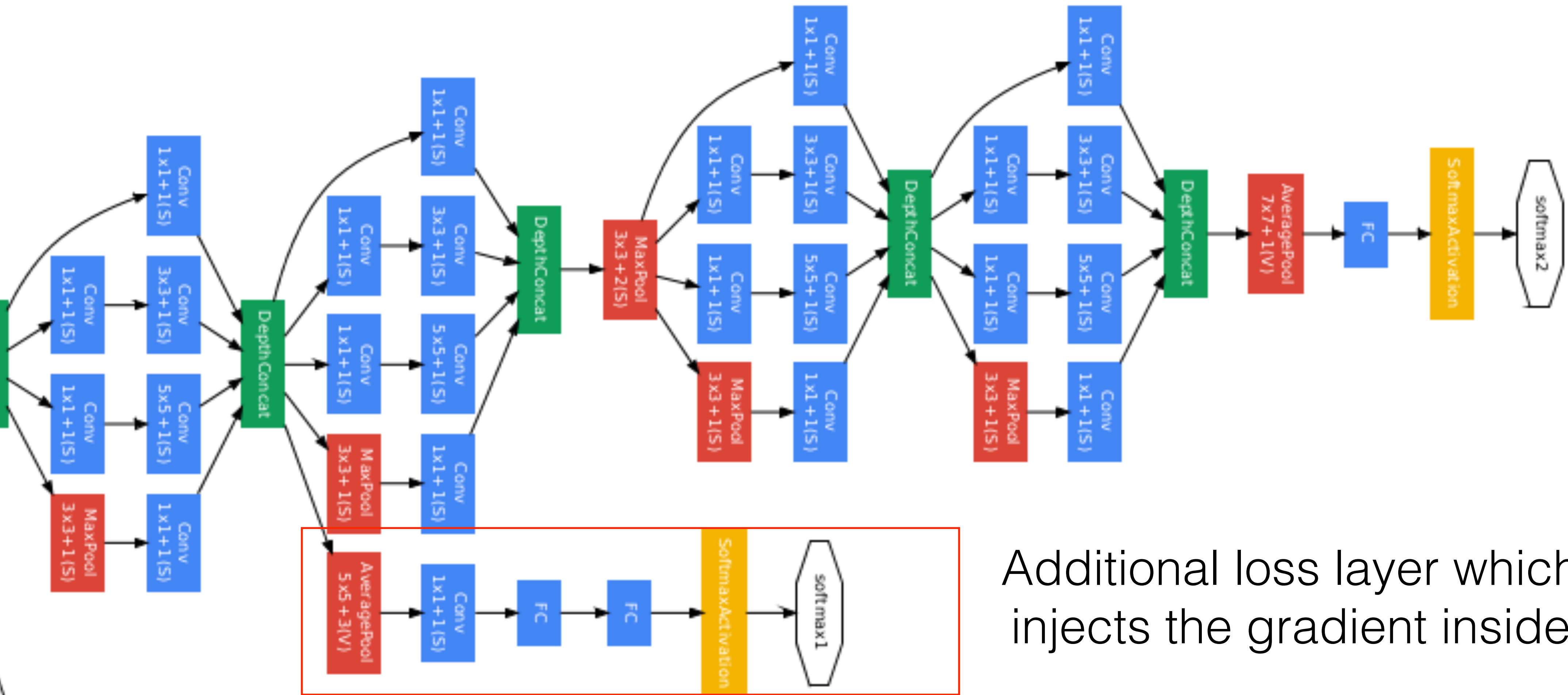
GoogLeNet



Inception module

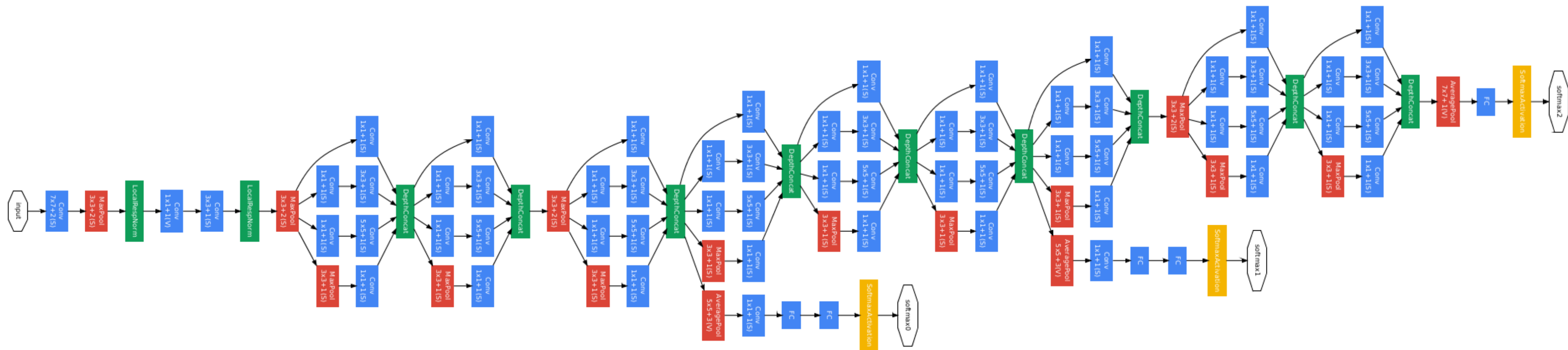
Szegedy et al. Going Deeper with Convolutions, CVPR, 2014
<https://arxiv.org/abs/1409.4842>

GoogLeNet



Szegedy et al. Going Deeper with Convolutions, CVPR, 2014
<https://arxiv.org/abs/1409.4842>

GoogLeNet



- 12x fewer parameters than AlexNet
- depth 22 layers
- training: few high-end GPU about a week

Szegedy et al. Going Deeper with Convolutions, CVPR, 2014
<https://arxiv.org/abs/1409.4842>

Classification results

AlexNet

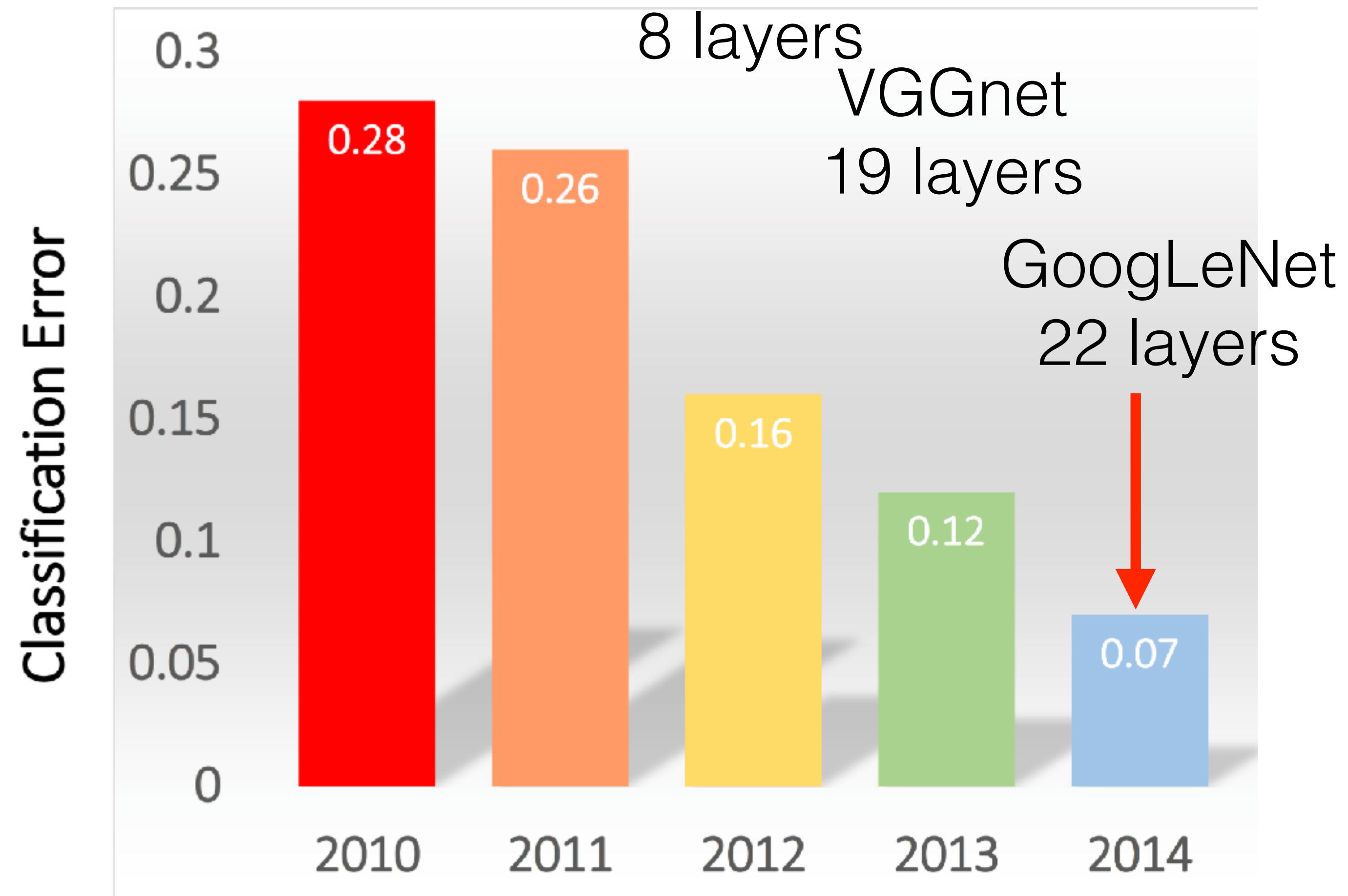
8 layers

VGGnet

19 layers

GoogLeNet

22 layers



Classification results

AlexNet

8 layers

VGGnet

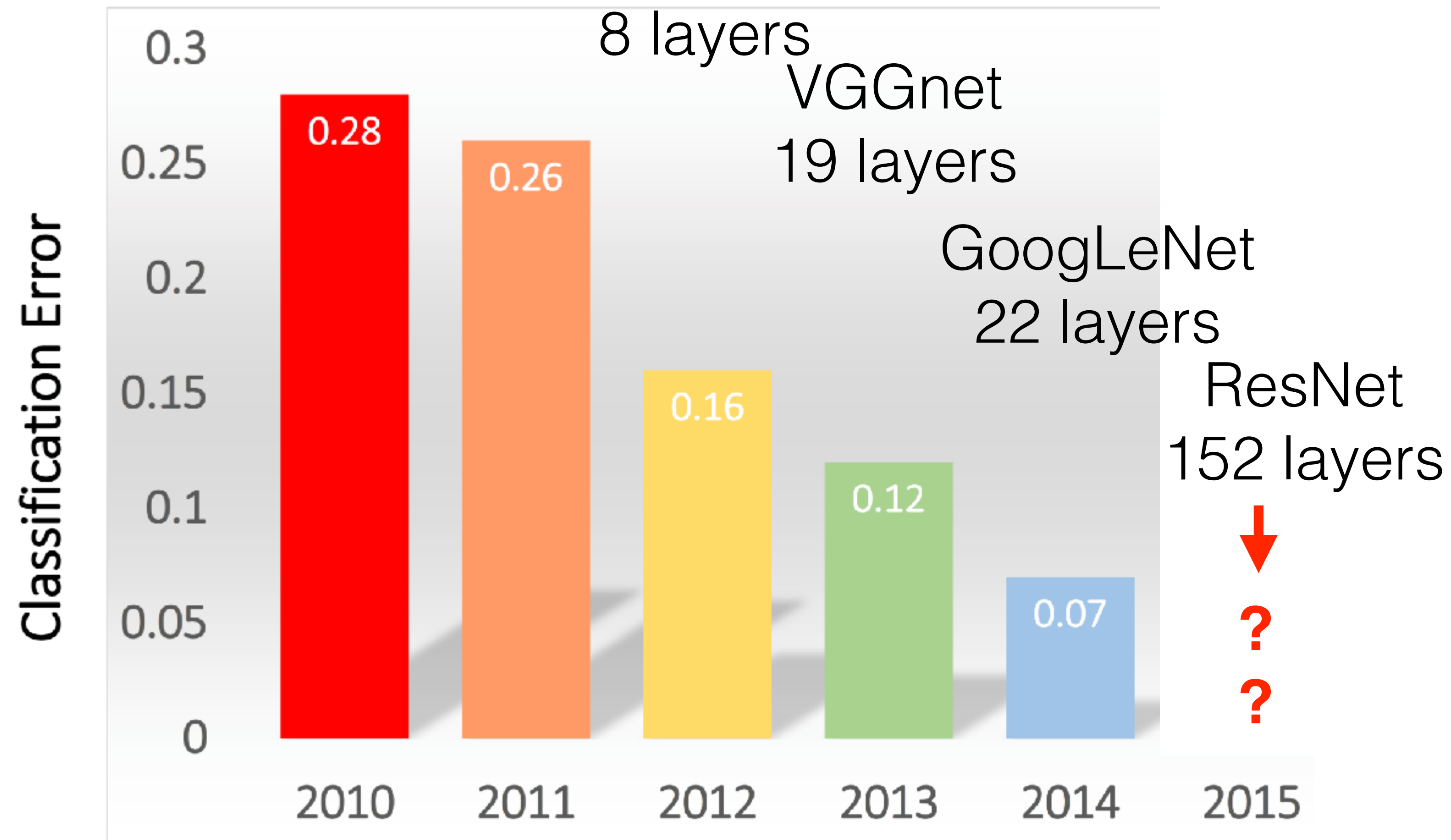
19 layers

GoogLeNet

22 layers

ResNet

152 layers



ResNet

Better results with smaller kernels + deeper architectures



Well said Leo, well said

- deeper ConvNet architectures yielded **higher training errors**. => **is it overfitting?**
- => no overfitting, but vanishing gradient !

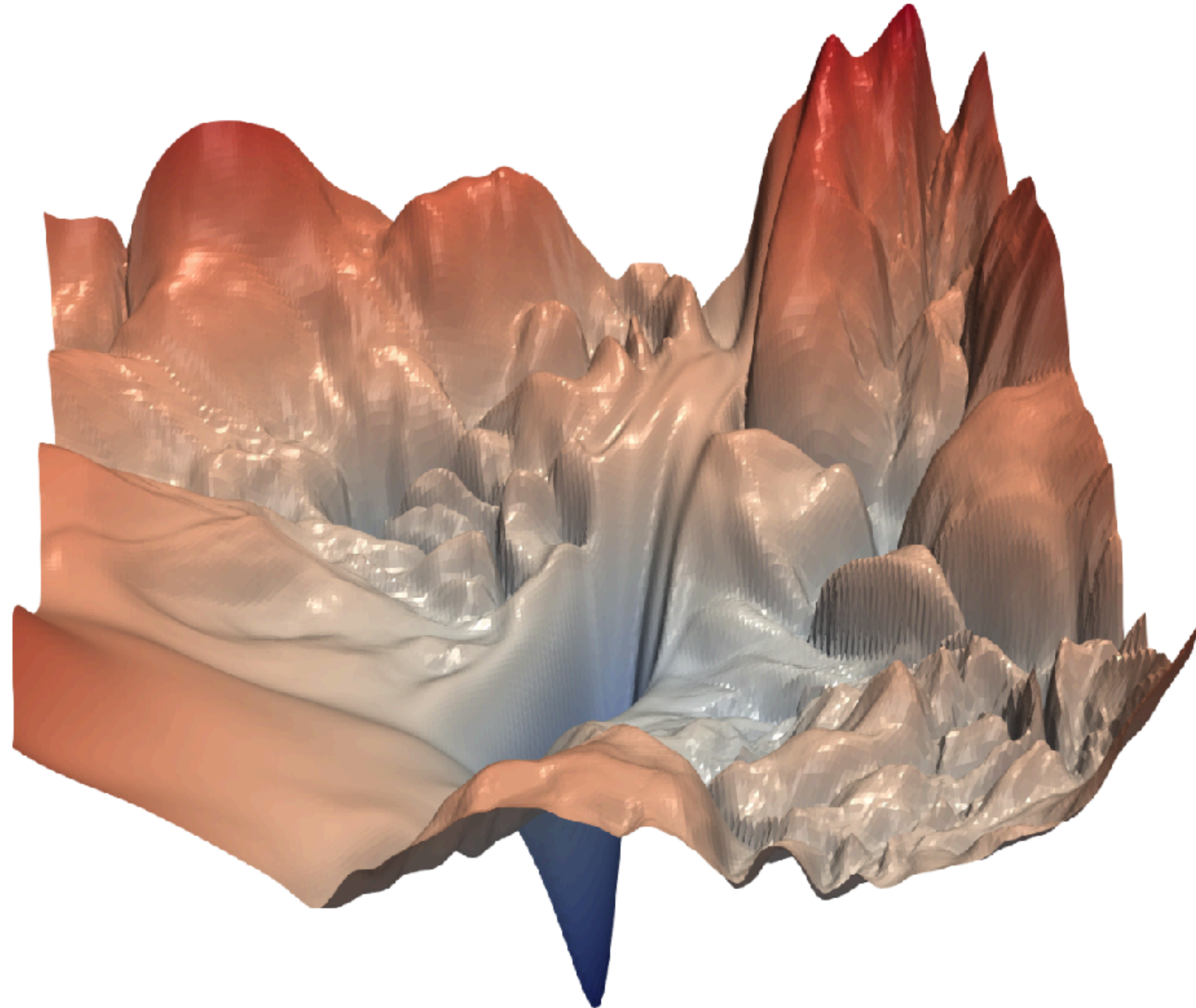
He et al. Going Deeper with Convolutions, CVPR, 2015

<https://arxiv.org/abs/1512.03385>

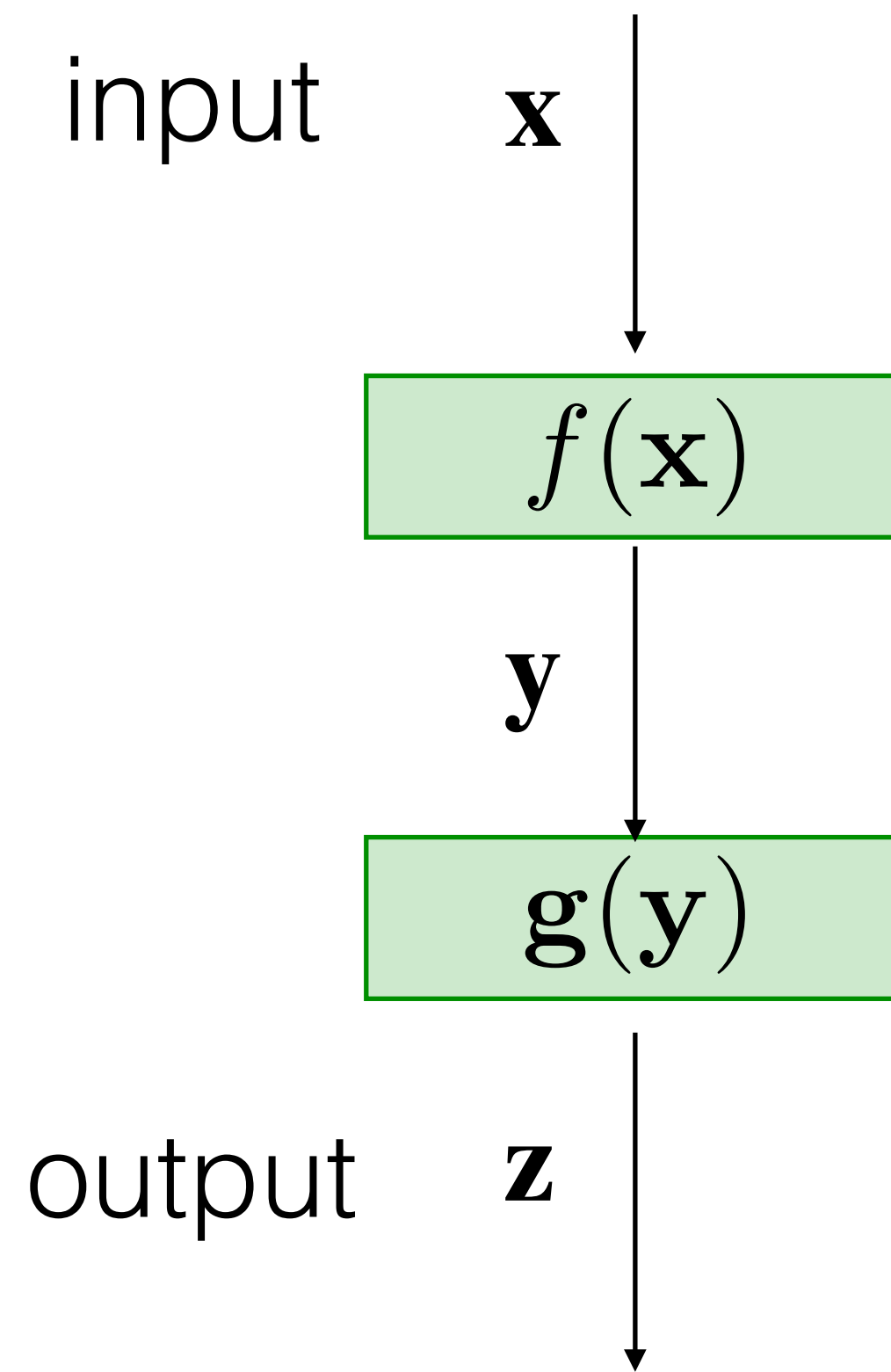
Visualizing Loss Landscape of Neural Nets

$$f(\alpha, \beta) = \mathcal{L}(\mathbf{w}^* + \alpha \mathbf{u} + \beta \mathbf{v})$$

for randomly chosen (and normalized) directions $\mathbf{u}, \mathbf{v}$



forward pass:

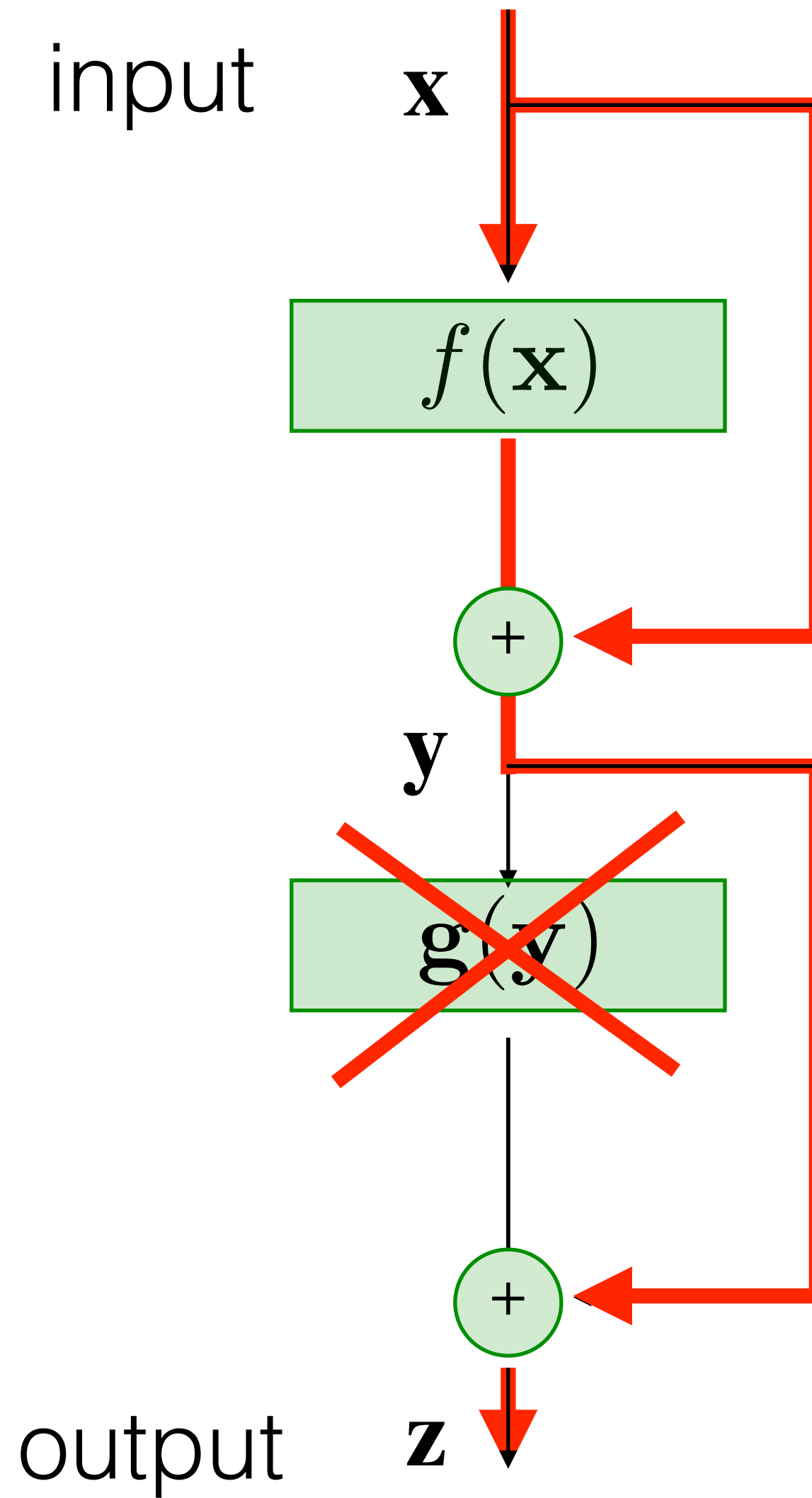


backward pass:

gradient $\frac{\partial \mathbf{z}}{\partial \mathbf{x}} = \begin{bmatrix} \frac{\partial \mathbf{z}}{\partial \mathbf{y}} & \frac{\partial \mathbf{y}}{\partial \mathbf{x}} \end{bmatrix} \approx \mathbf{0}$ then gradient is always zero

if any local gradient is zero

forward pass:



backward pass:

gradient $\frac{\partial \mathbf{z}}{\partial \mathbf{x}} = \left(\frac{\partial \mathbf{z}}{\partial \mathbf{y}} + 1 \right) \left(\frac{\partial \mathbf{y}}{\partial \mathbf{x}} + 1 \right) \neq \mathbf{0}$

