



AI and Deep Learning

4-2 Convolution Neural Network I

Zhonglei Wang

WISE, SOE, CHOW

XMU

(Version v1)

Contents

1. Basic structure

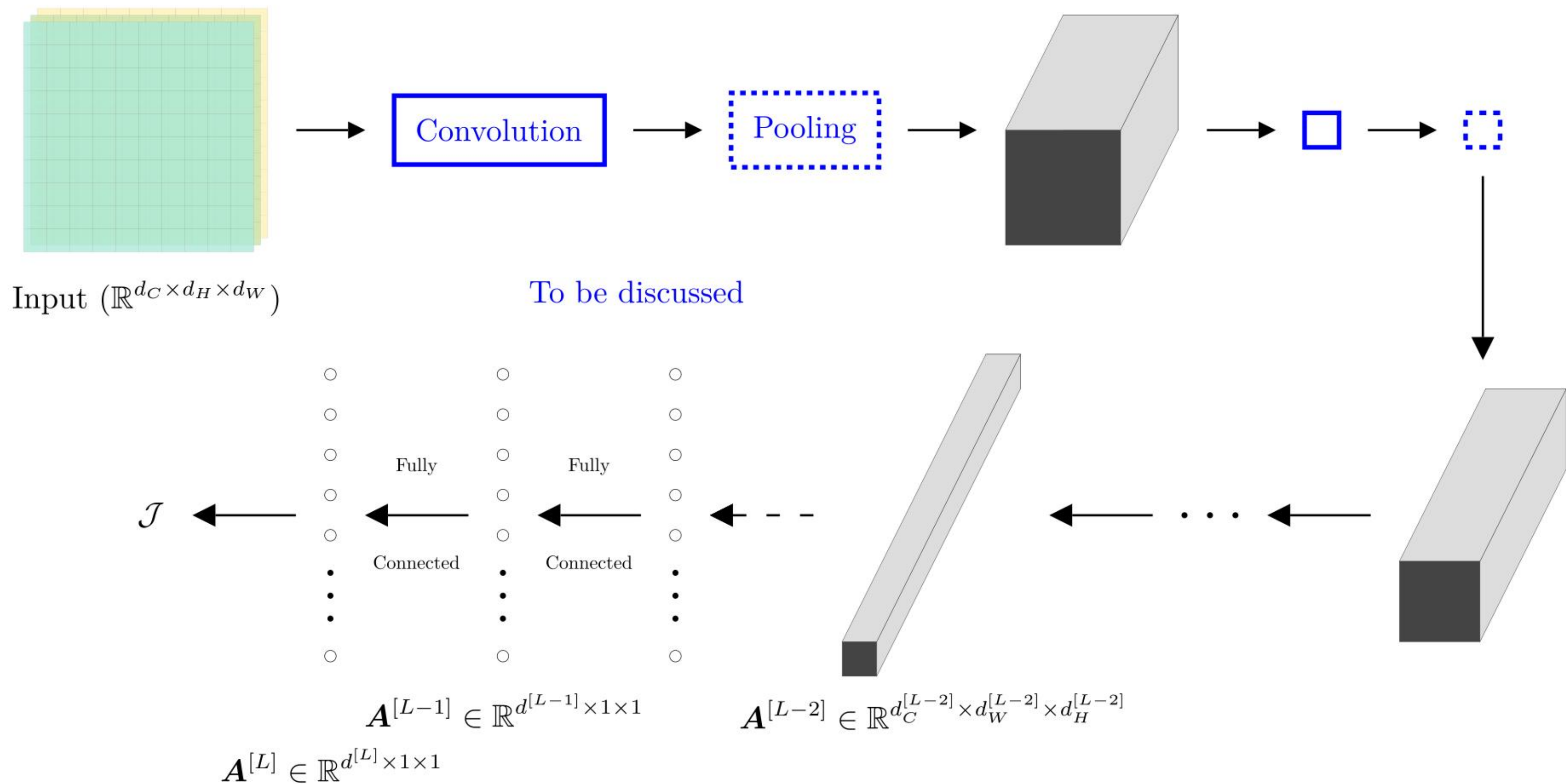
2. Forward propagation

3. Backpropagation

4. Understanding convolution networks

Background

1. Given an image, why not use a fully connected neural network?
 - Number of parameters is **overwhelmingly** large, easily leading to overfitting
 - ▷ To overcome this problem, we may require more training data, more memory and faster computation algorithms
 - ▷ Almost impossible in practice
 - Overlooks the spatial structure
 - Not location (translation) invariant
2. Consider a new network, convolution neural network, for image processing



Summary

1. Generally, $d_C^{[l]}$ increases as l increases, but $d_W^{[l]}$ and $d_H^{[l]}$ decrease
2. The vectorization step does not involve any new model parameter
3. Typically, there are three “layers” for CNN
 - Convolution layers (CONV)
 - Pooling layers (POOL)
 - Fully connected layers (FC)