$\approx \mathbf{0}$ $\approx \mathbf{0}$

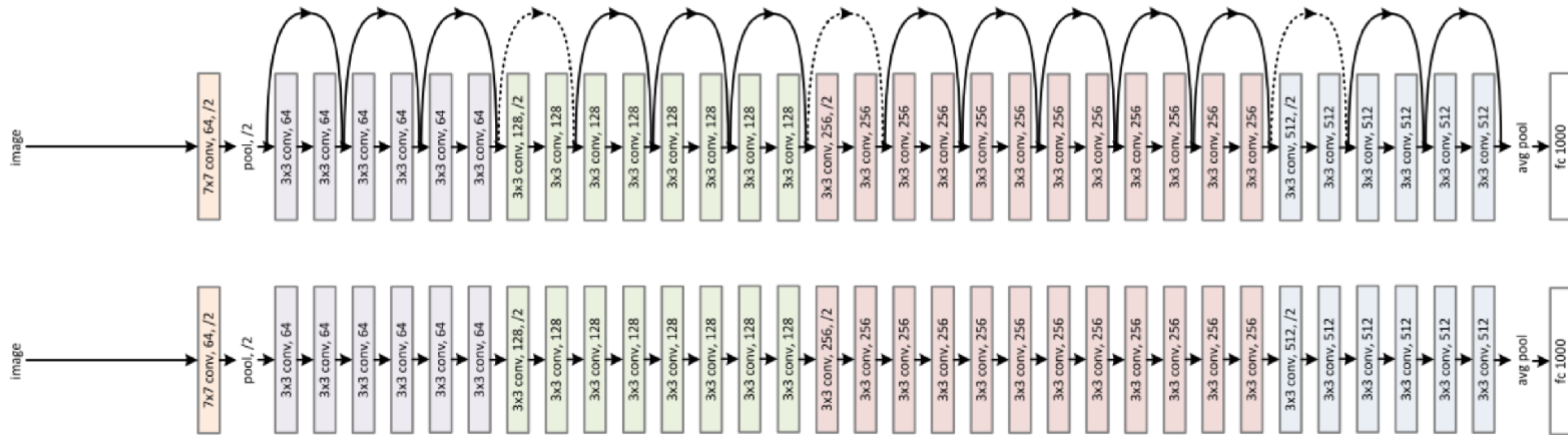
Red arrows point from the terms $\frac{\partial \mathbf{z}}{\partial \mathbf{y}}$ and $\frac{\partial \mathbf{y}}{\partial \mathbf{x}}$ in the equation to the text $\approx \mathbf{0}$ below them, indicating that these local gradients can be zero.

if any local gradient is zero

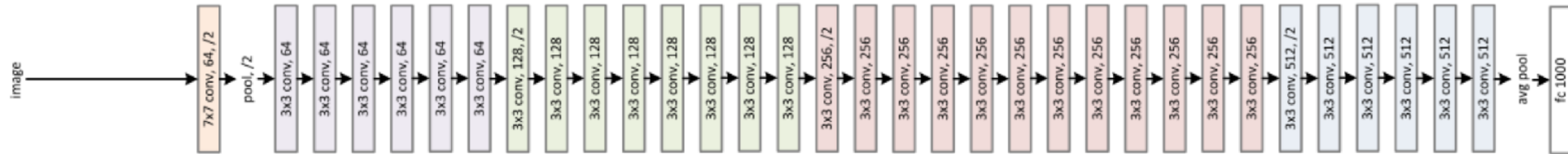
the gradient can still flow through another path

ResNet: deep ConvNet with skip connections

ResNet

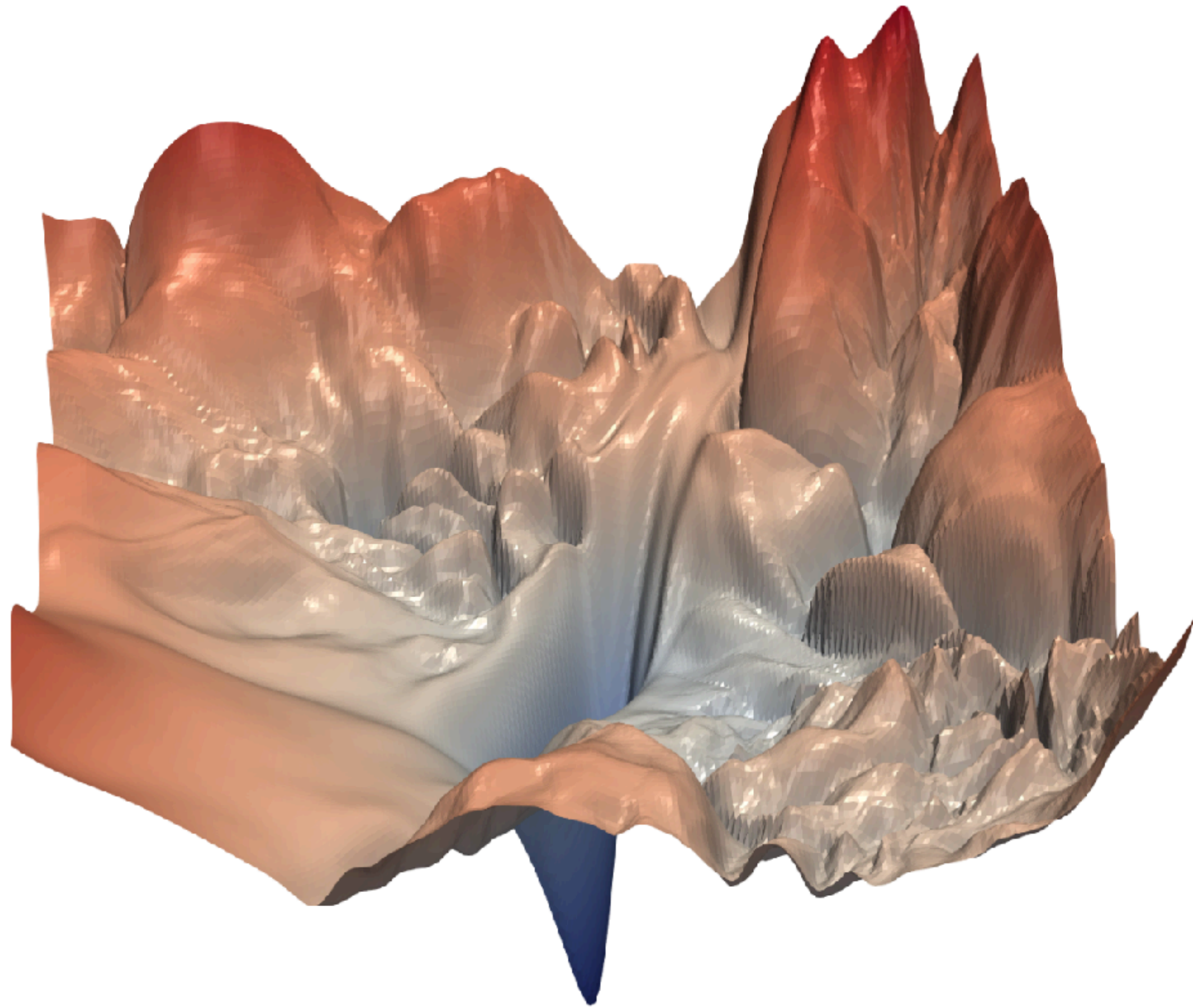


ResNet-NS



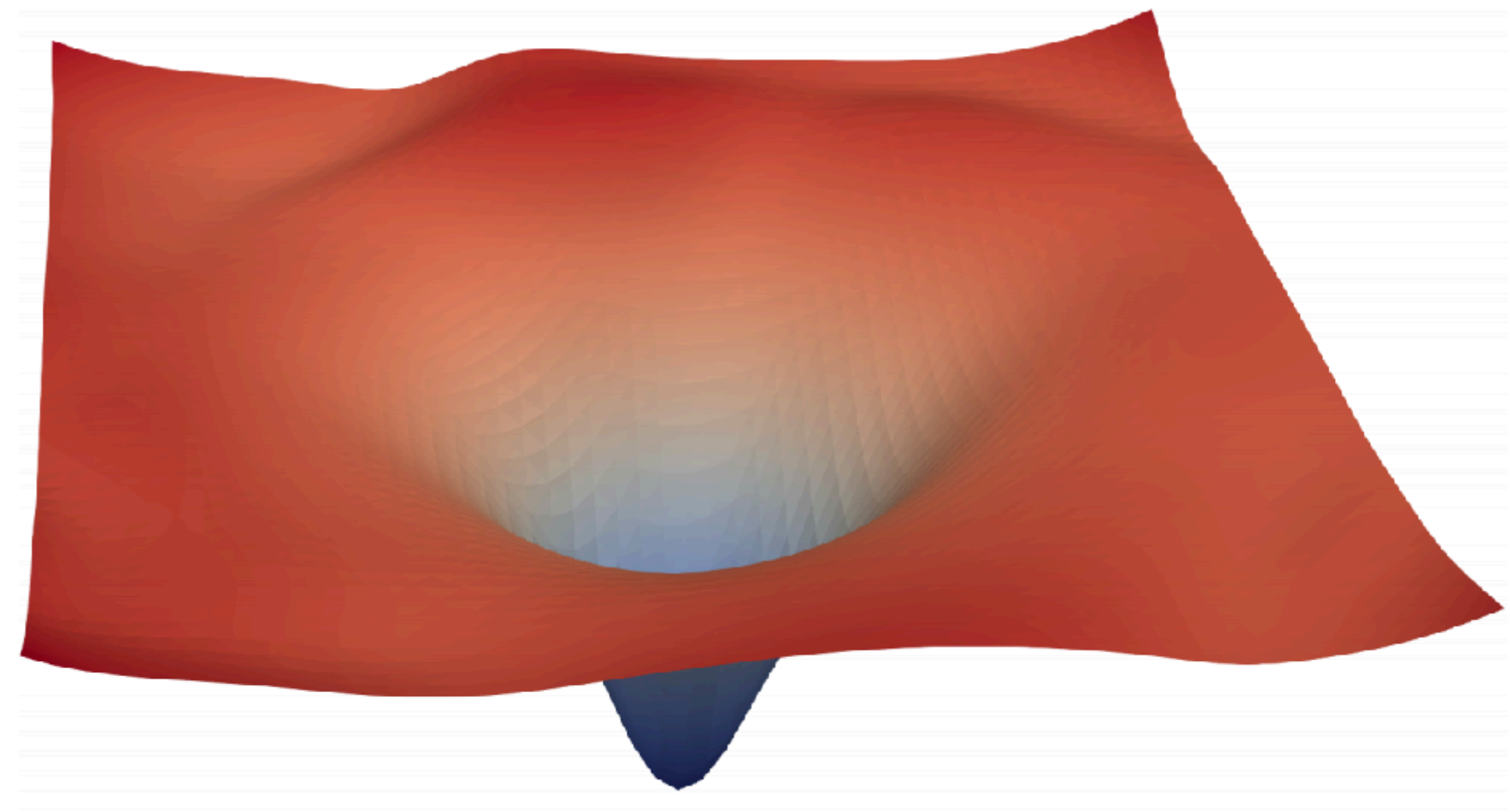
Visualizing Loss Landscape of Neural Nets

[Li et al, NIPS, 2018] <https://arxiv.org/pdf/1712.09913.pdf>



(a) without skip connections

ResNet-NS

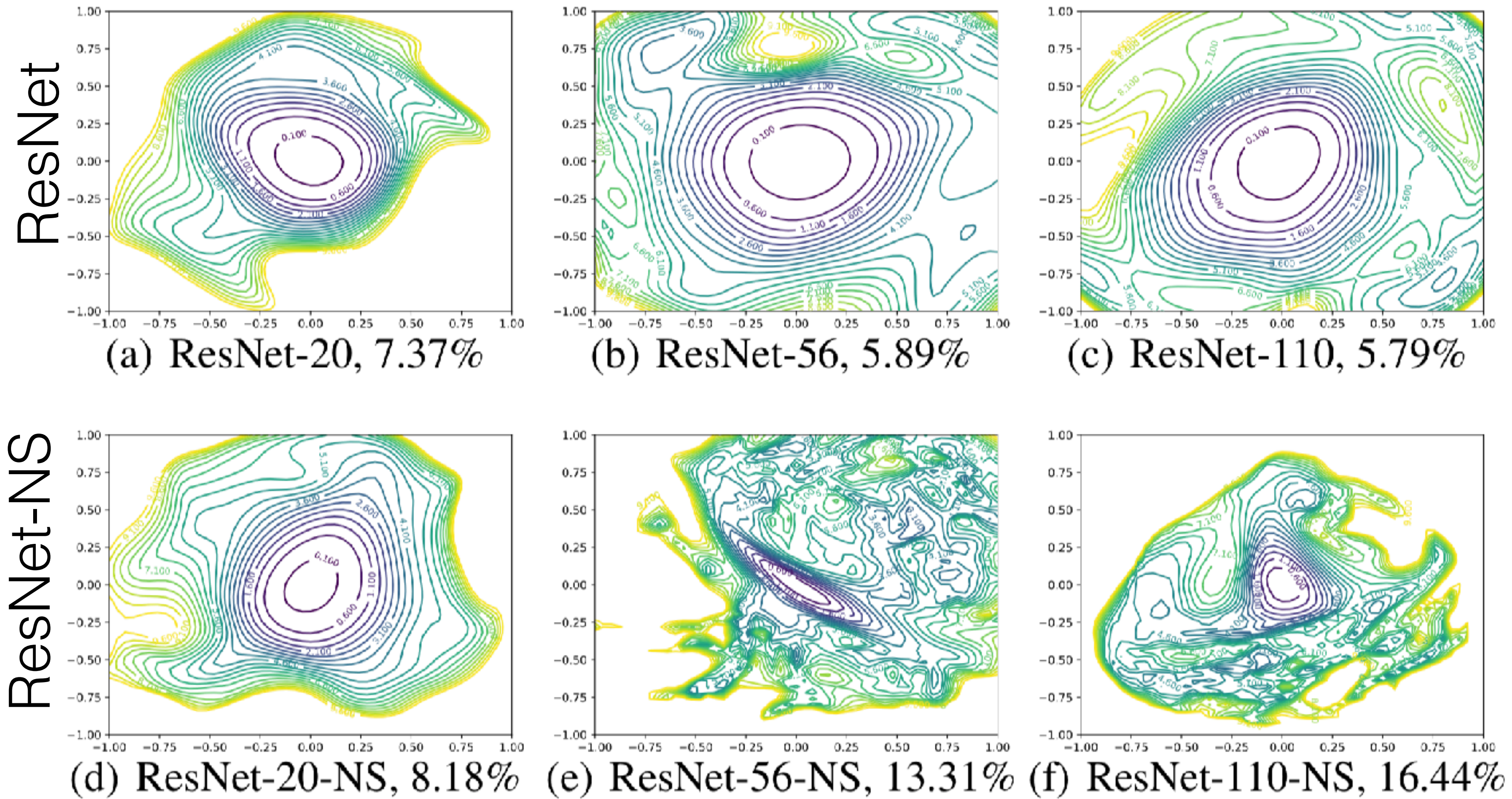


(b) with skip connections

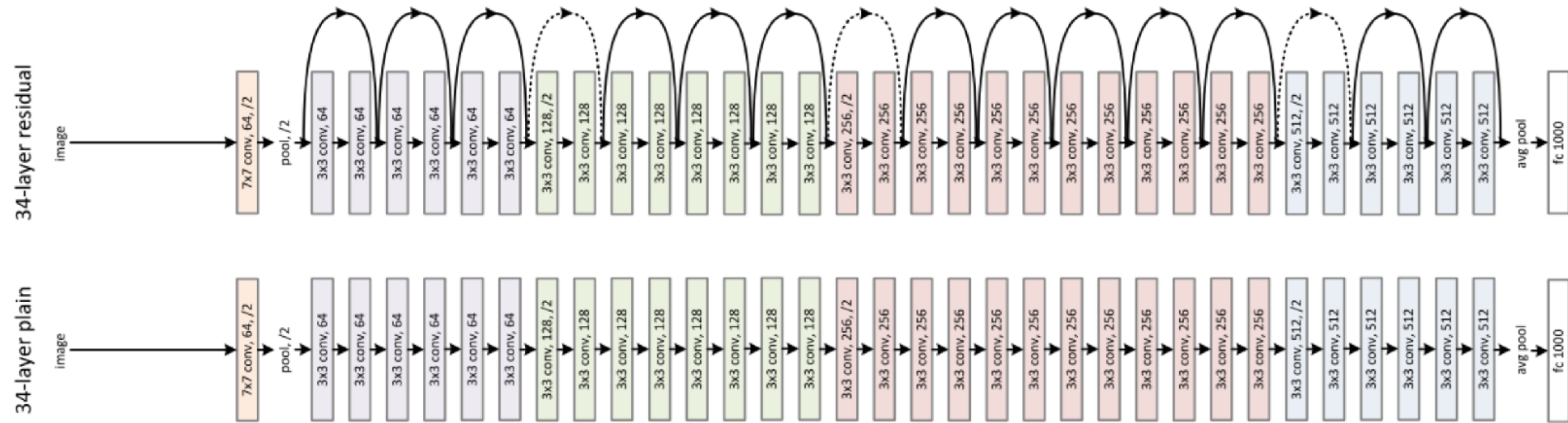
ResNet

ResNet: deep ConvNet with skip connections

[Li et al, NIPS, 2018] <https://arxiv.org/pdf/1712.09913.pdf>

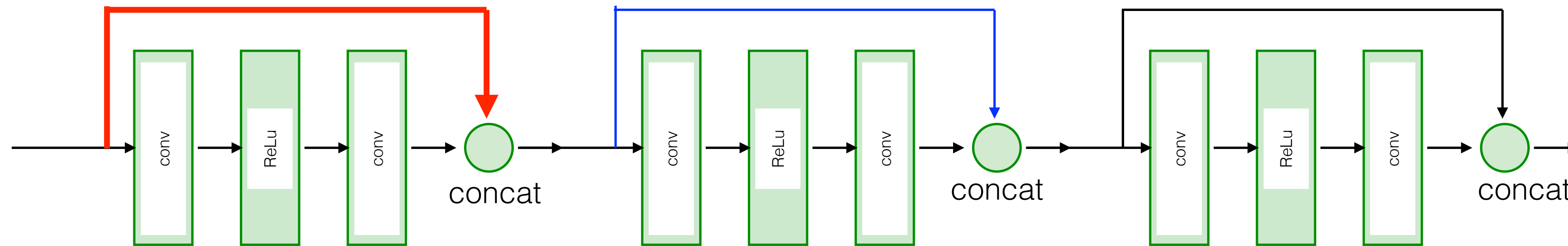


ResNet: deep ConvNet with skip connections



- **1k+ layers** possible
- **Initialization** with **zero** weights is meaningful
- **Better gradient flow:** many independent paths, opt-friendly landscape
- **Robustness** wrt **noise** and layer removal (if some important feature is not detected in a particular layer it can be detected in following layers)

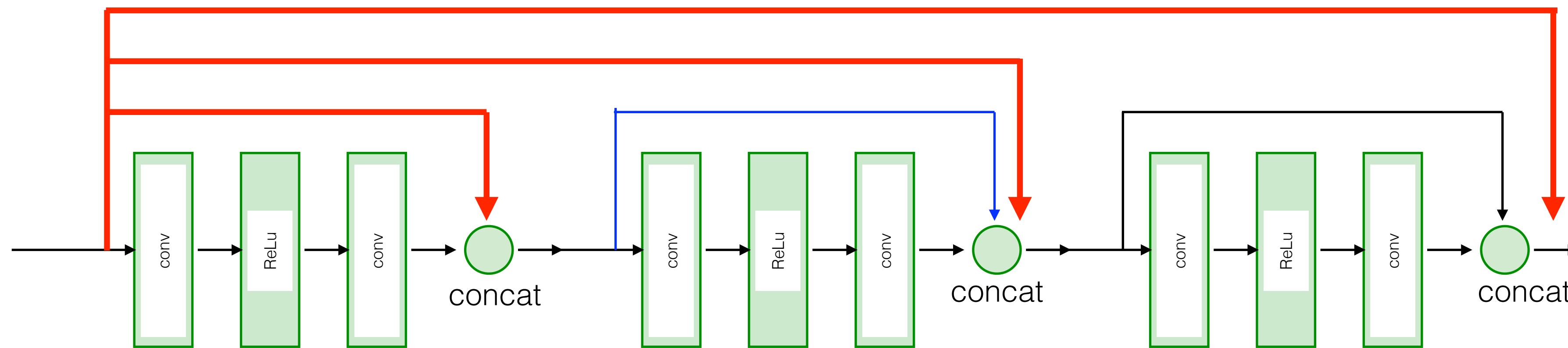
ResNet=>DenseNet



- Directly propagate each feature map to all following layers

Huang, Densely Connected Convolutional Networks, CVPR 2017. <https://arxiv.org/abs/1608.06993>

DenseNet



- Directly propagate each feature map to all following layers

Huang, Densely Connected Convolutional Networks, CVPR 2017.
<https://arxiv.org/abs/1608.06993>