Notation

1. l : Layer index
2. $d_C^{[l]}$: Number of channels of the l th layer
3. $d_H^{[l]}$: Height of the “image” of the l th layer
4. $d_W^{[l]}$: Width of the “image” of the l th layer
5. $f^{[l]}$: Kernel size associated with the l th layer ($f_1^{[l]} = f_2^{[l]}$ by default)
6. $p^{[l]}$: Padding associated with the l th layer ($p_1^{[l]} = p_2^{[l]}$ by default)
7. $s^{[l]}$: Stride associated with the l th layer ($s_1^{[l]} = s_2^{[l]}$ by default)
8. No dilation

Forward propagation for a CONV layer

1. $\mathbf{Z}^{[l]} \in \mathbb{R}^{n \times d_C^{[l]} \times d_H^{[l]} \times d_W^{[l]}}$: Output of linear transformation

$$\mathbf{Z}^{[l]} = \mathbf{A}^{[l-1]} \ast \mathbf{W}^{[l]} + \mathbf{b}^{[l]}$$

- $\mathbf{A}^{[l-1]} \in \mathbb{R}^{n \times d_C^{[l-1]} \times d_H^{[l-1]} \times d_W^{[l-1]}}$: Input for the l th layer
- $\mathbf{W}^{[l]} \in \mathbb{R}^{d_C^{[l]} \times d_C^{[l-1]} \times f^{[l]} \times f^{[l]}}$: Kernels
- $\mathbf{b}^{[l]} \in \mathbb{R}^{d_C^{[l]} \times 1 \times 1}$: Bias term
- “ $\ast$ ”: Convolution for each sample and each channel
- “ $+$ ”: Common bias for each channel

2. $\mathbf{A}^{[l]} \in \mathbb{R}^{n \times d_C^{[l]} \times d_H^{[l]} \times d_W^{[l]}}$: Output of activation

$$\mathbf{A}^{[l]} = \sigma^{[l]}(\mathbf{Z}^{[l]})$$

- $\sigma^{[l]}(\cdot)$: Activation function for the l th layer

Forward propagation for a CONV layer

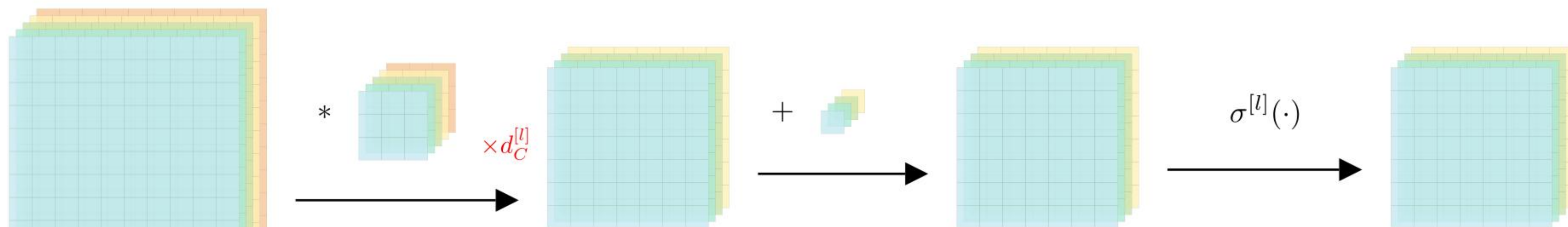
1. For simplicity, we show forward- and back-propagation for $n = 1$

- $\mathbf{A}^{[l-1]} \in \mathbb{R}^{d_C^{[l-1]} \times d_H^{[l-1]} \times d_W^{[l-1]}}$
- $\mathbf{Z}^{[l]} \in \mathbb{R}^{d_C^{[l]} \times d_H^{[l]} \times d_W^{[l]}}$

2. Forward propagation:

$$\begin{aligned}\mathbf{Z}^{[l]} &= \mathbf{A}^{[l-1]} \ast \mathbf{W}^{[l]} + \mathbf{b}^{[l]}, \\ \mathbf{A}^{[l]} &= \sigma^{[l]}(\mathbf{Z}^{[l]})\end{aligned}$$

Visualization



$$\mathbf{A}^{[l-1]} \in \mathbb{R}^{d_C^{[l-1]} \times d_H^{[l-1]} \times d_W^{[l-1]}}$$

$$\mathbf{Z}_0^{[l]} = \mathbb{R}^{d_C^{[l]} \times d_H^{[l]} \times d_W^{[l]}}$$

$$\mathbf{Z}^{[l]} \in \mathbb{R}^{d_C^{[l]} \times d_H^{[l]} \times d_W^{[l]}}$$

$$\mathbf{W}^{[l]} \in \mathbb{R}^{d_C^{[l]} \times d_C^{[l-1]} \times f^{[l]} \times f^{[l]}}$$

$$\mathbf{b}^{[l]} \in \mathbb{R}^{d_C^{[l]} \times 1 \times 1}$$

$$\mathbf{A}^{[l]} \in \mathbb{R}^{d_C^{[l]} \times d_H^{[l]} \times d_W^{[l]}}$$