Classification results

AlexNet

8 layers

VGGnet

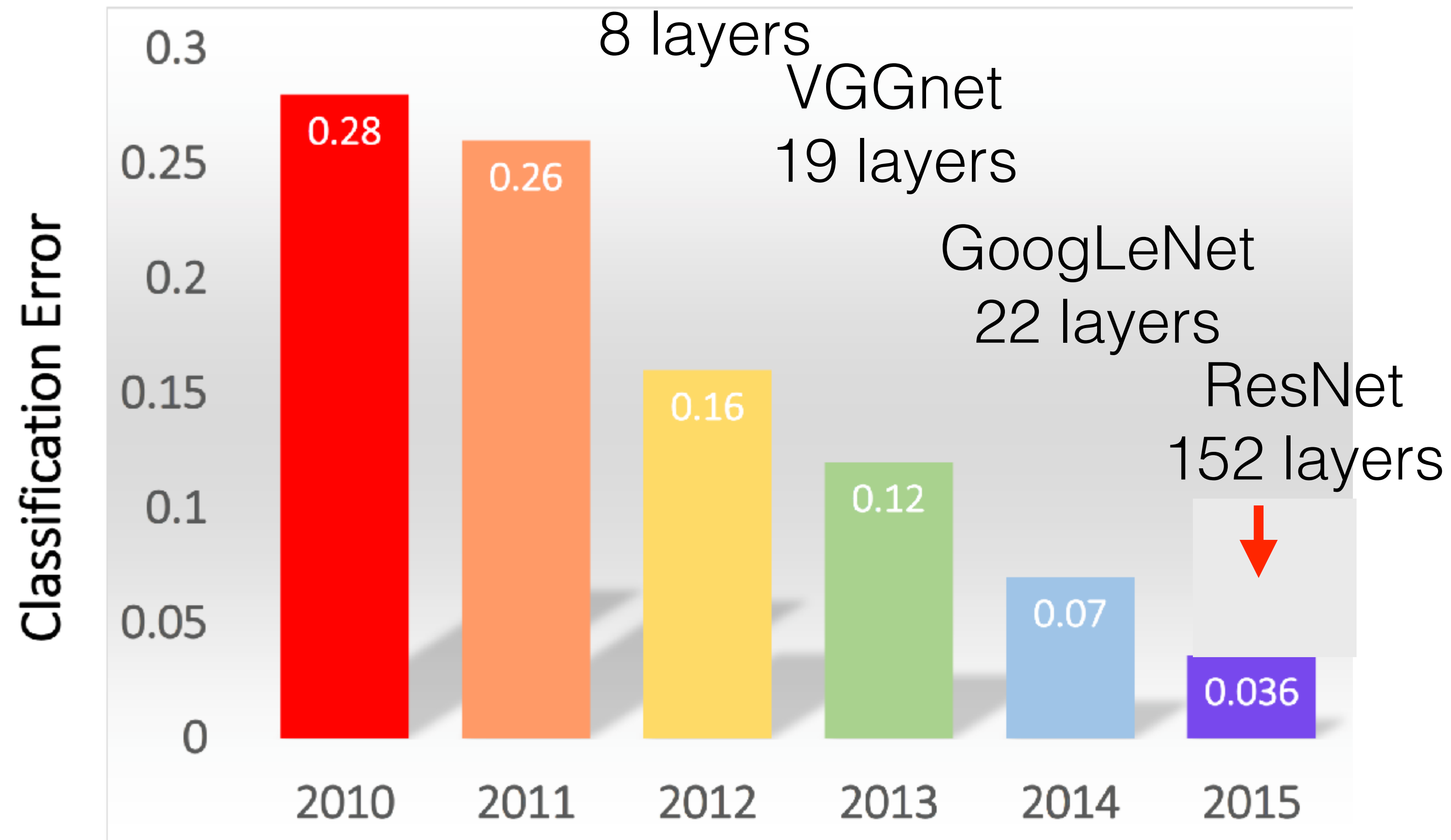
19 layers

GoogLeNet

22 layers

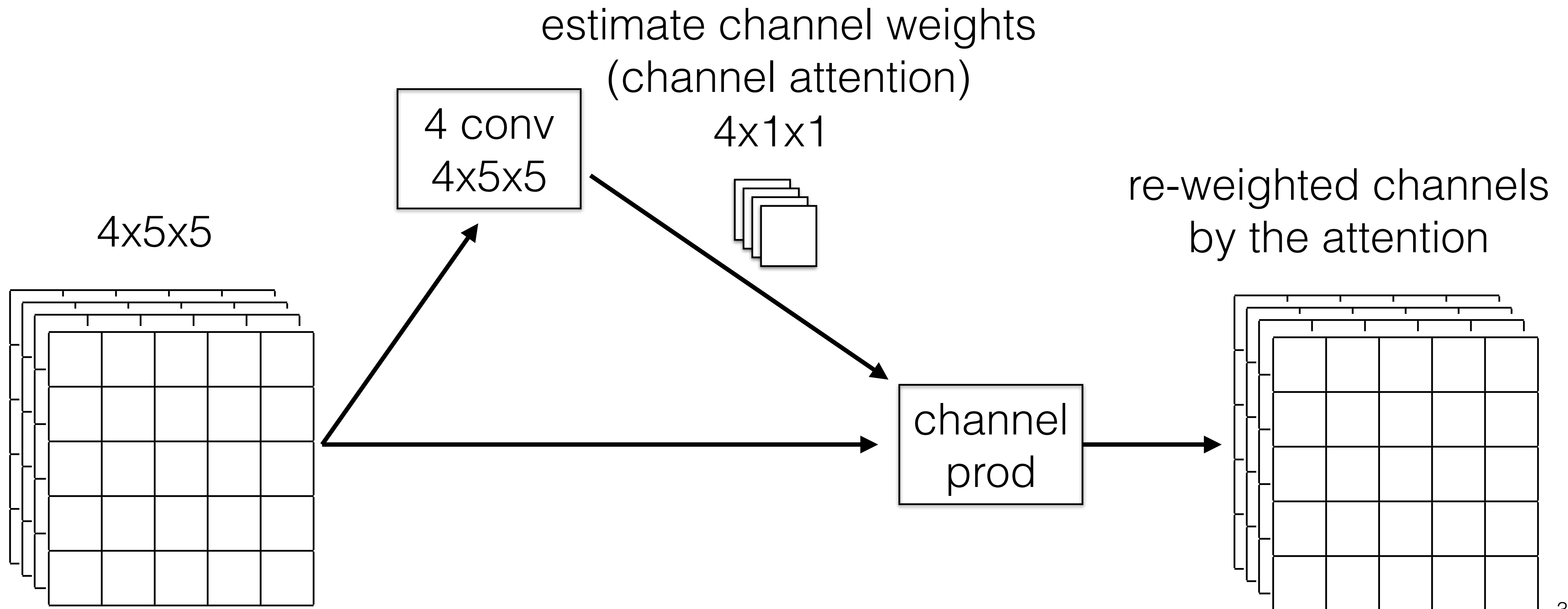
ResNet

152 layers



Human error around 5%

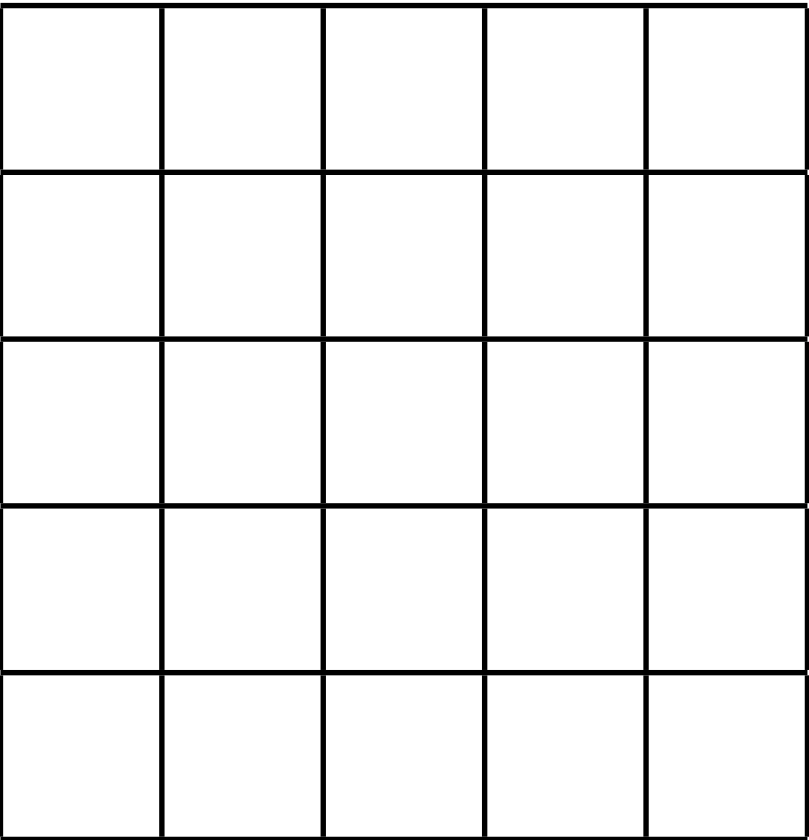
Channel attention



Spatial attention

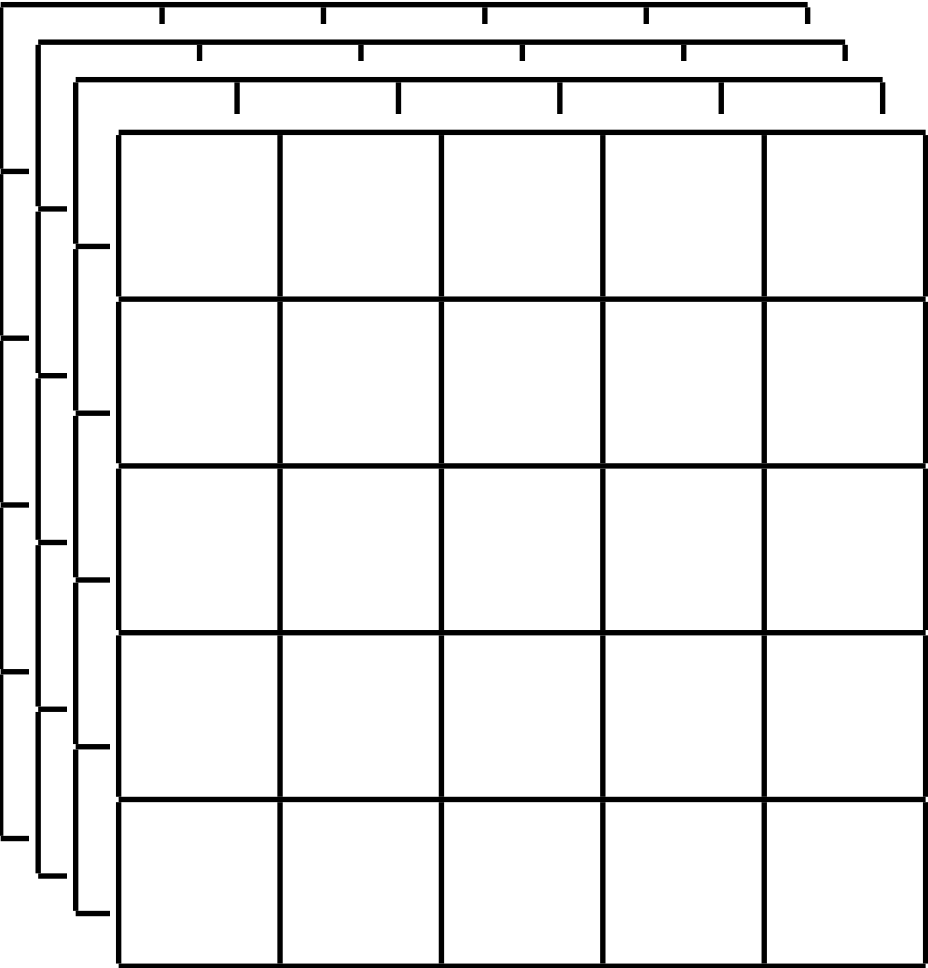
1x5x5

estimate pixel weights
(spatial attention)



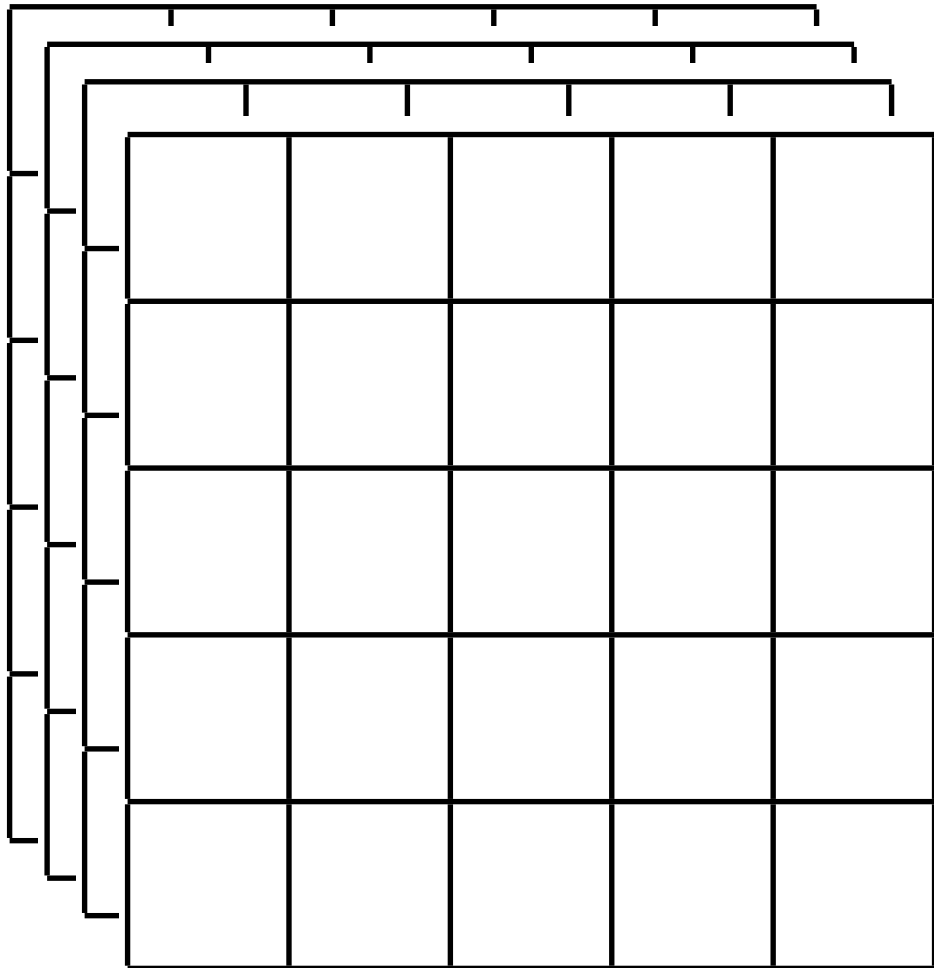
1 conv
4x1x1

4x5x5

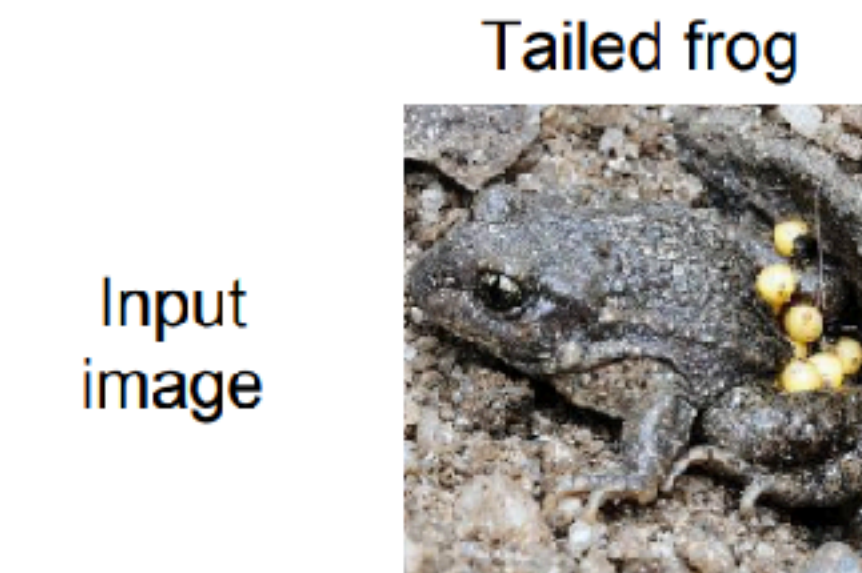


re-weighted pixels
by spatial attention

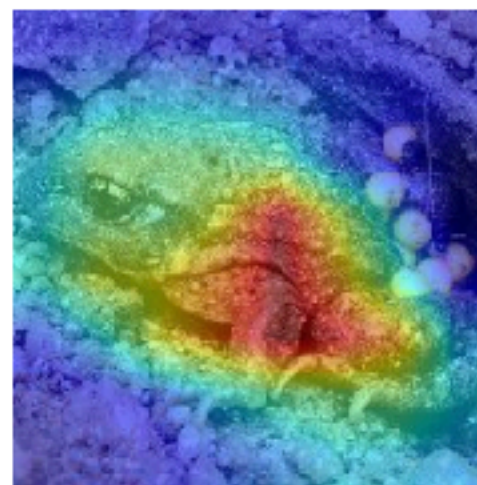
pixel
prod



The attention map of classified object is better localized

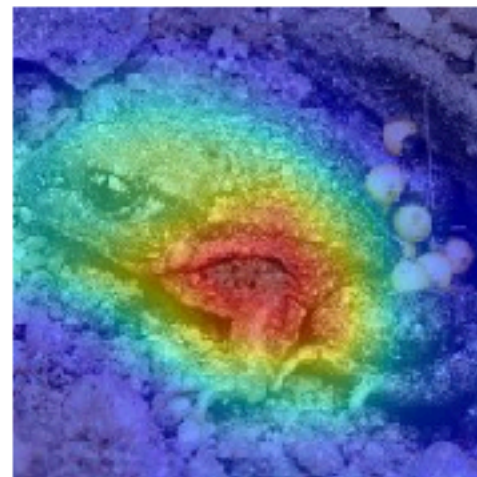


ResNet50



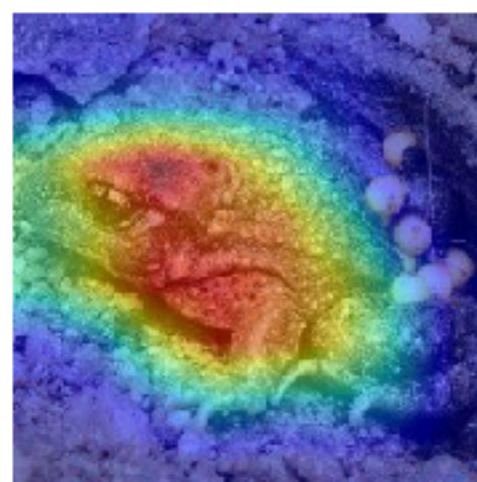
P=0.80736

ResNet50
+ SE



P=0.87240

ResNet50
+ CBAM



P = 0.96340

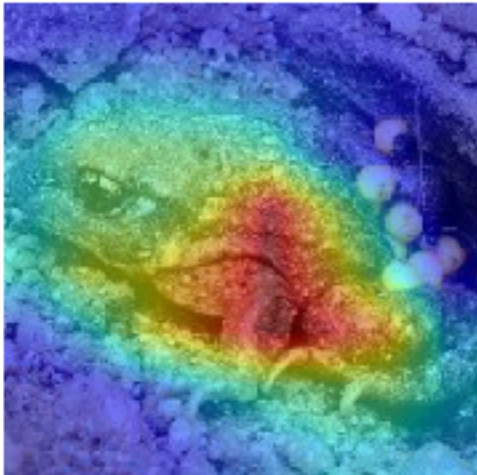
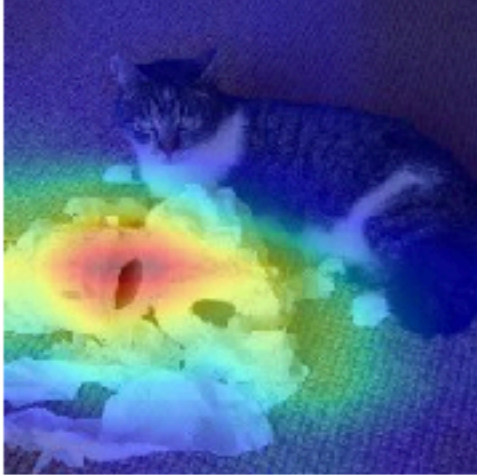
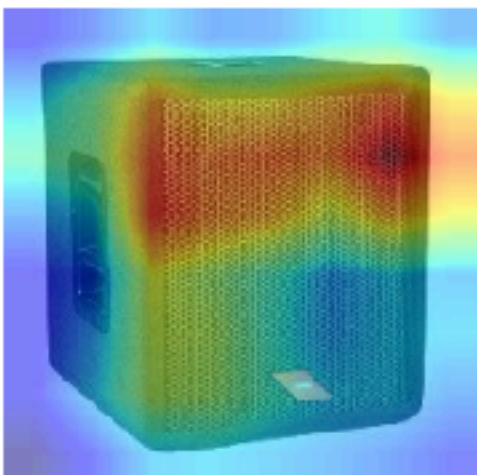
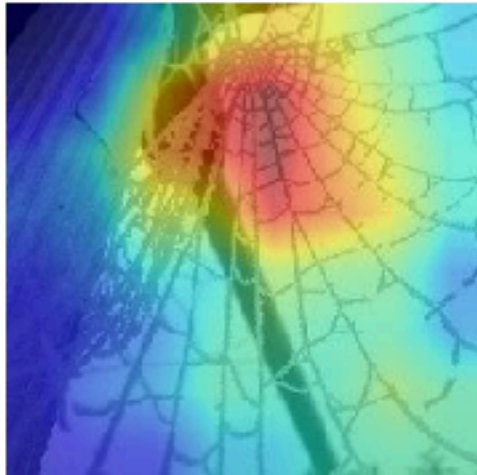
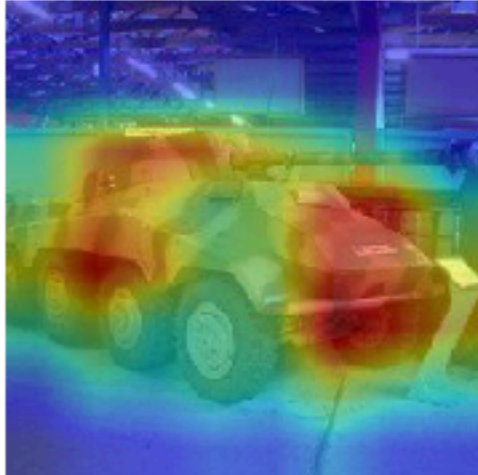
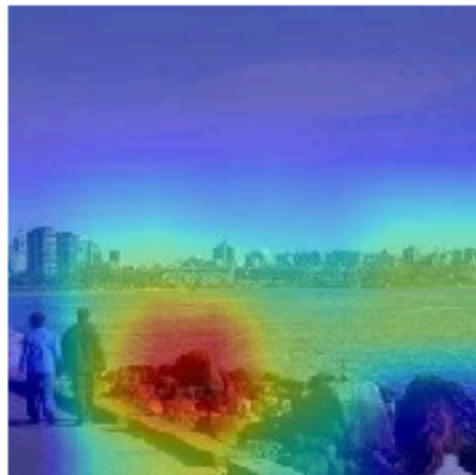
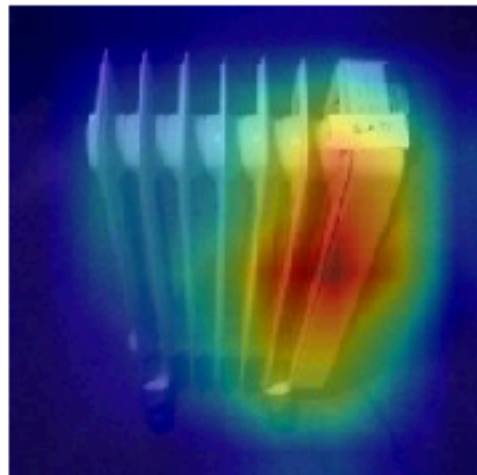
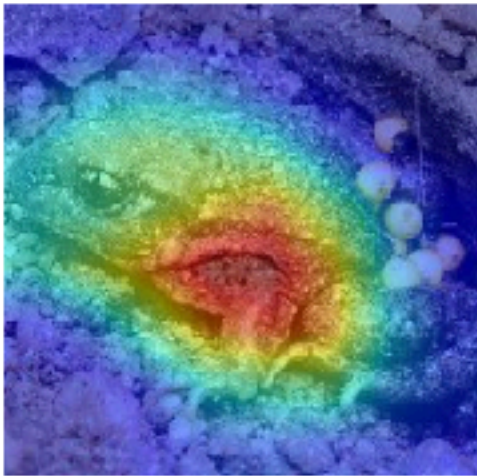
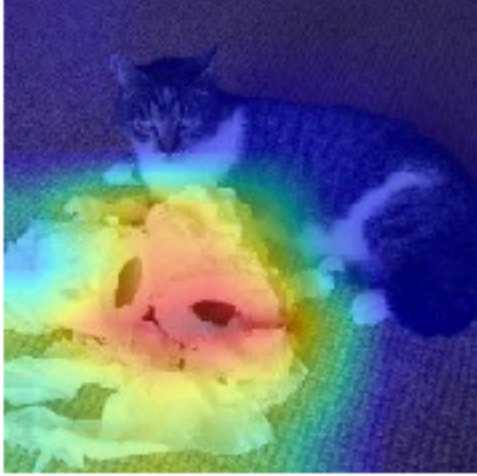
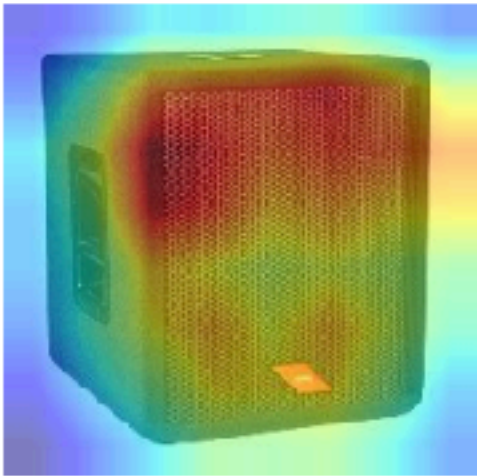
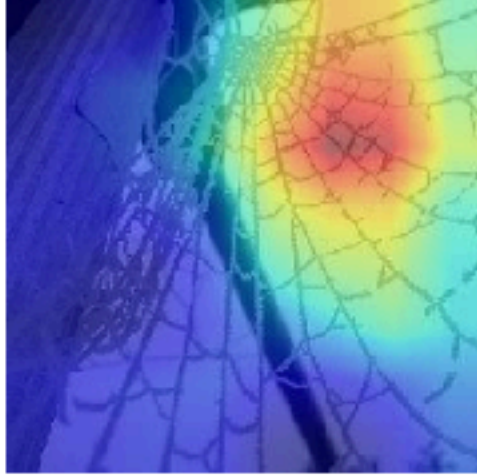
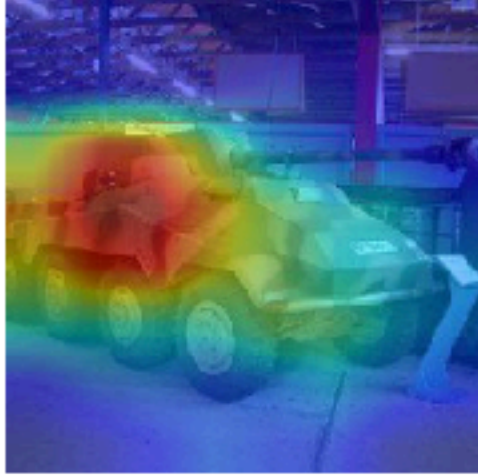
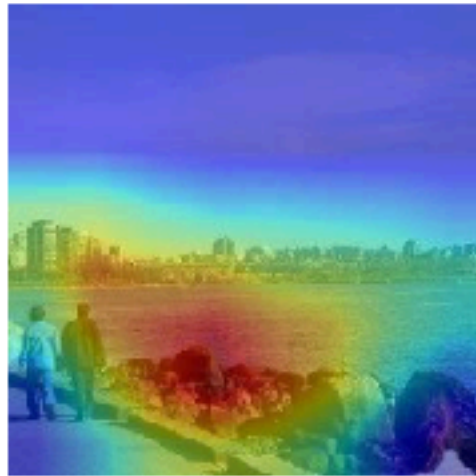
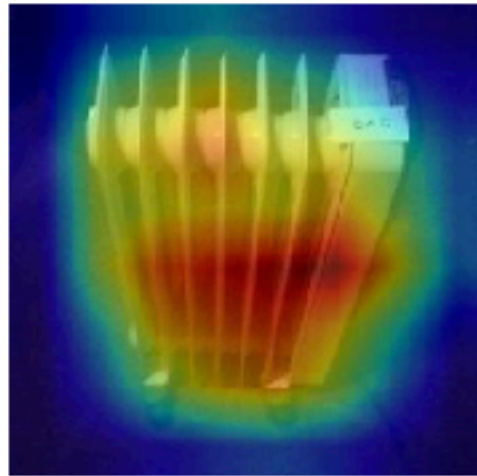
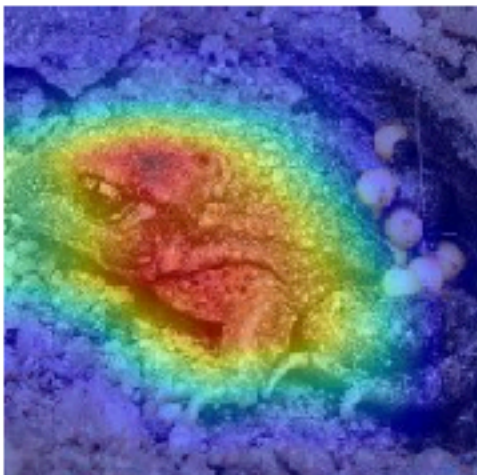
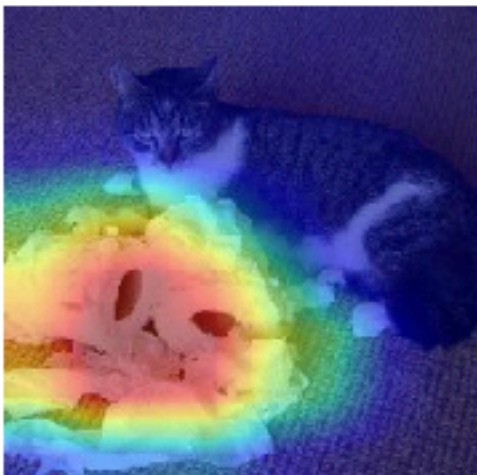
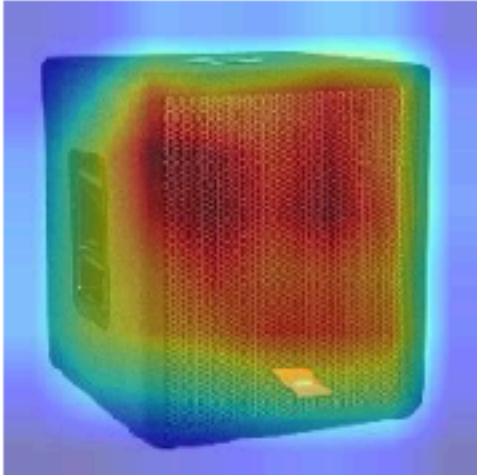
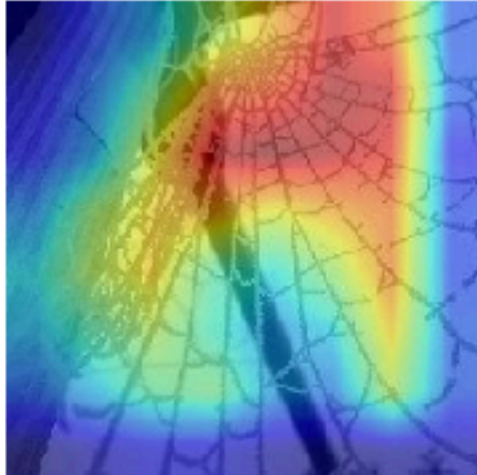
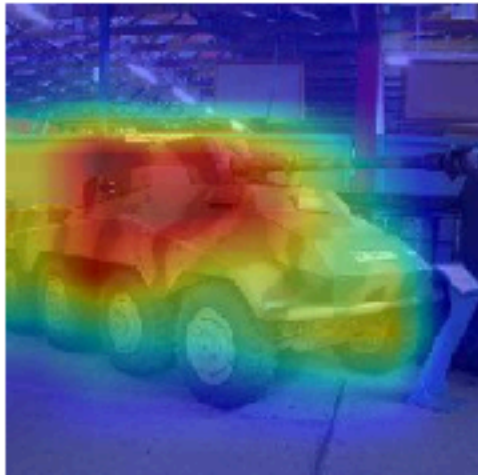
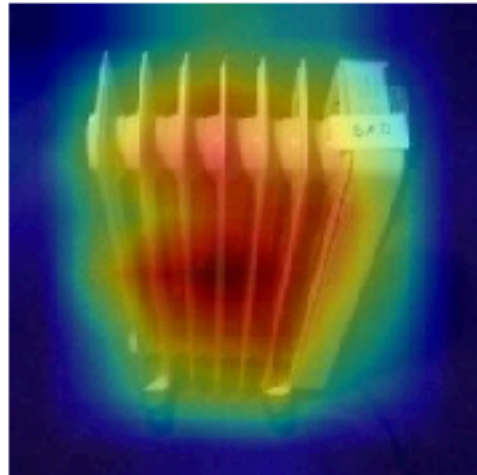
w/o attention

with channel attention

with spatial attention

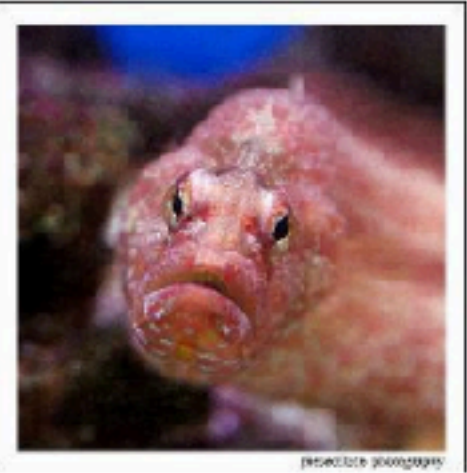
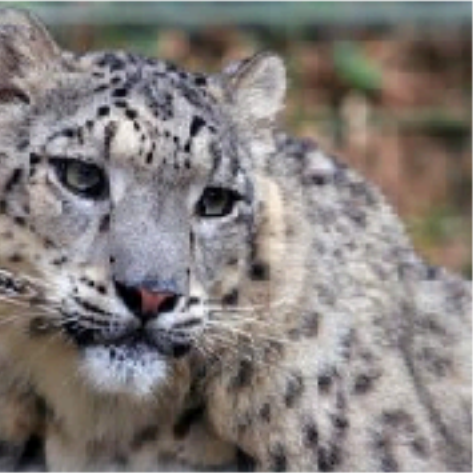
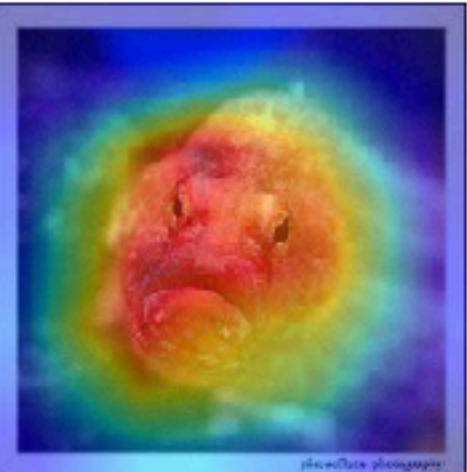
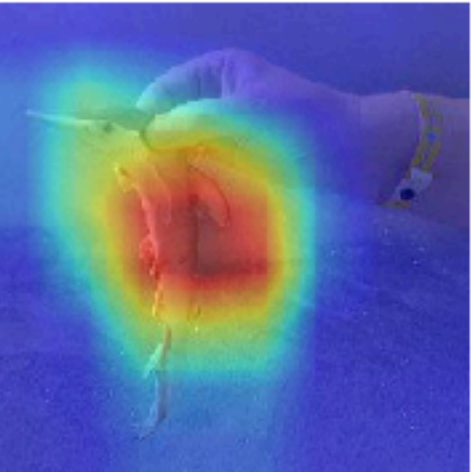
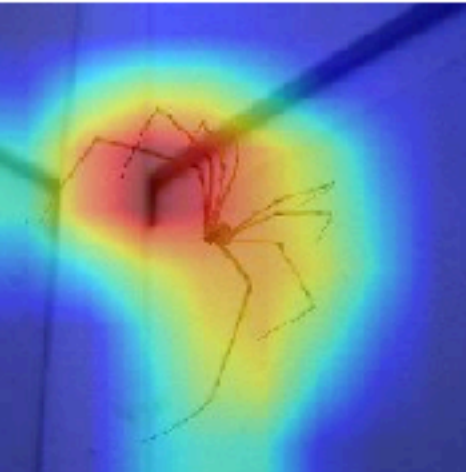
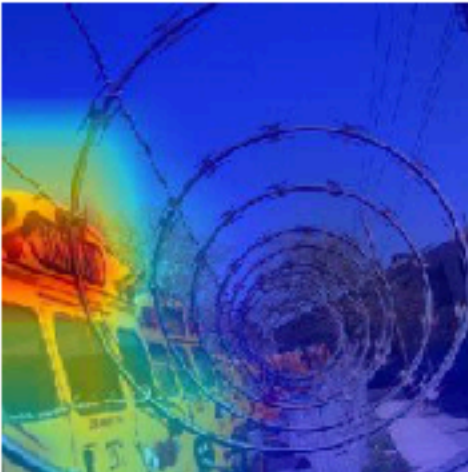
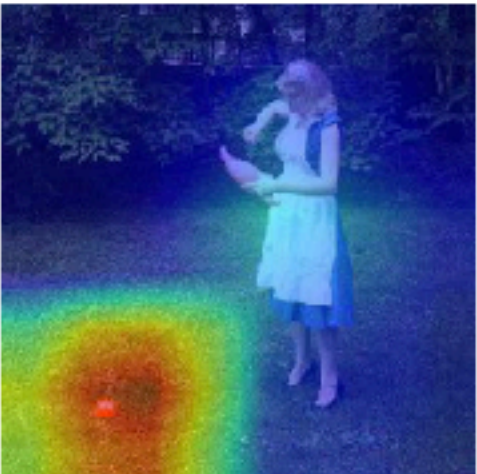
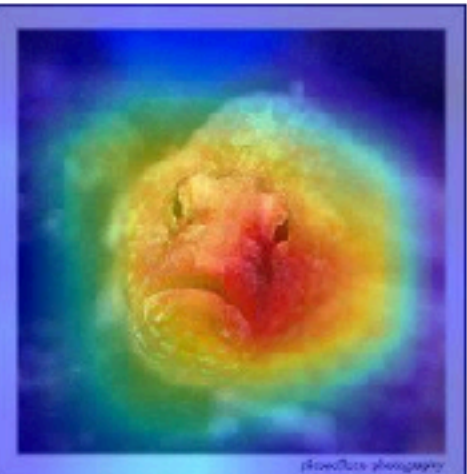
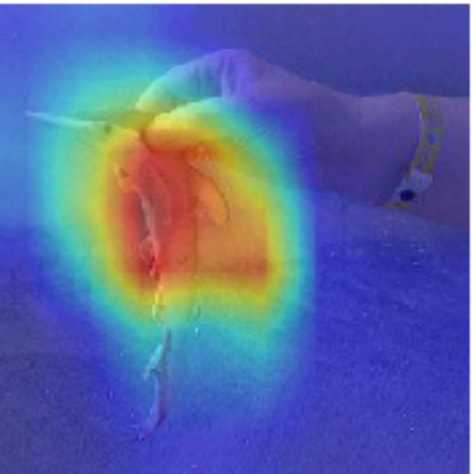
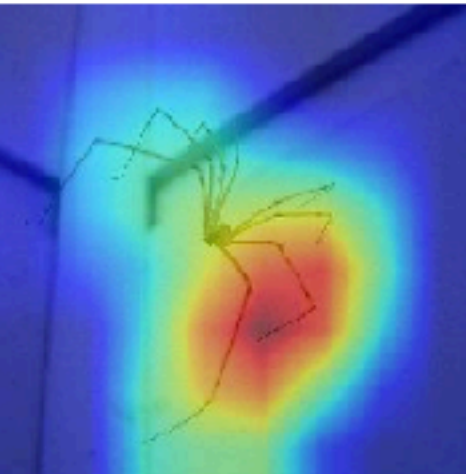
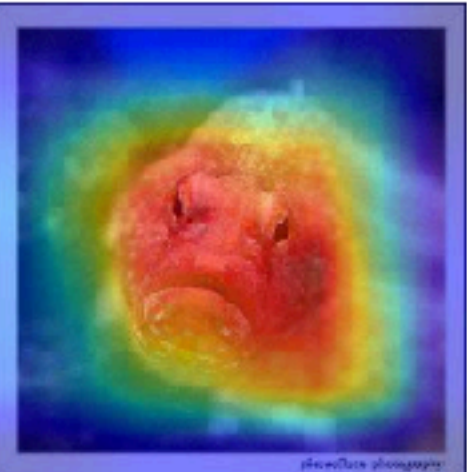
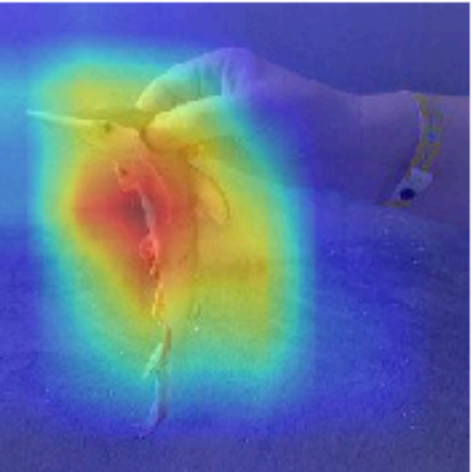
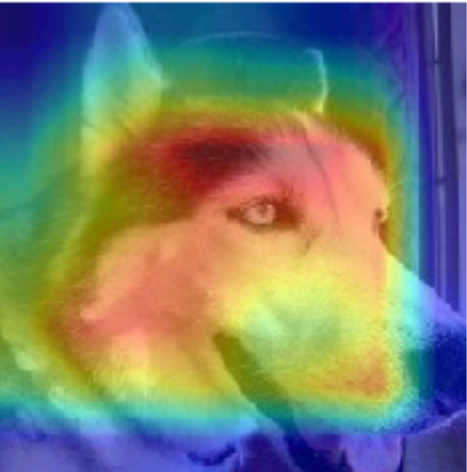
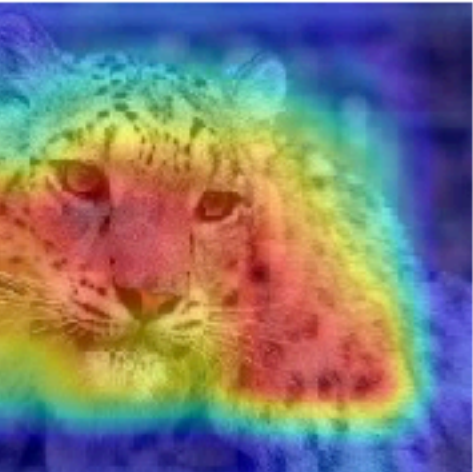
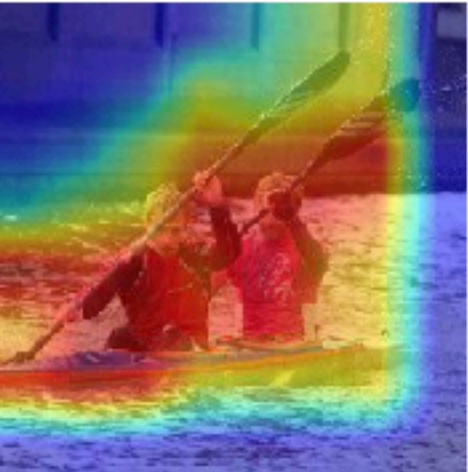
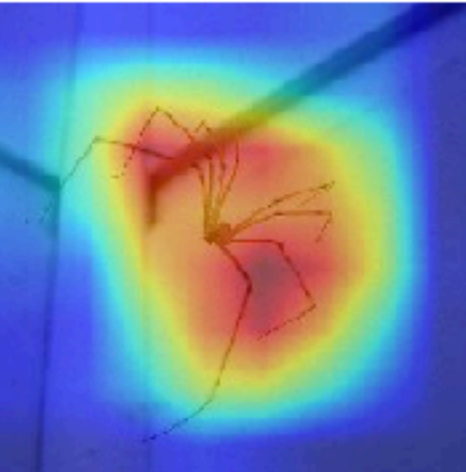
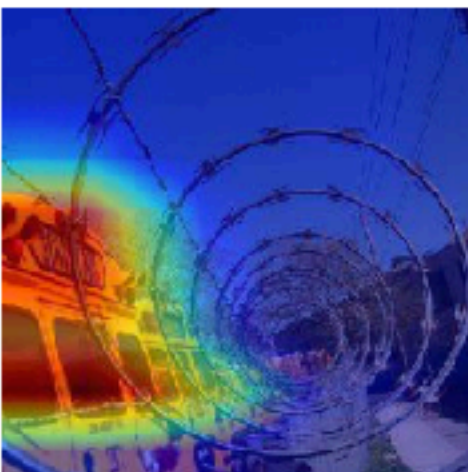
Attention modules [Woo et al, ECCV, 2018]

<https://arxiv.org/pdf/1807.06521v2.pdf>

	Tailed frog	Toilet tissue	Loudspeaker	Spider web	American egret	Tank	Seawall	Space heater
Input image								
ResNet50								
	P=0.80736	P=0.11857	P=0.65681	P=0.22357	P=0.64185	P=0.14763	P=0.92236	P=0.01176
ResNet50 + SE								
	P=0.87240	P=0.14643	P=0.77550	P=0.25093	P=0.70827	P=0.15367	P=0.97166	P=0.26611
ResNet50 + CBAM								
	P = 0.96340	P = 0.19994	P = 0.93707	P = 0.35248	P = 0.87490	P = 0.53005	P = 0.99085	P = 0.59662

Attention modules [Woo et al, ECCV, 2018]

<https://arxiv.org/pdf/1807.06521v2.pdf>

	Croquet ball	Eel	Hammerhead	Eskimo dog	Snow leopard	Boat paddle	Daddy longlegs	School bus
Input image								
ResNet50								
	P=0.58732	P=0.08126	P=0.67128	P=0.56834	P=0.85873	P=0.62856	P=0.68065	P=0.07429
ResNet50 + SE								
	P=0.89962	P=0.14804	P=0.72659	P=0.61595	P=0.96575	P=0.79829	P=0.73723	P=0.92250
ResNet50 + CBAM								
	P = 0.96039	P = 0.59790	P = 0.84387	P = 0.71000	P = 0.98482	P = 0.90806	P = 0.78636	P = 0.98567

GradCAM [Selvaraju et al, ICCV, 2018]



Classification results

AlexNet

8 layers

VGGnet

19 layers

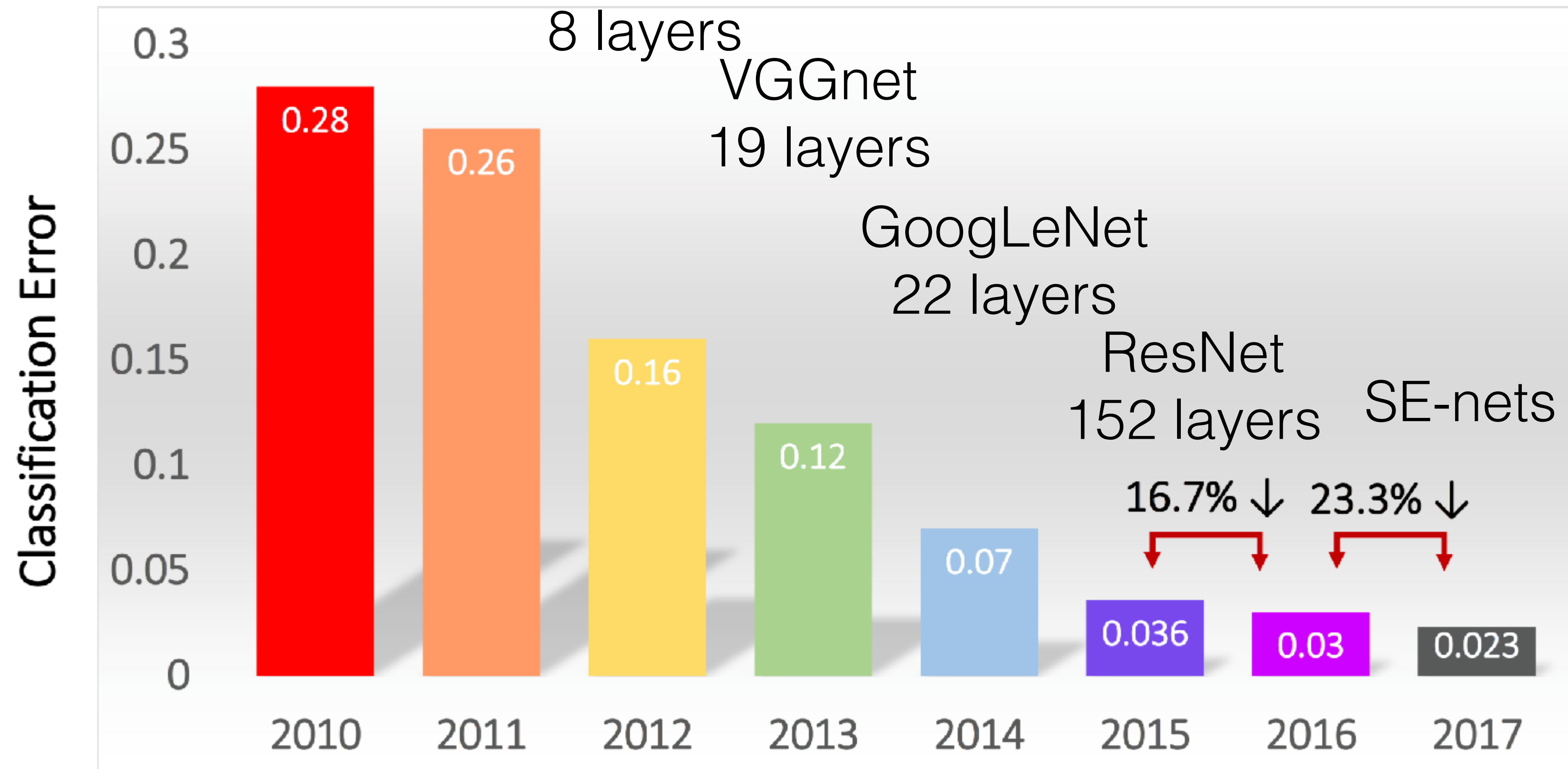
GoogLeNet

22 layers

ResNet

152 layers

SE-nets



Summary classification

- Injecting gradient
- Skip connections
- Receptive field (“one 7x7” vs “three 3x3 convs”)
- Spatial and Channel Attention

Outline

- Architectures of classification networks
- Architectures of segmentation networks
- Architectures of regression networks
- Architectures of detection networks
- Architectures of regression networks
- Architectures of feature matching networks

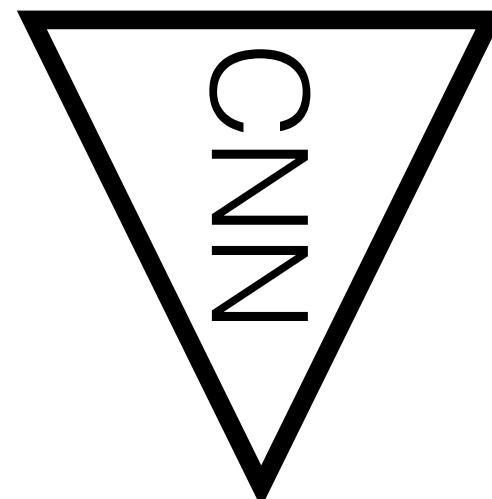
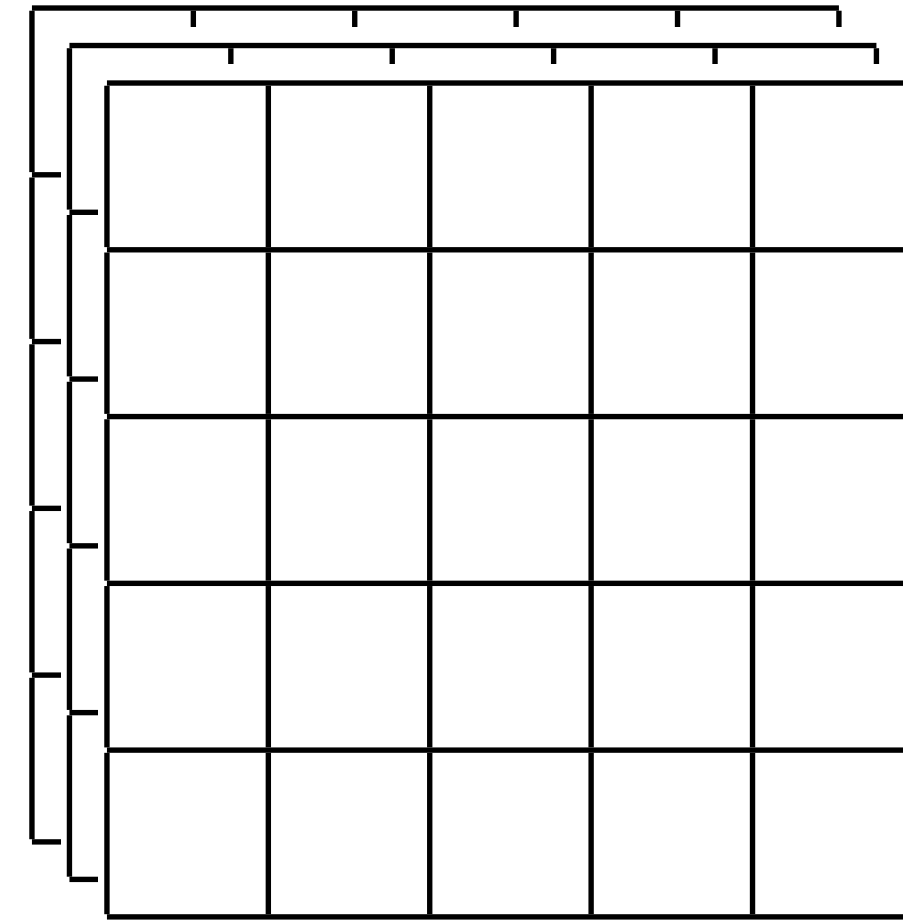
Semantic segmentation



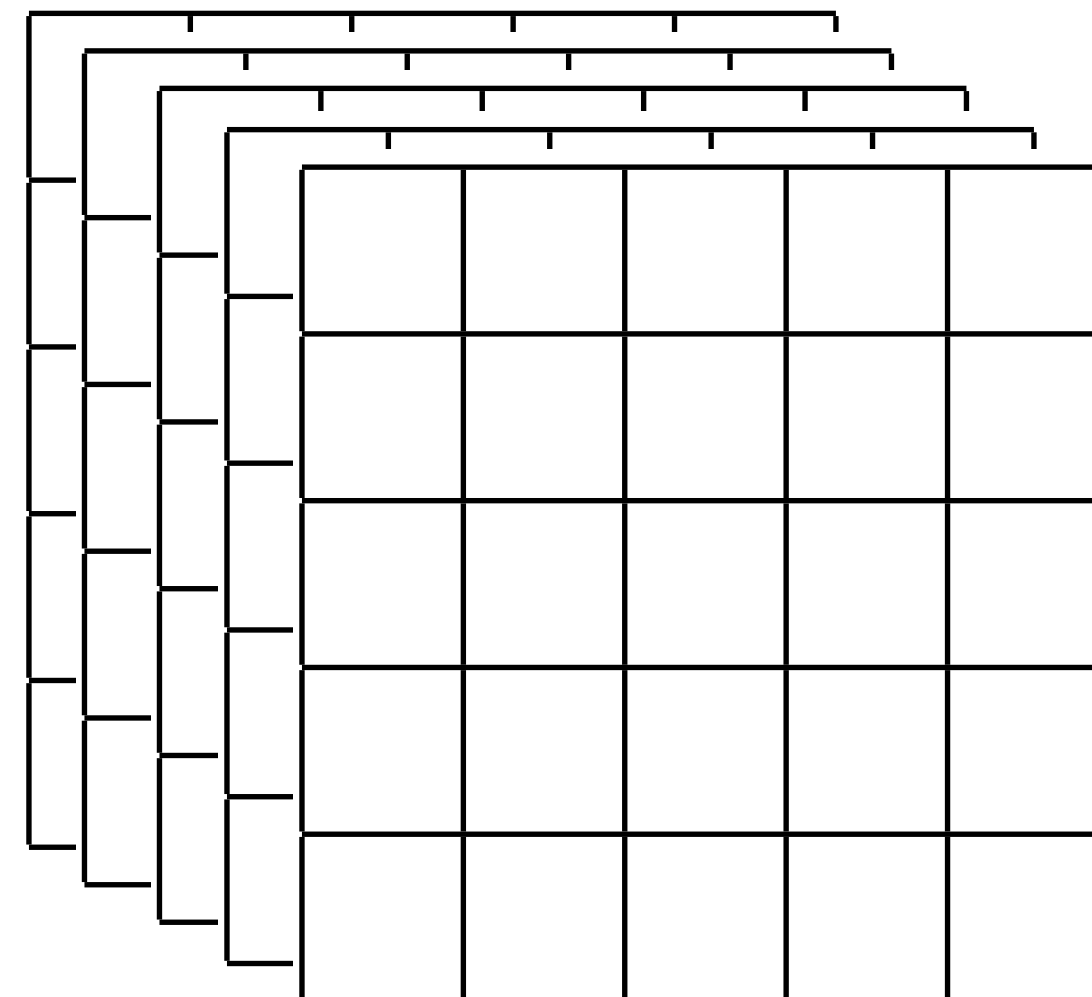
- road
- sidewalk
- pedestrian
- traffic sign
- trees
- sky

Semantic segmentation

RGB image
(HxWx3)



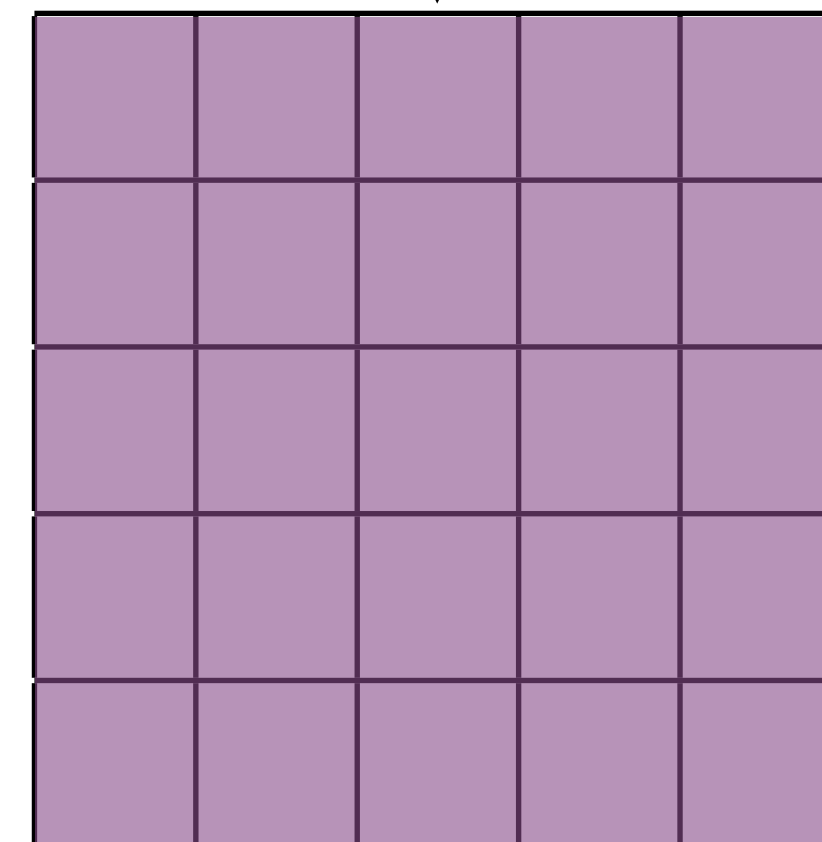
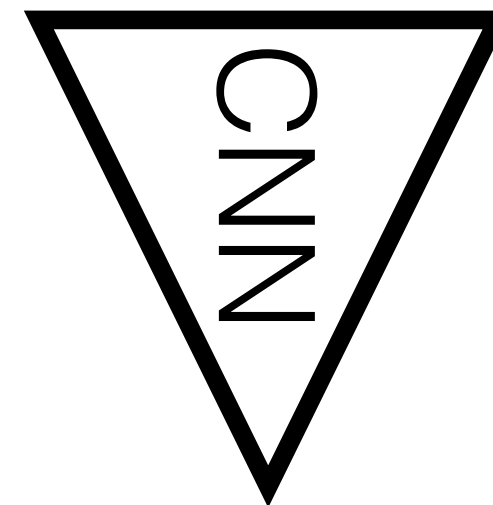
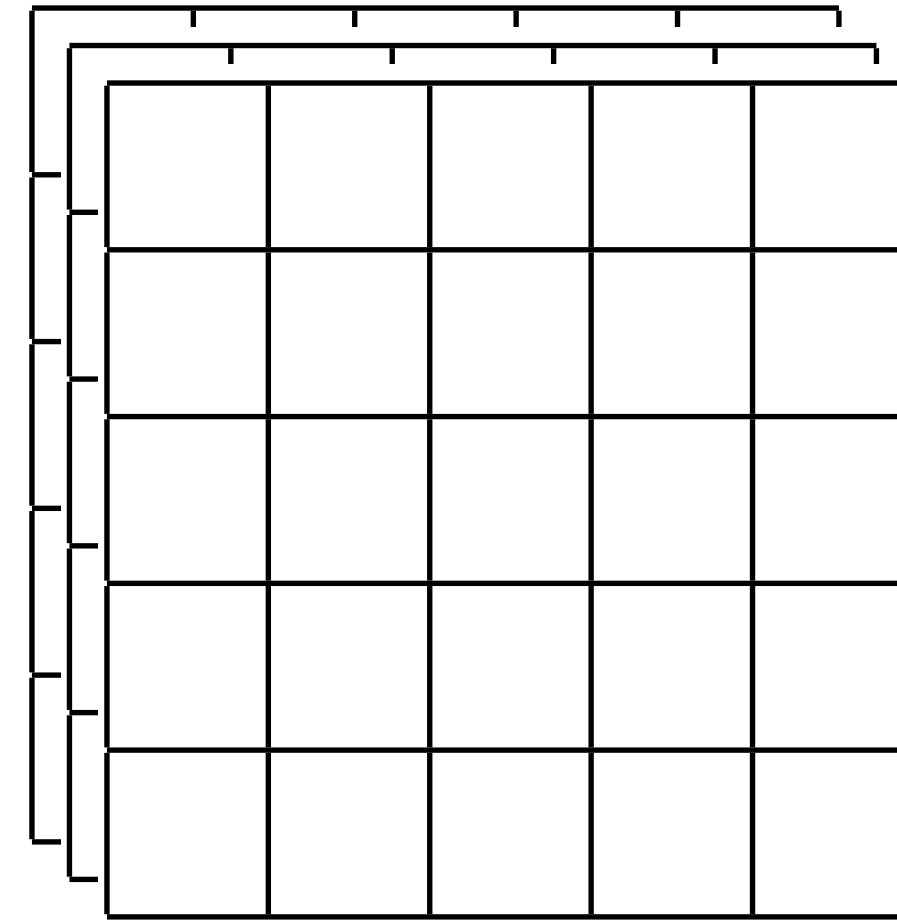
pixel-wise probs
(HxWxN)



- road
- sideway
- pedestrian
- traffic sign
- trees
- sky

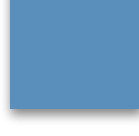
Semantic segmentation

RGB image
(HxWx3)



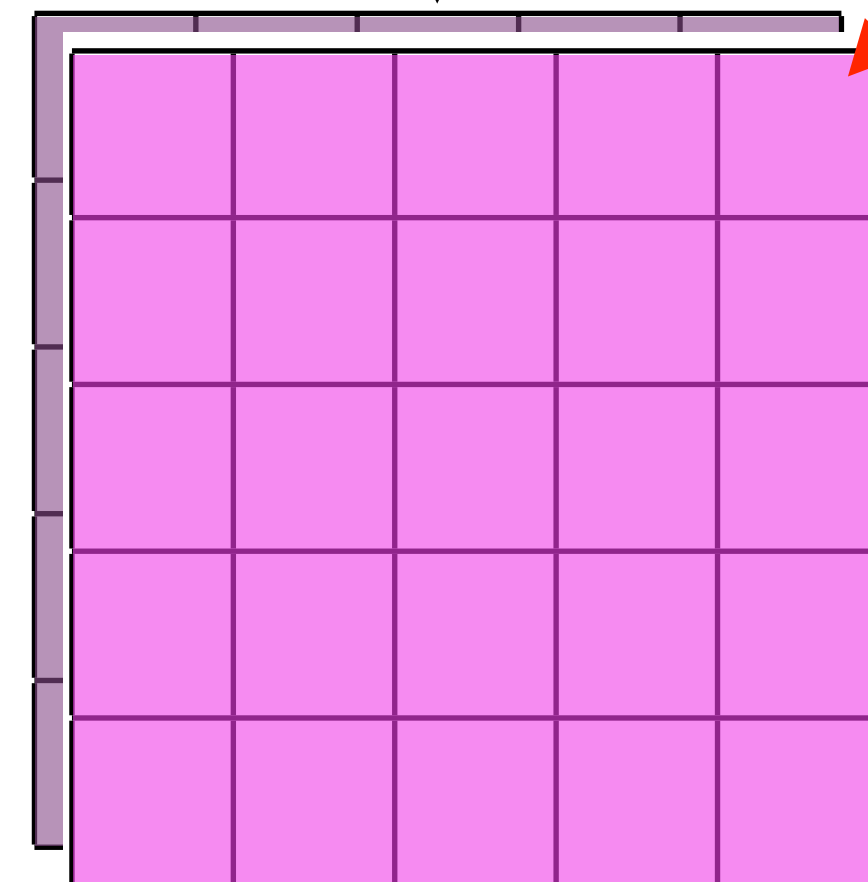
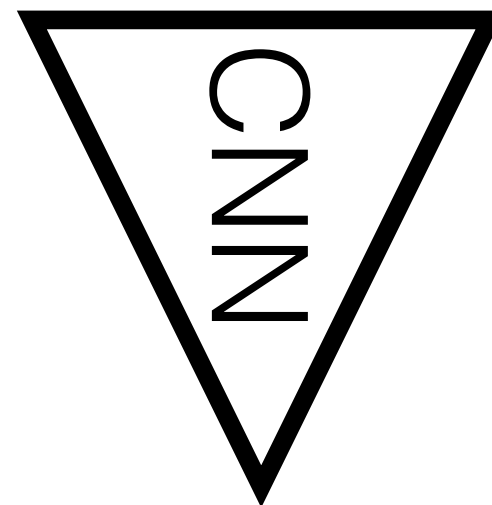
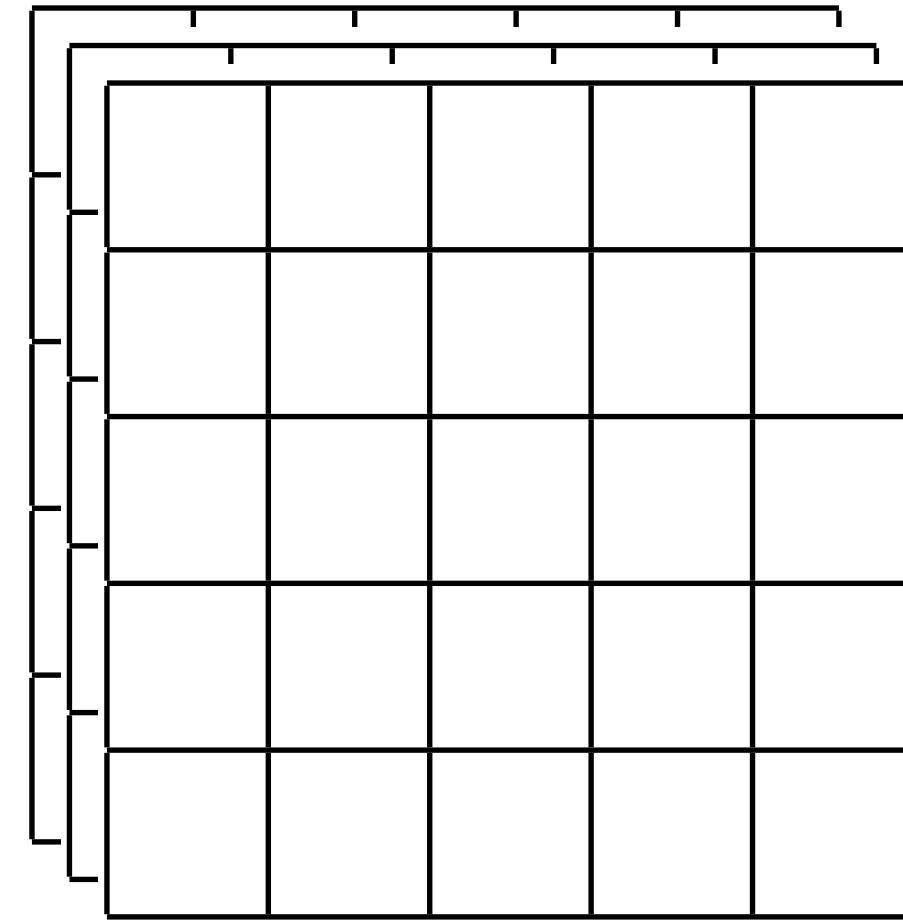
pixel-wise probability
of being **road**

channel 1

-  road
-  sidewalk
-  pedestrian
-  traffic sign
-  trees
-  sky

Semantic segmentation

RGB image
(HxWx3)



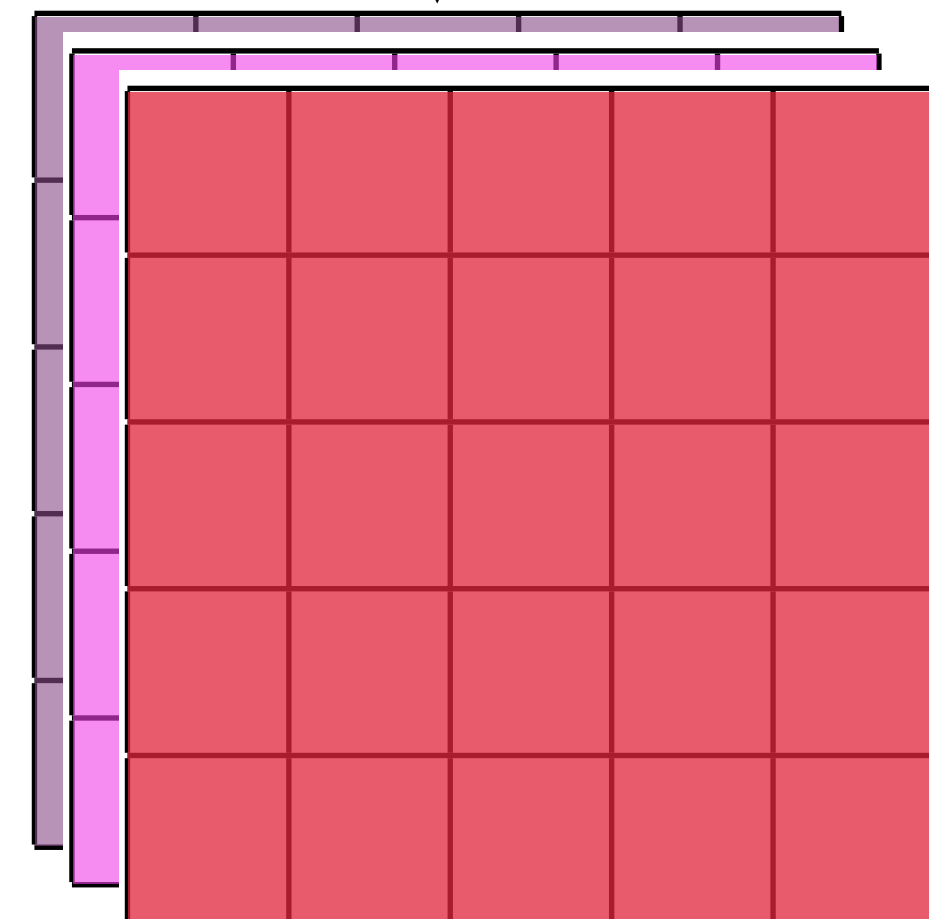
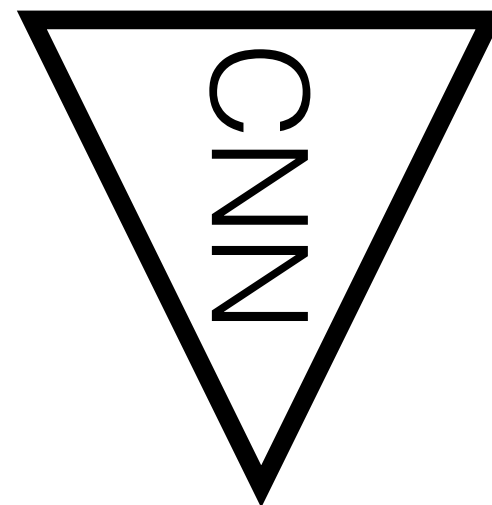
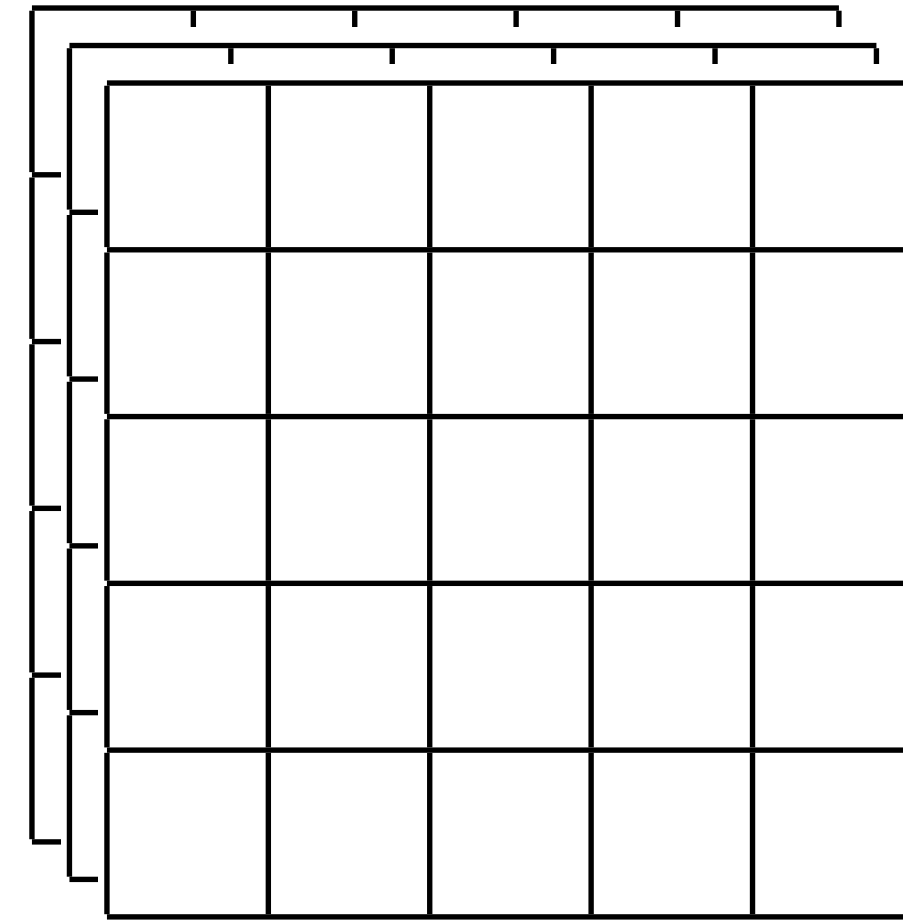
- road
- sideway
- pedestrian
- traffic sign
- trees
- sky

pixel-wise probability
of being **sideway**

channel 2

Semantic segmentation

RGB image
(HxWx3)



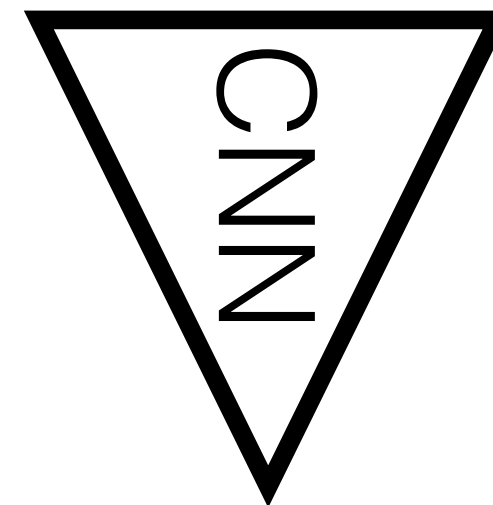
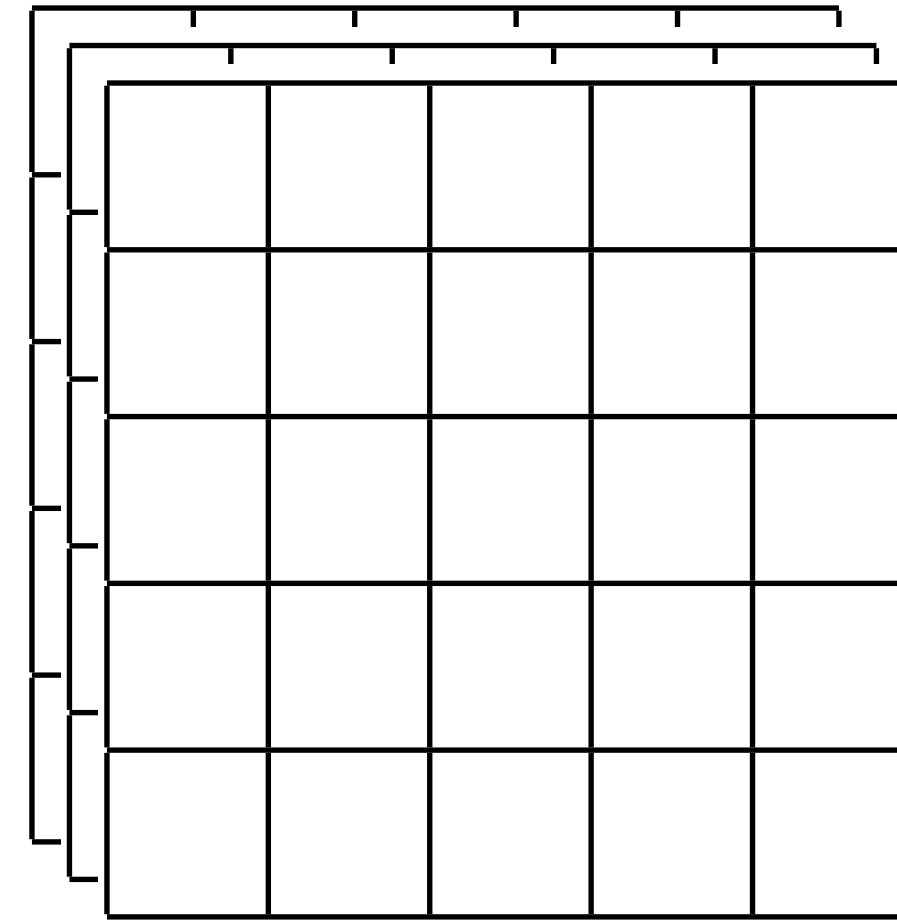
- road
- sideway
- pedestrian
- traffic sign
- trees
- sky

pixel-wise probability
of being **pedestrian**

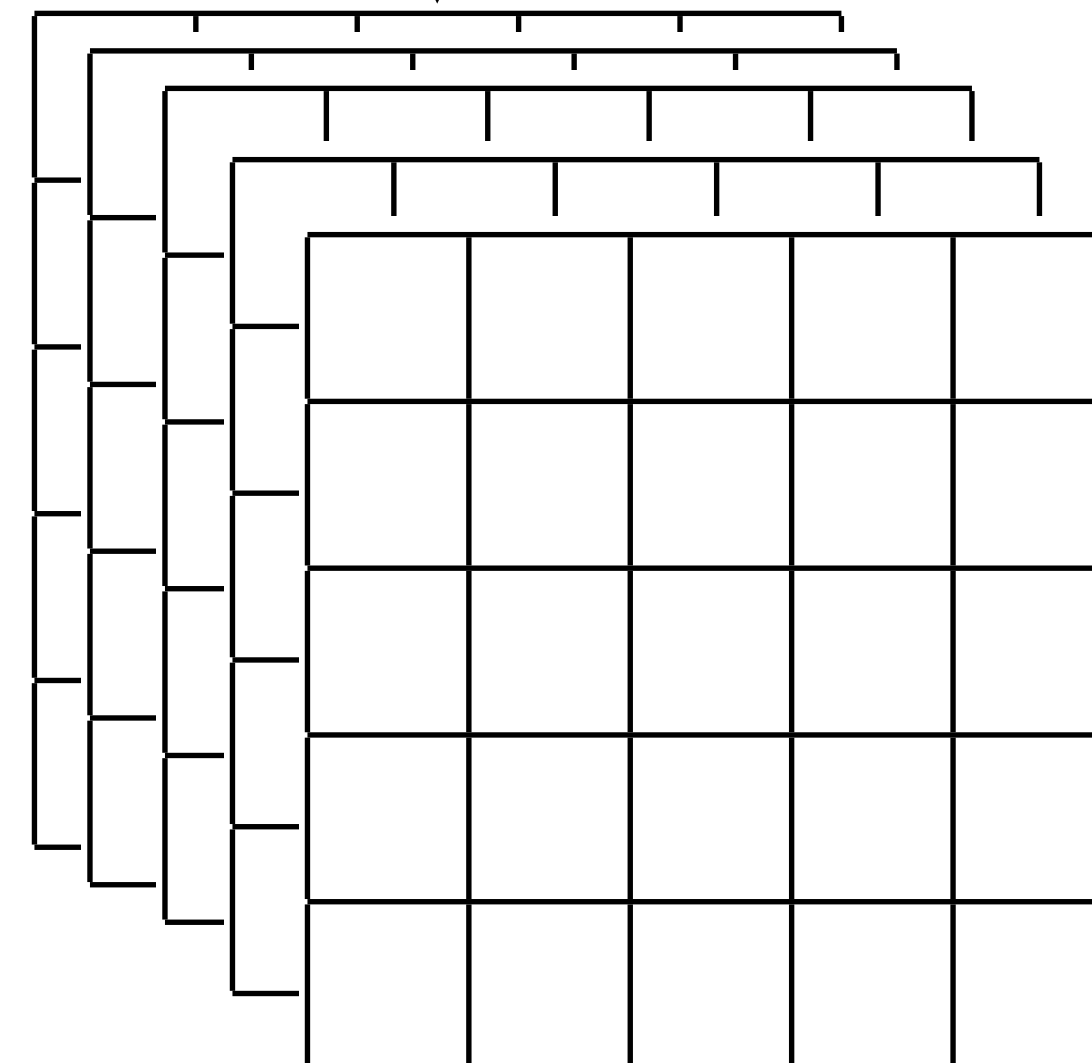
channel 3

Semantic segmentation

RGB image
(HxWx3)



pixel-wise probs
(HxWxN)

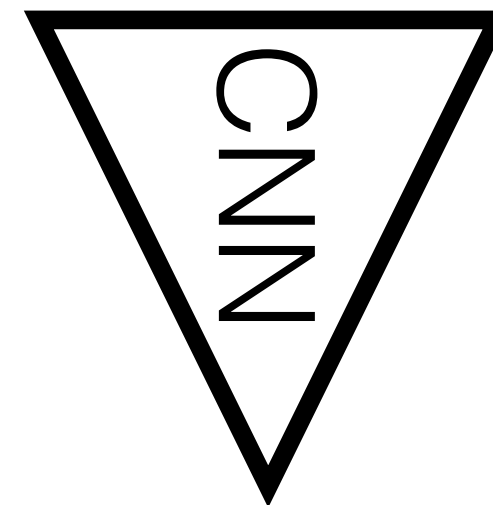
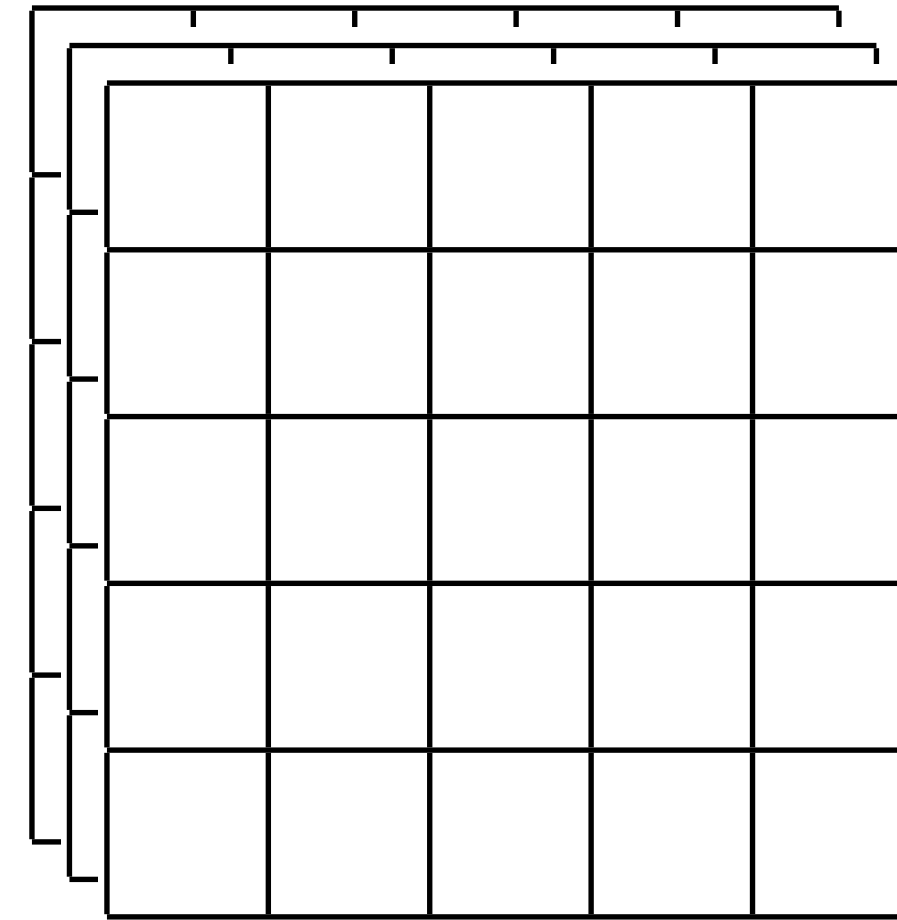


- road
- sideway
- pedestrian
- traffic sign
- trees
- sky

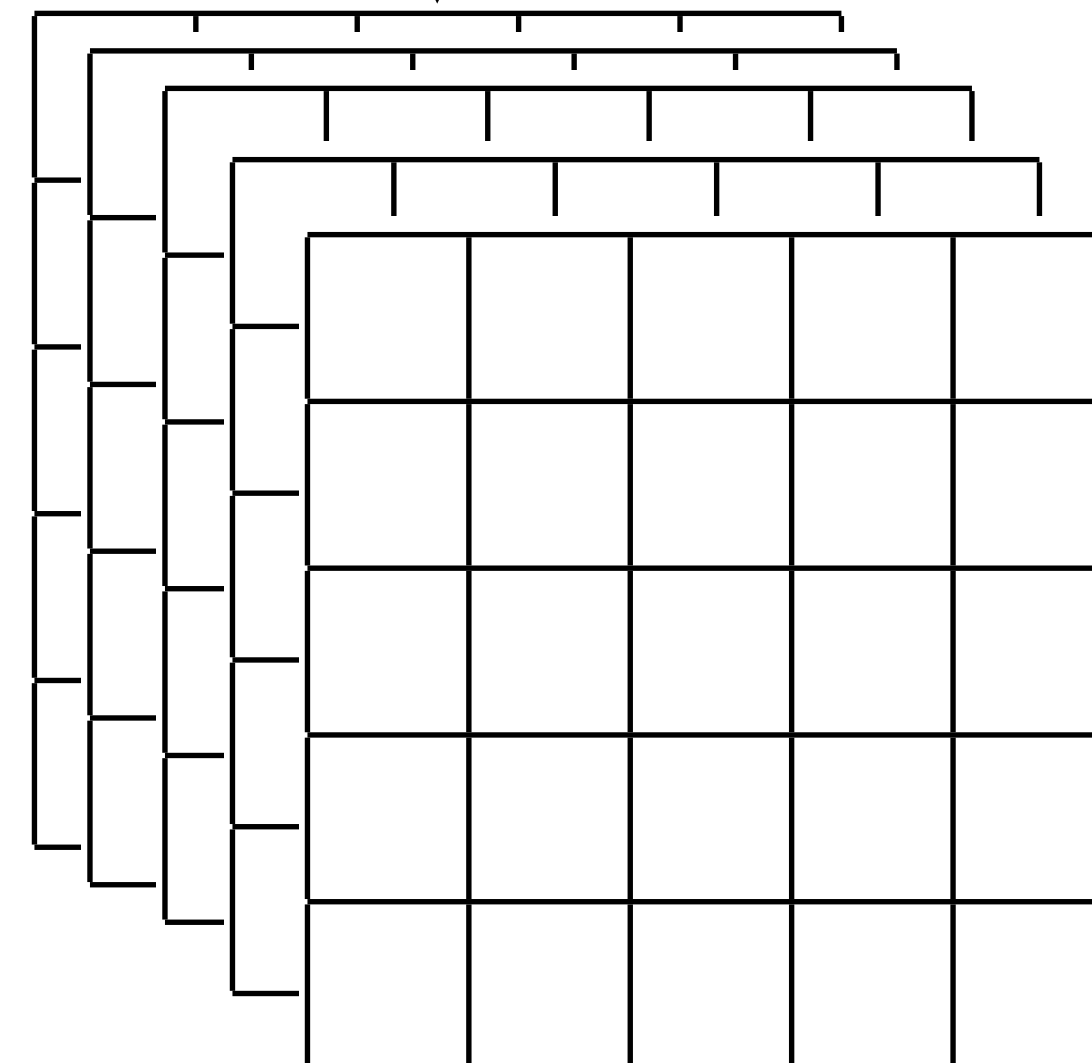
as many output channels
as semantic labels

Semantic segmentation

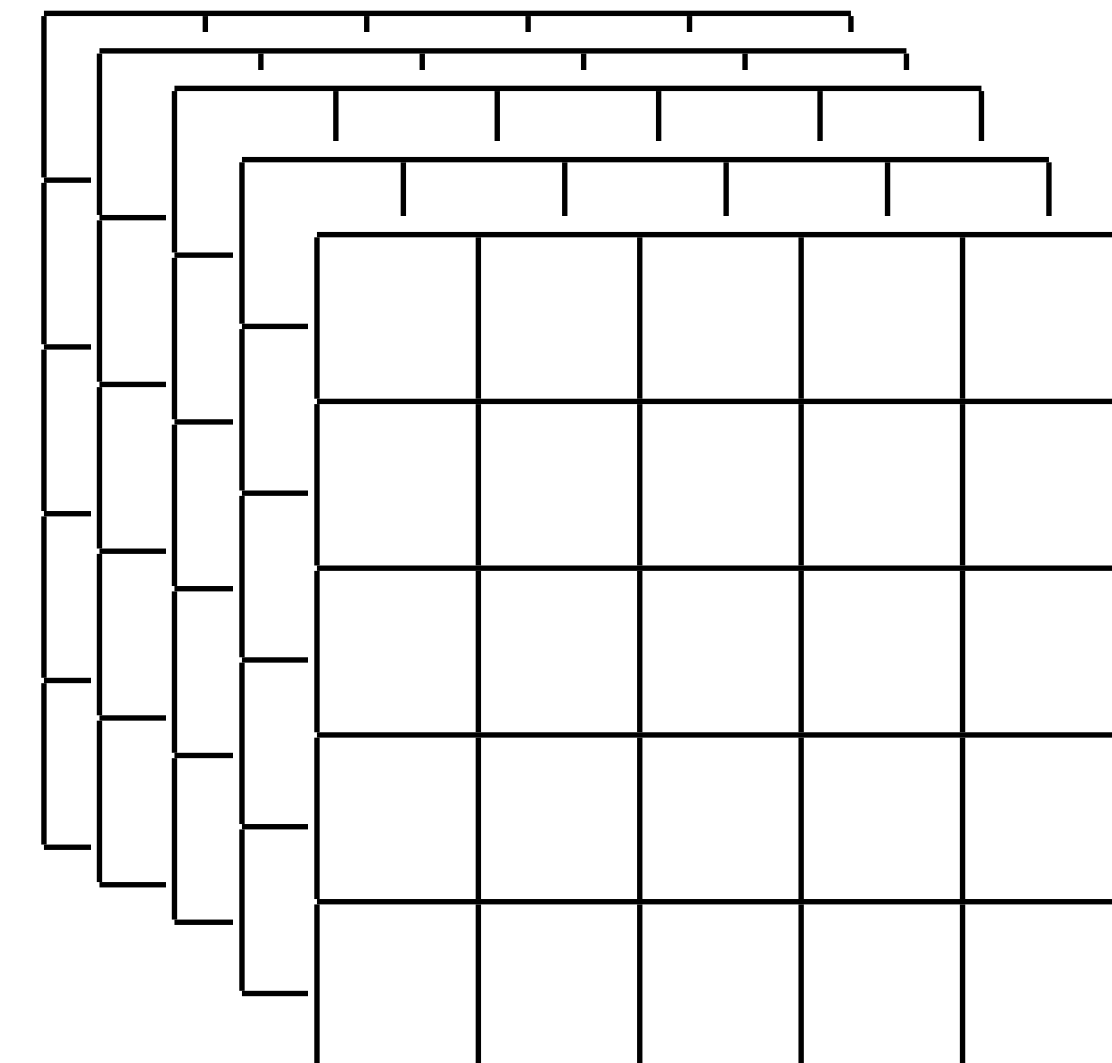
RGB image
($H \times W \times 3$)



pixel-wise probs
($H \times W \times N$)



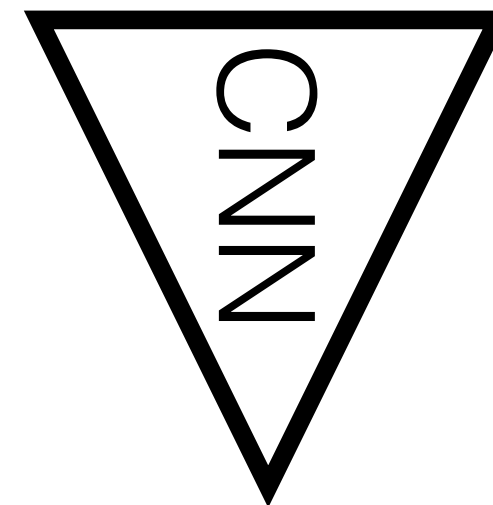
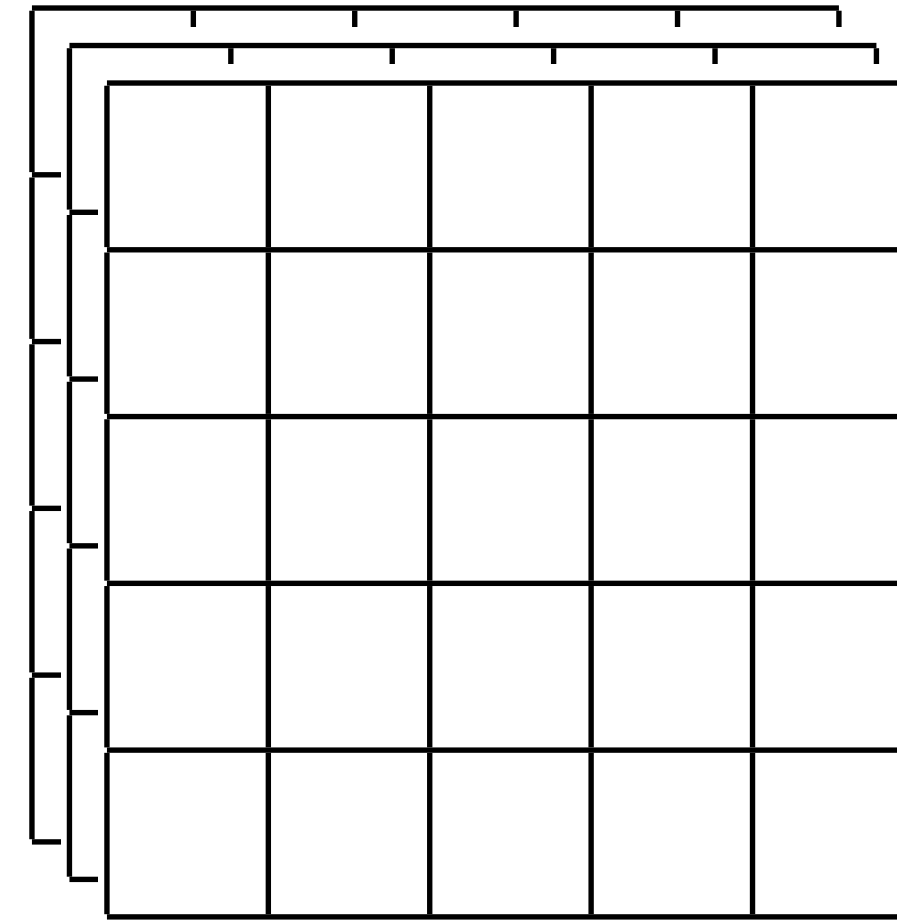
ground truth (1-hot encoding)



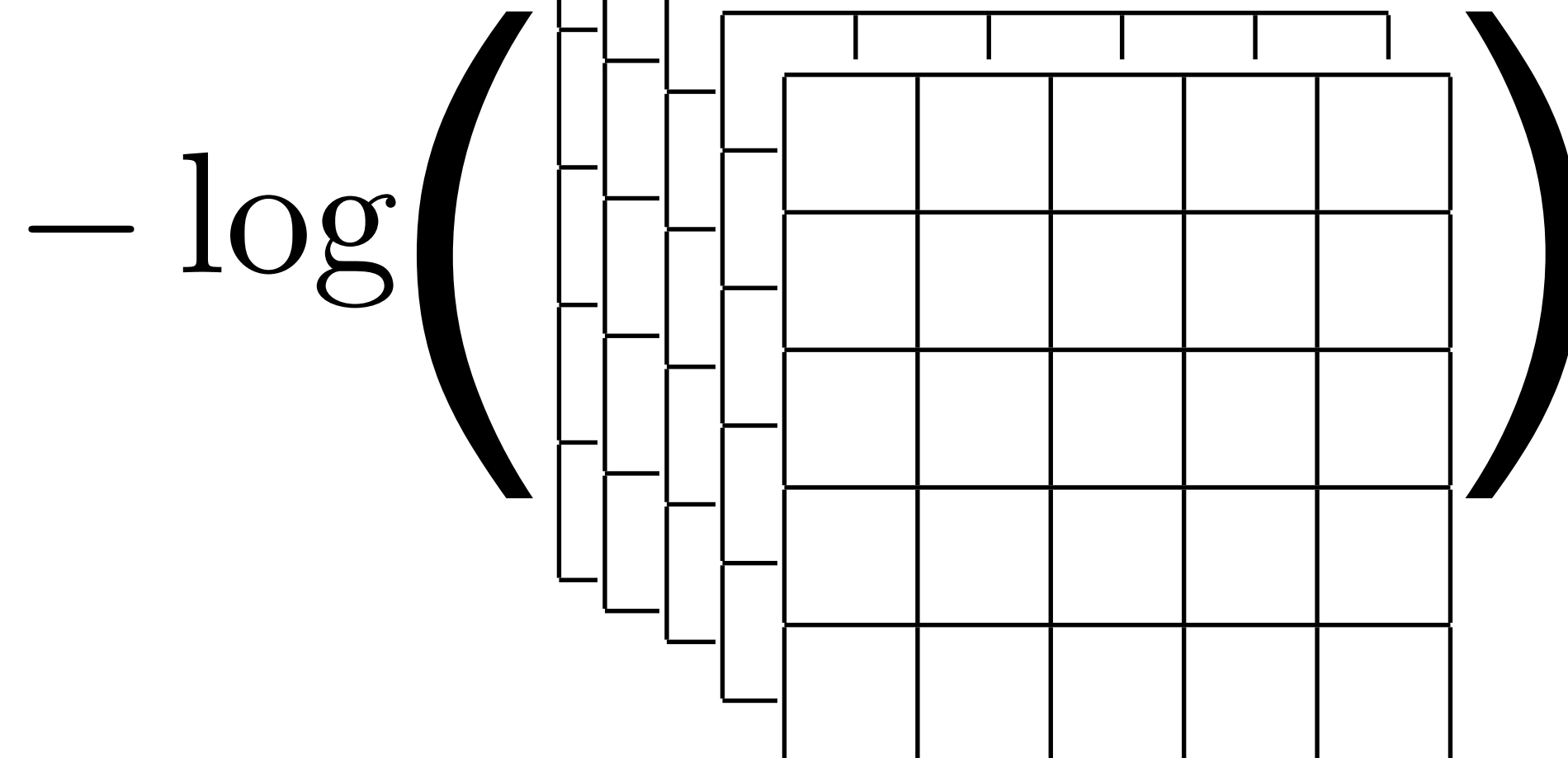
annotations
($H \times W \times N$)

Semantic segmentation

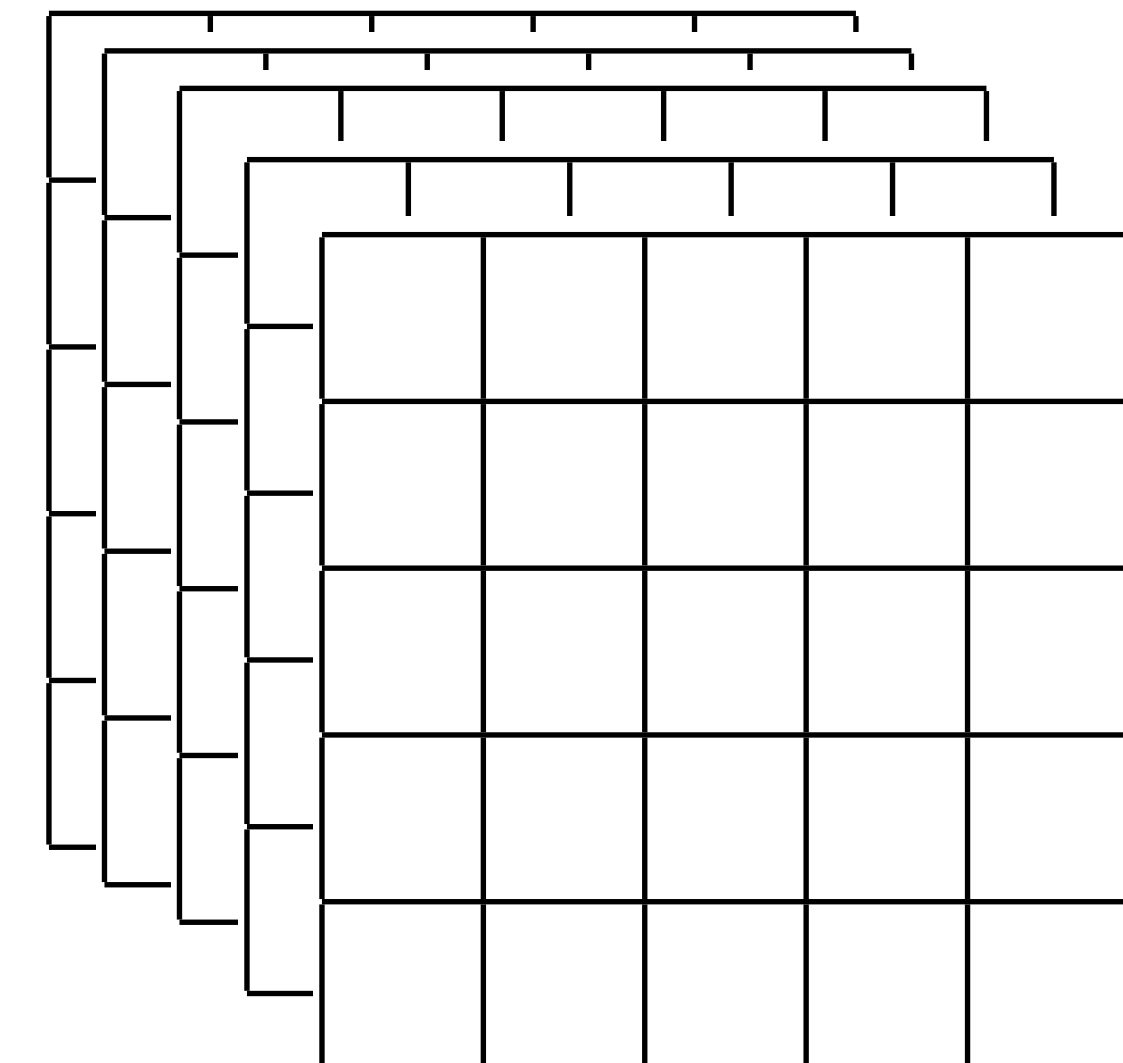
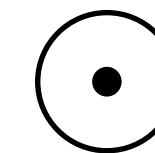
RGB image
(HxWx3)



pixel-wise CE-loss
(HxWxN)



ground truth (1-hot encoding)

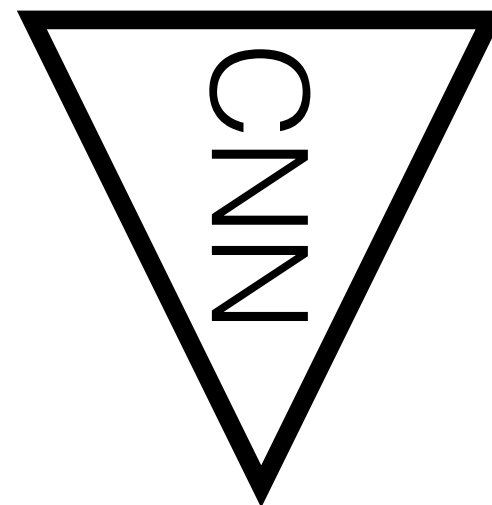
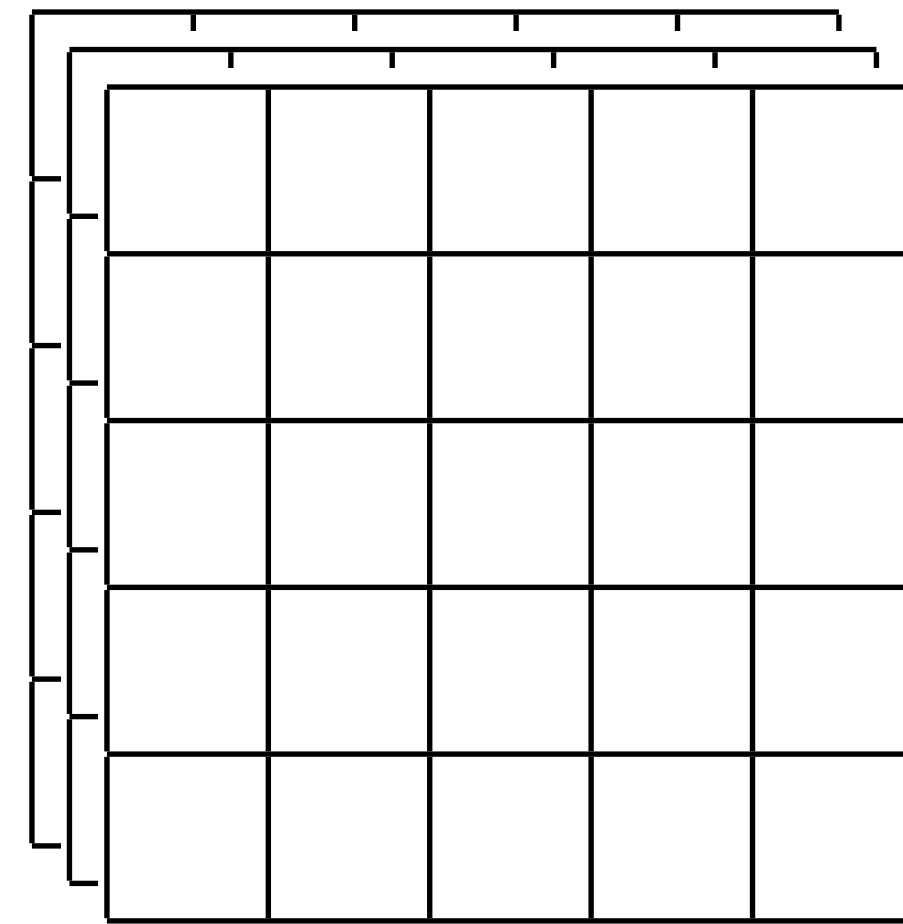


annotations
(HxWxN)

Semantic segmentation

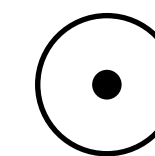
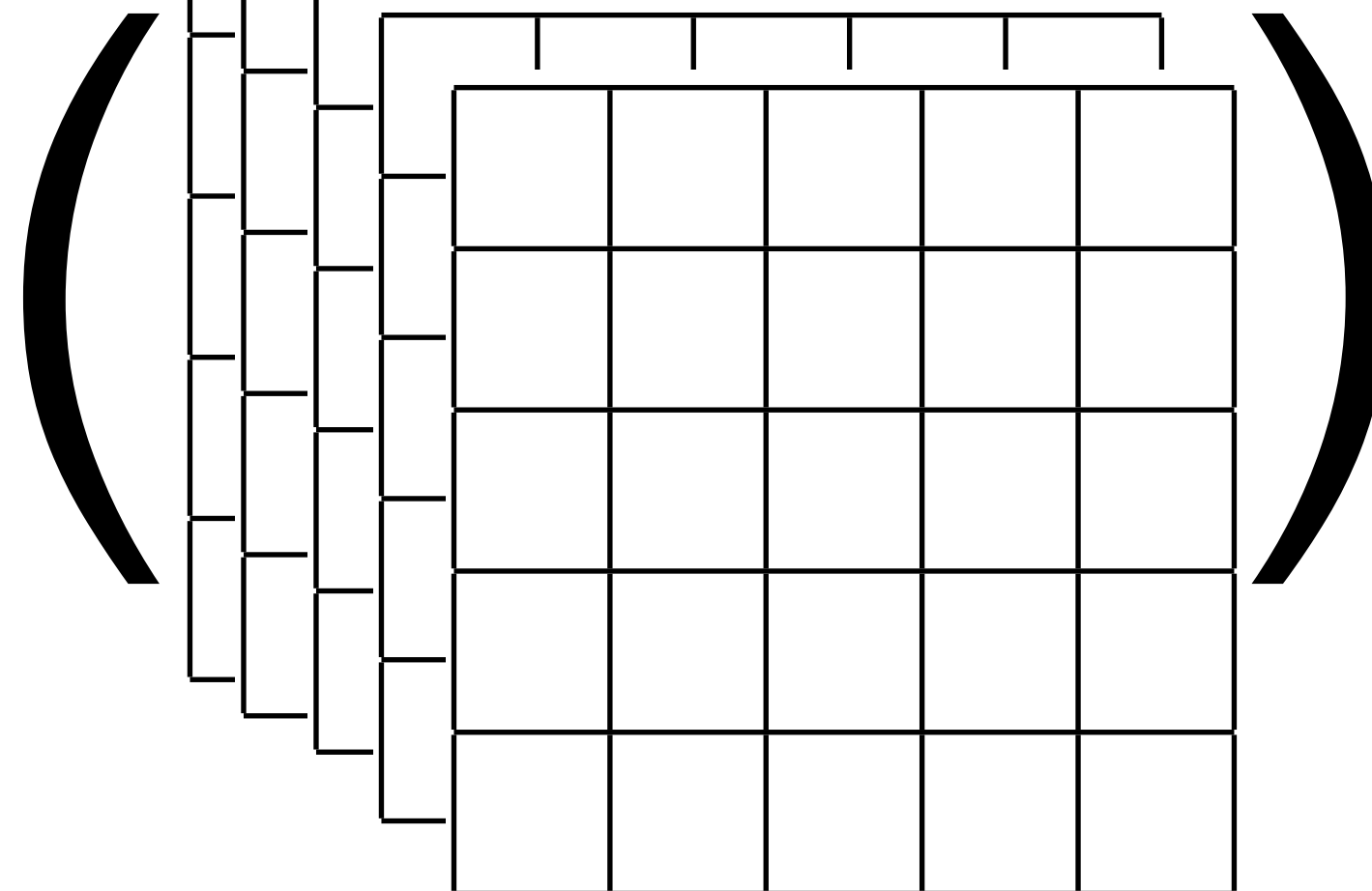
Pixel-wise classifier with the loss summed over all pixels

RGB image
(HxWx3)

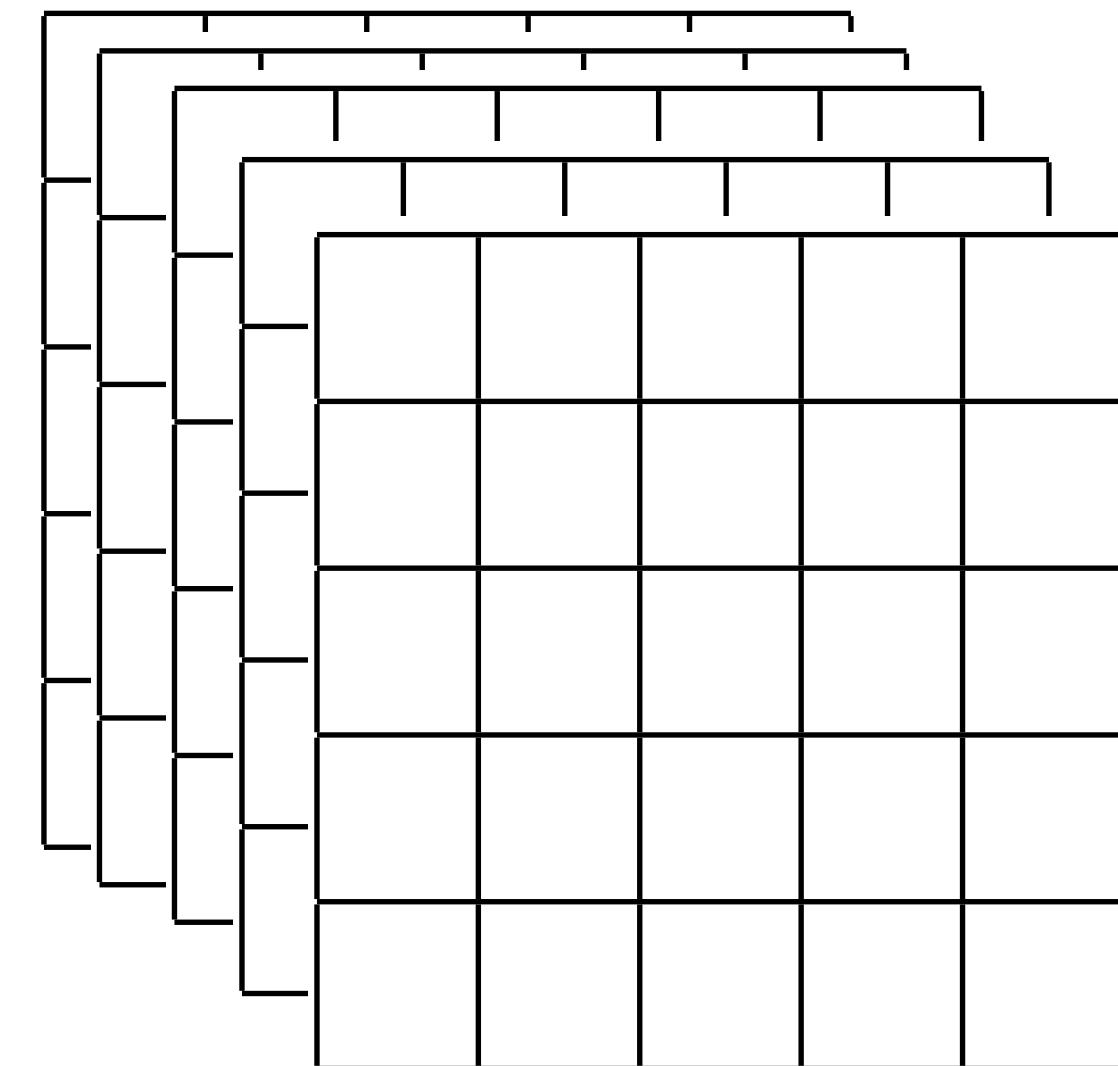


pixel-wise CE-loss
(HxWxN)

$\sum_{\text{pixels channels}} -\log$



ground truth (1-hot encoding)



annotations
(HxWxN)

U-net architecture



Mirror the usual classification network

[Noh et al ICCV 2015] <https://arxiv.org/pdf/1505.04366.pdf>

U-net architecture



[Noh et al ICCV 2015] <https://arxiv.org/pdf/1505.04366.pdf>

U-net architecture

high spatial resolution
 $N \times N \times c$

image embedding

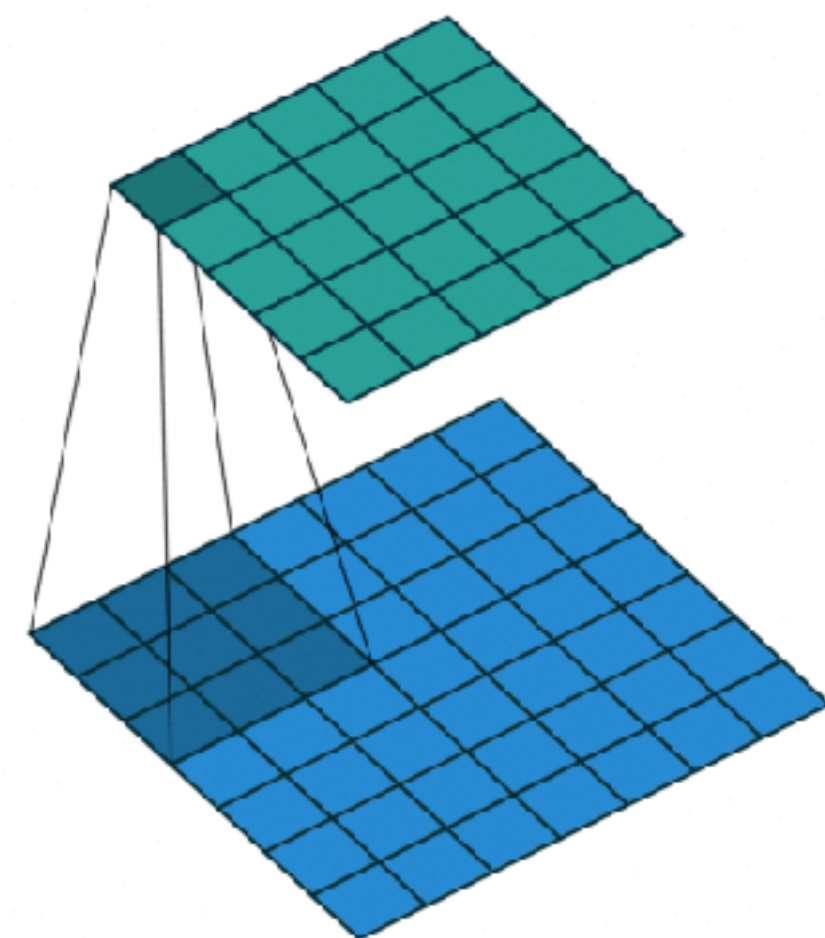
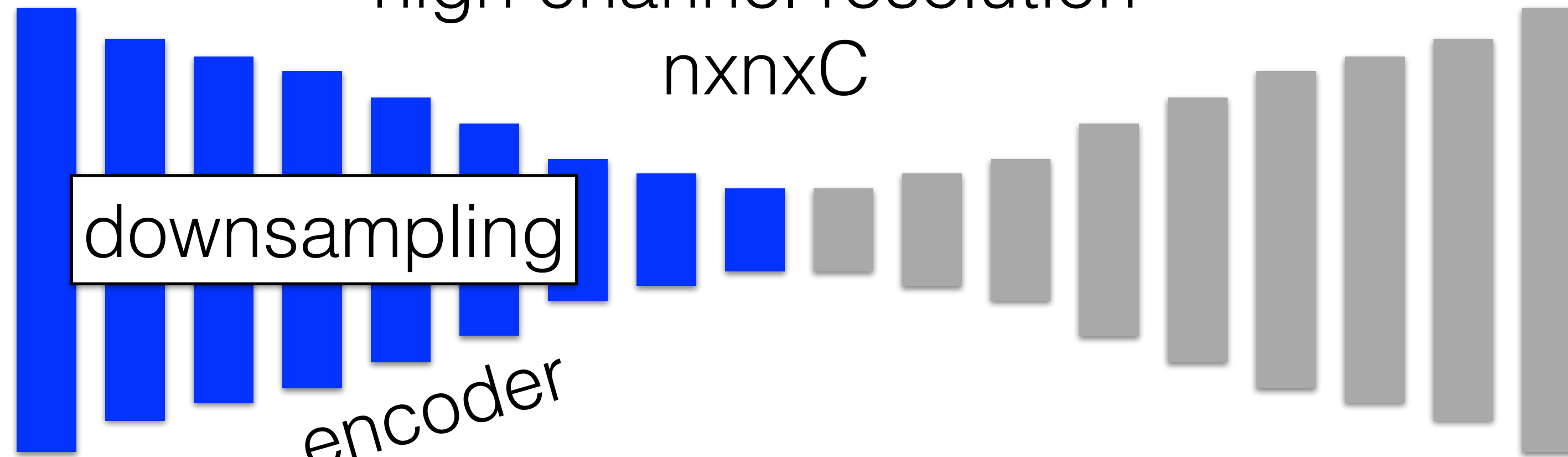
high channel resolution
 $n \times n \times C$

image



downsampling

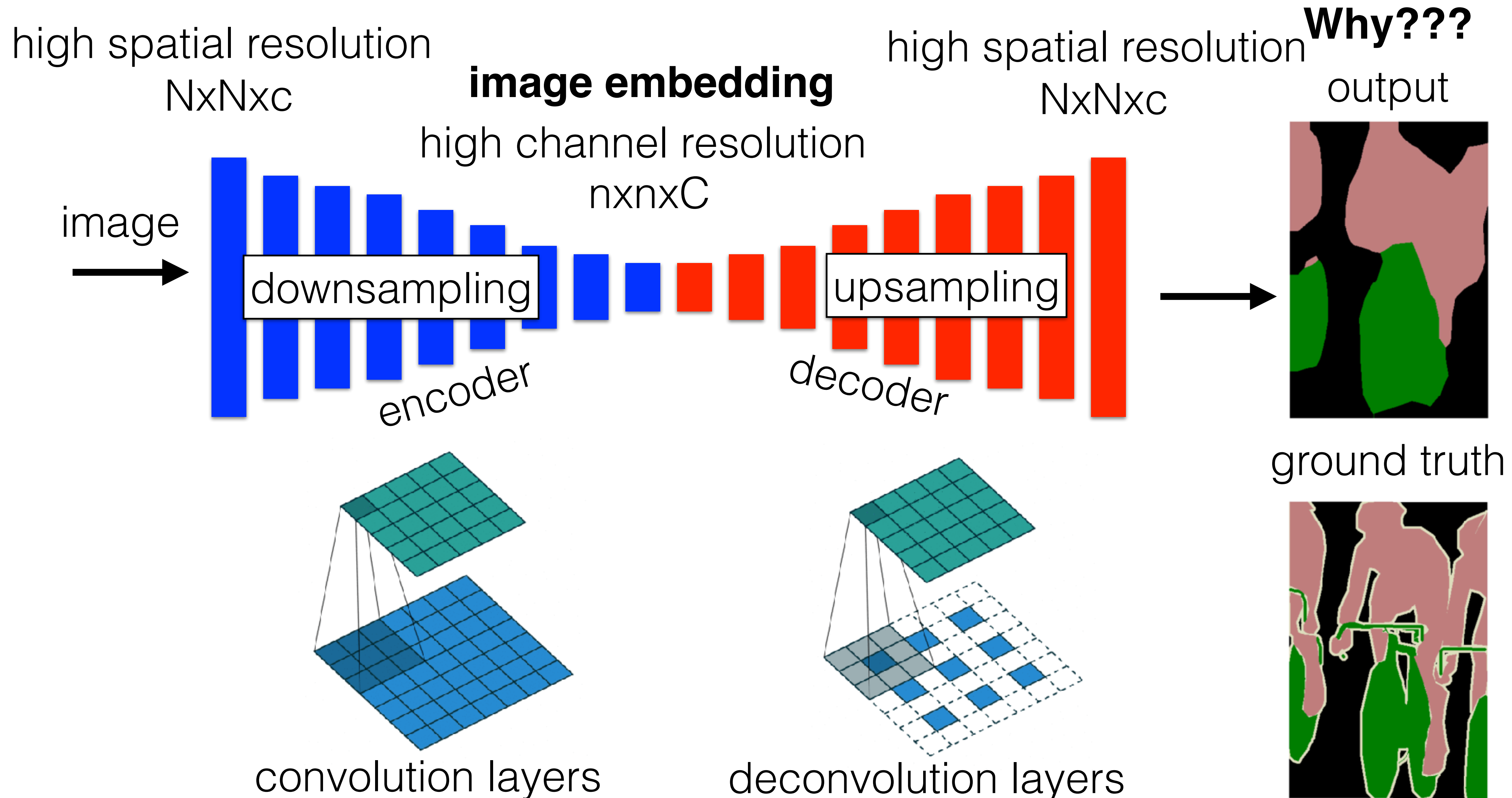
encoder



convolution layers

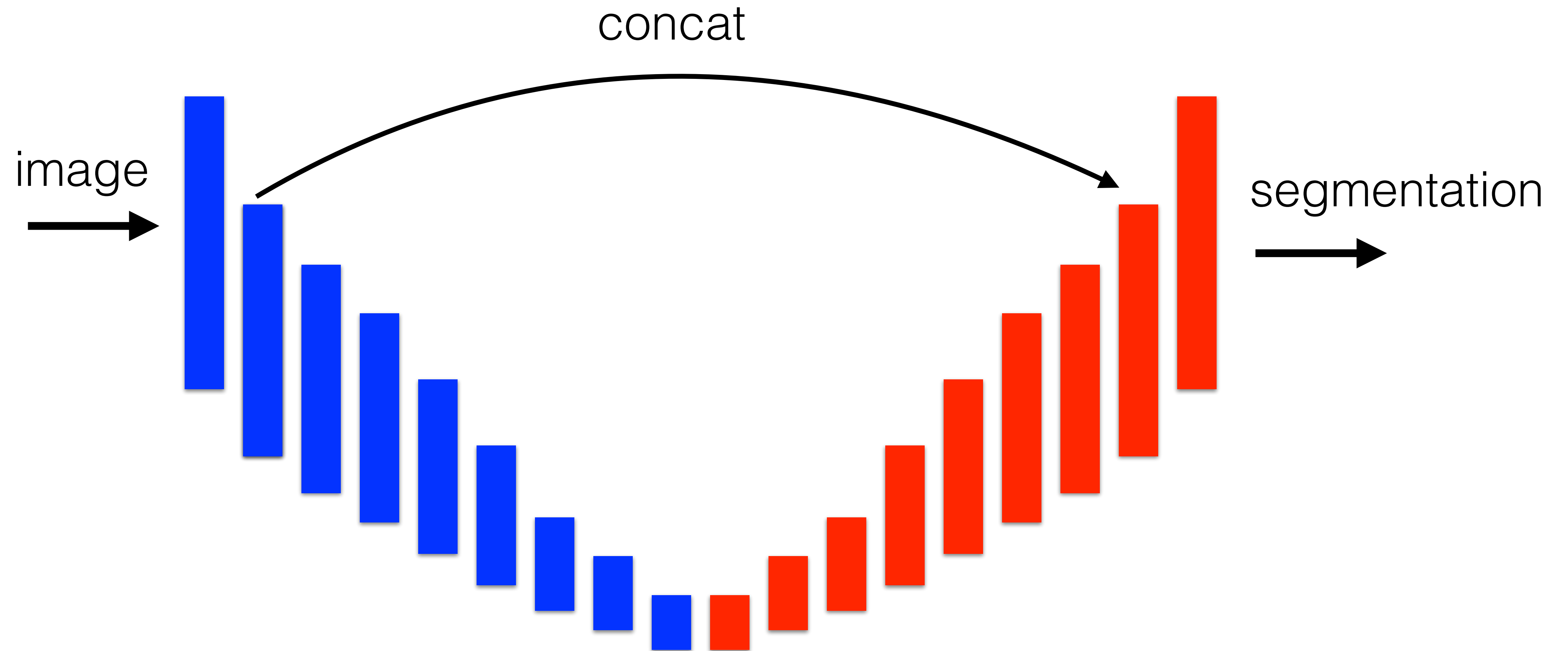
[Noh et al ICCV 2015] <https://arxiv.org/pdf/1505.04366.pdf>

U-net architecture



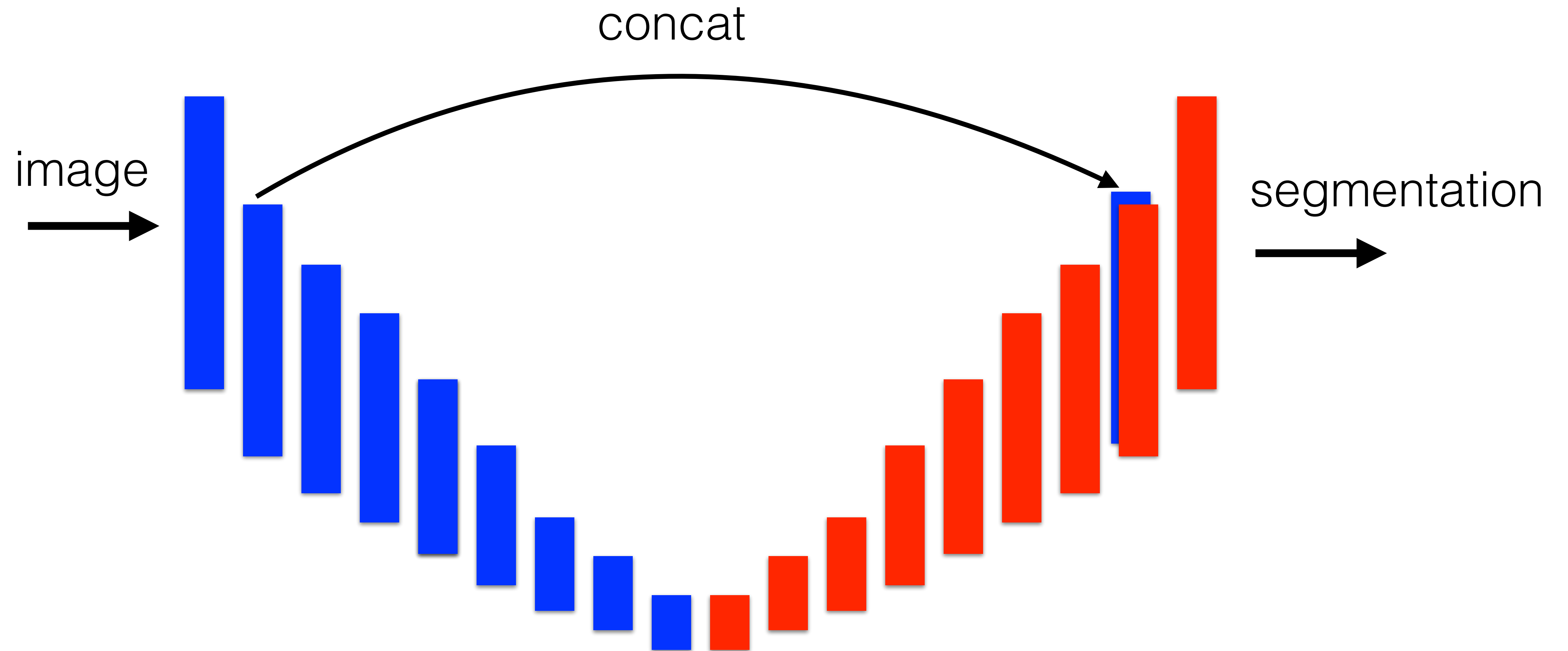
[Noh et al ICCV 2015] <https://arxiv.org/pdf/1505.04366.pdf>

U-net architecture



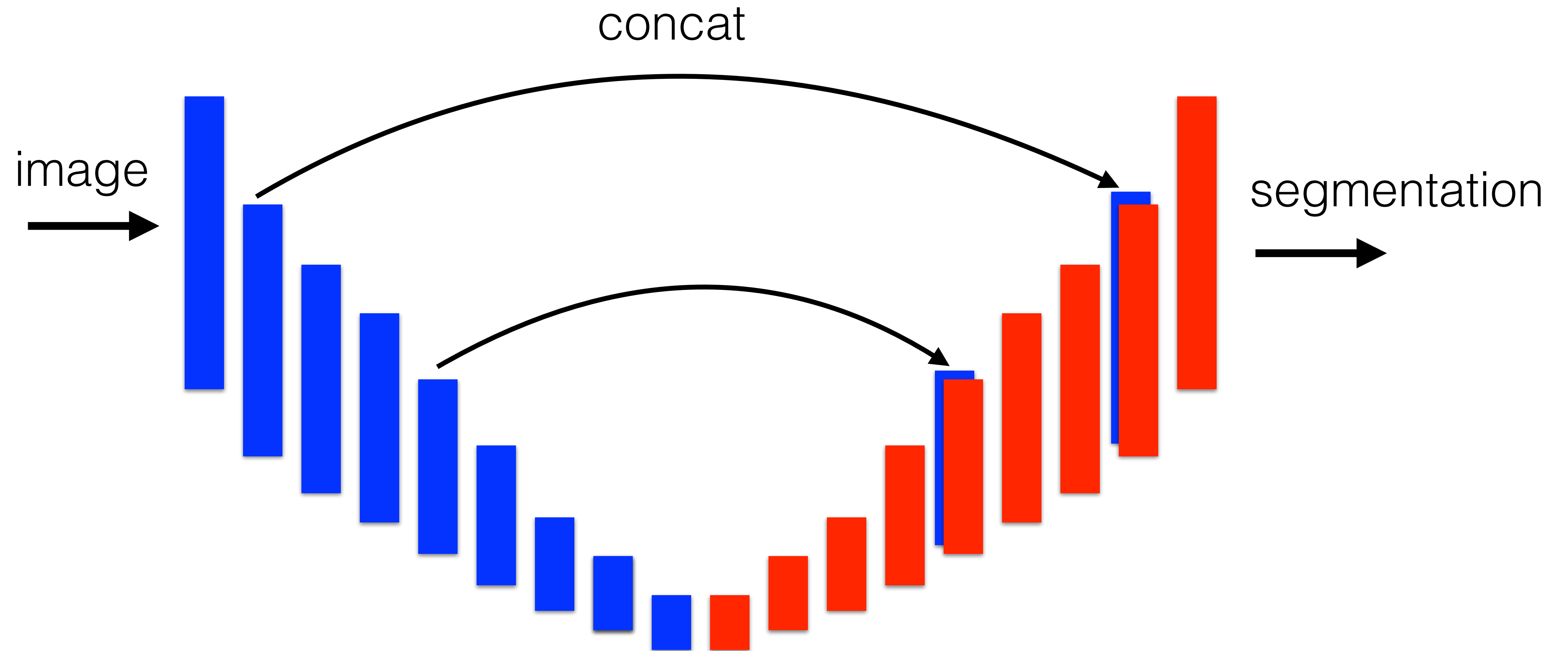
[Noh et al ICCV 2015] <https://arxiv.org/pdf/1505.04366.pdf>

U-net architecture



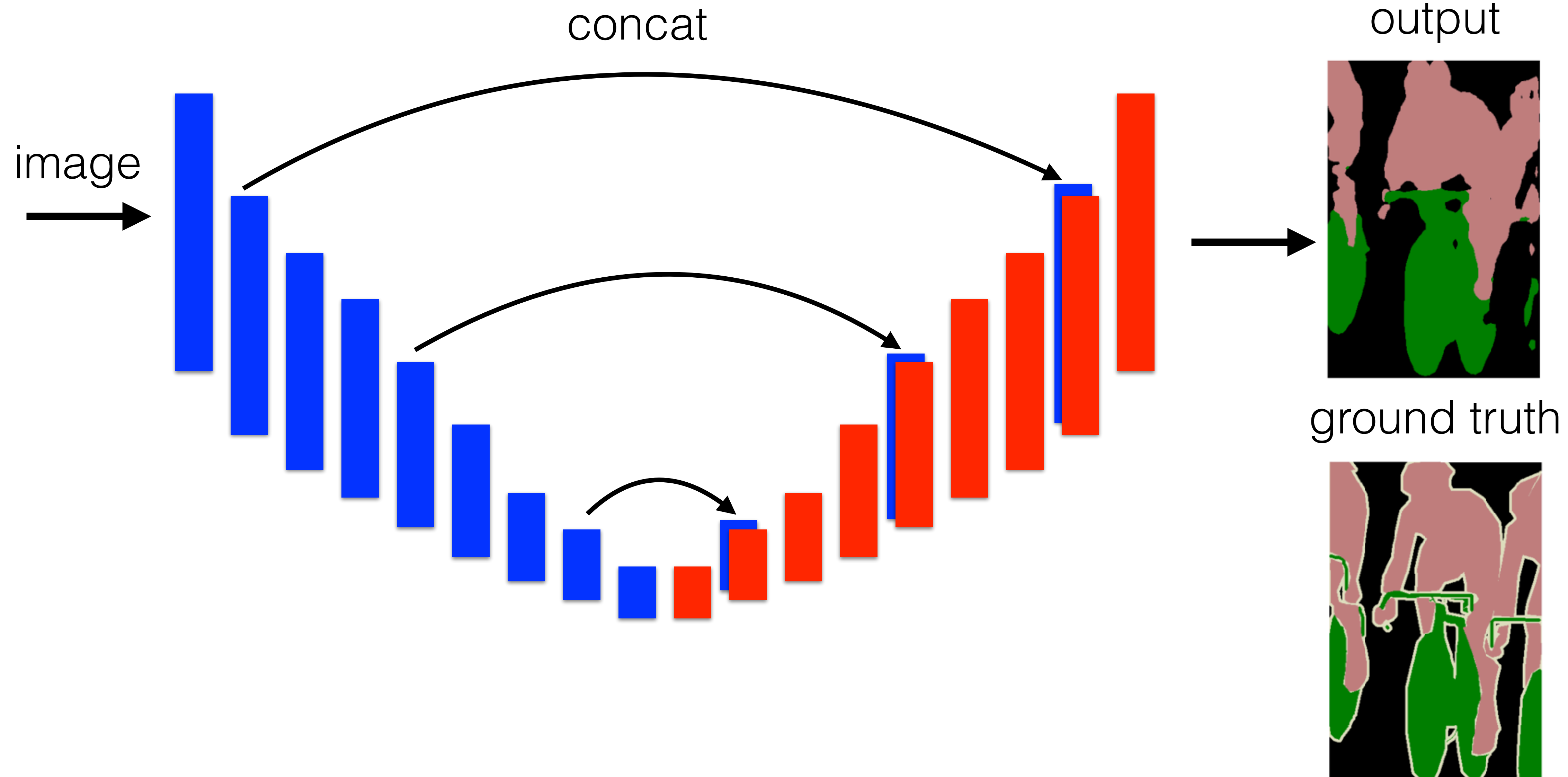
[Noh et al ICCV 2015] <https://arxiv.org/pdf/1505.04366.pdf>

U-net architecture



[Noh et al ICCV 2015] <https://arxiv.org/pdf/1505.04366.pdf>

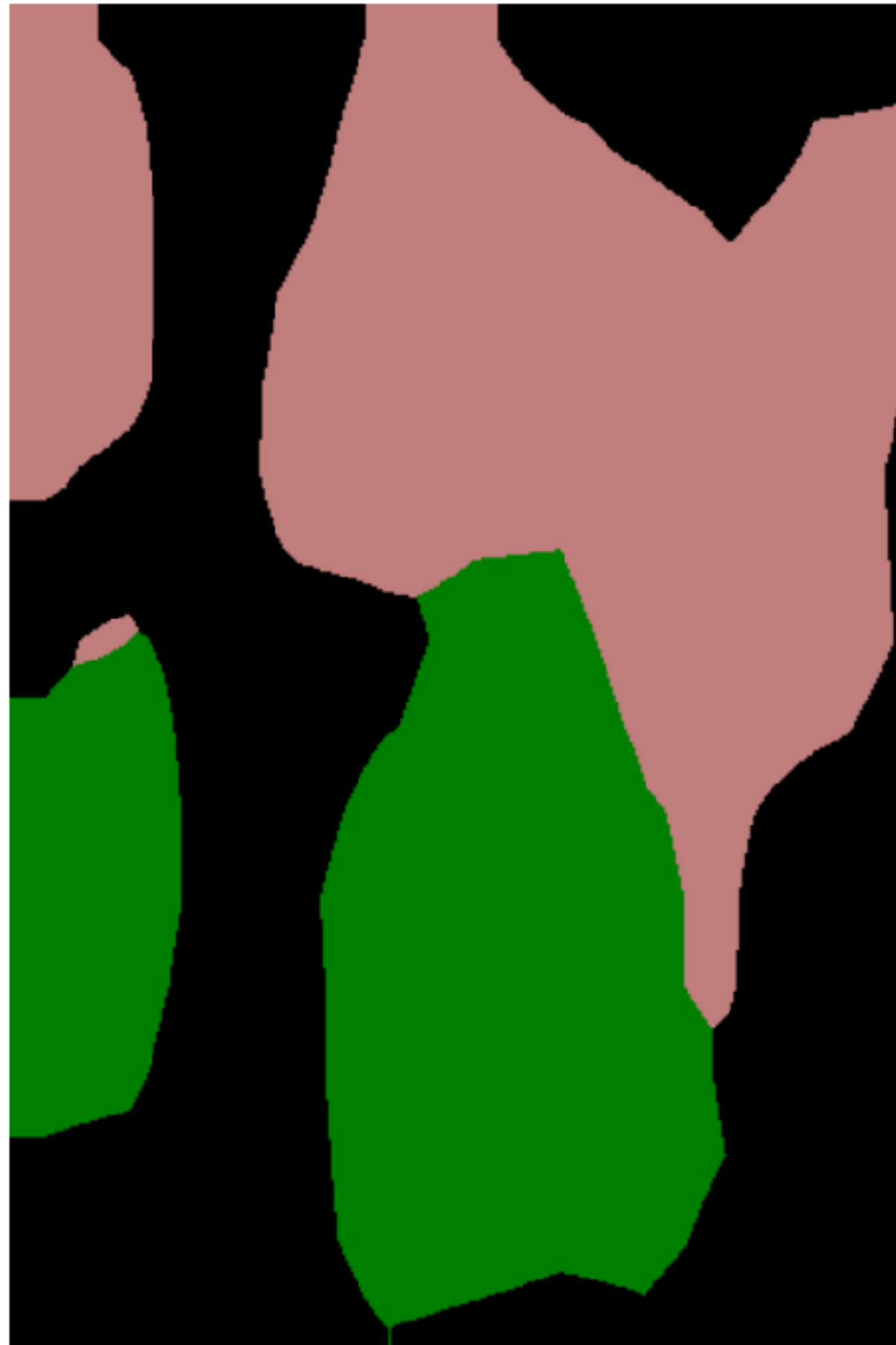
U-net architecture



[Noh et al ICCV 2015] <https://arxiv.org/pdf/1505.04366.pdf>

U-net architecture

no skip connections



with skip connections



ground truth



[Noh et al ICCV 2015] <https://arxiv.org/pdf/1505.04366.pdf>

Segmentation architectures

- **FCN** (Fully Convolutional Network): LongCVPR 2015, FCN was one of the pioneering architectures for end-to-end pixel-wise prediction.
- **U-Net**: symmetric architecture and skip connections, which helps in capturing both global and local features.
- **DeepLab**: DeepLab uses dilated convolutions to enlarge the field of view and capture multi-scale information. DeepLabv3 is an improved version.
- **MobileNet**: lightweight architecture designed for real-time semantic segmentation. It is known for its efficiency and speed.
- **LinkNet**: LinkNet employs a skip connection structure with a series of encoder and decoder blocks for segmentation tasks.
- **HRNet** (High-Resolution Network): HRNet focuses on maintaining high-resolution representations throughout the network, which can be beneficial for capturing fine details.
- **ViTS**: Vision transformer extended for segmentation task
- **SAM**: Segment anything architecture from Facebook 2023