Pooling

1. Pooling is typically used after convolution layers to achieve

- Reduce height and width for each **channel** (downsampling), and computation efficiency is improved
- Achieves robustness by neglecting useless or repeated information
- Increase the receptive field
- Achieve robustness against small translations, rotations, and other distortions in the input
- Prevent overfitting, especially when combined with dropout

Pooling

1. Two types of pooling:

- Average pooling
- Max pooling

2. Notice:

- Pooling is conducted for **each channel**
- Pooling does not involve new model parameters

Average pooling

1. Input size: 6×6
2. Kernel size: 2×2
3. Stride: 2

Average pooling

Input

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

Kernel

0.25	0.25
0.25	0.25

Result

4.25	6.0	6.25
3.0	6.0	2.25
3.25	3.75	3.0

Max pooling

1. Input size: 6×6
2. Kernel size: 2×2
3. Stride: 2

Max pooling

Input

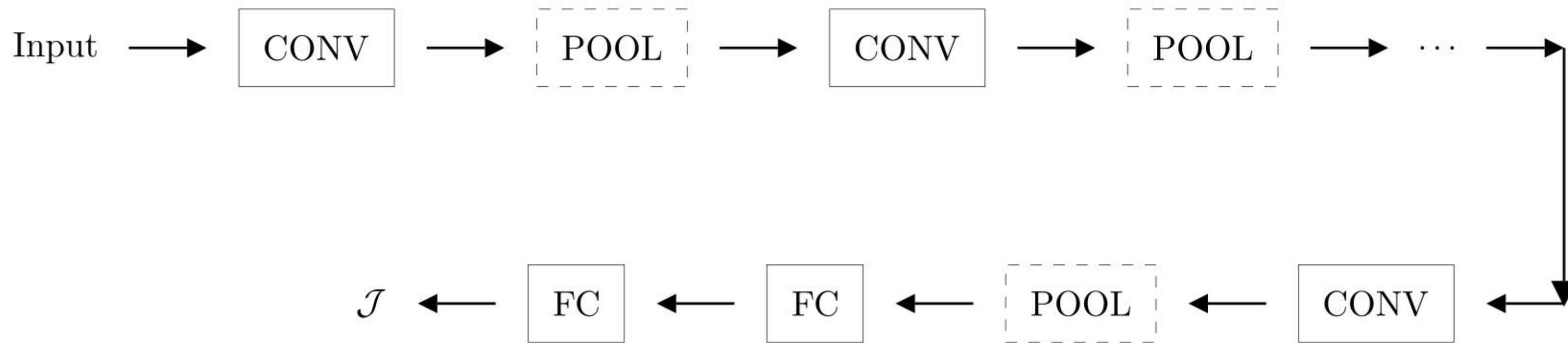
3	6	5	4	8	9
1	7	9	6	8	0
5	0	9	6	2	0
5	2	6	3	7	0
9	0	3	2	3	1
3	1	3	7	1	7

Kernel

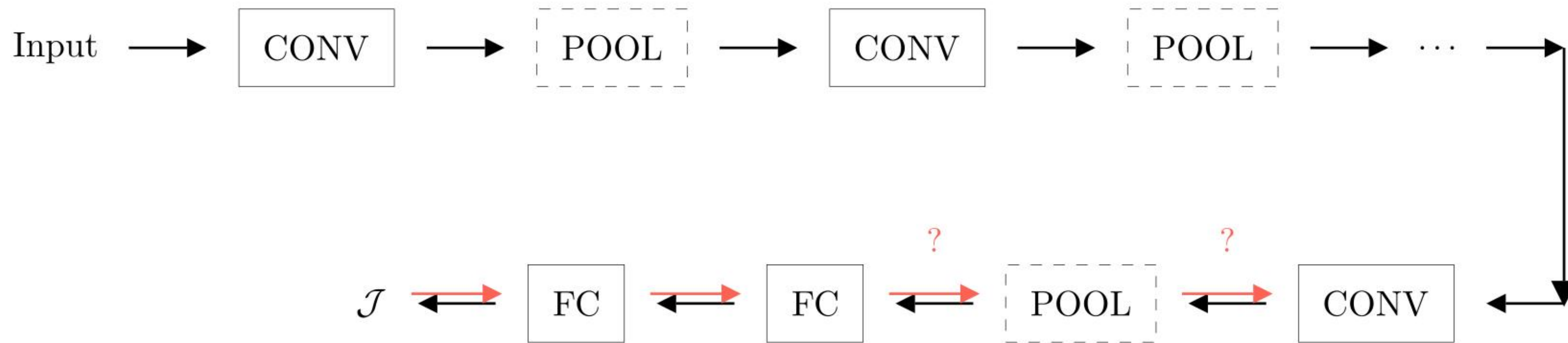
Result

7	9	9
5	9	7
9	7	7

Forward propagation



Backpropagation



1. We have learnt backpropagation for fully connected layers
2. Thus, it remains to show the backpropagation for
 - The pooling layer (POOL)
 - The convolution layer (CONV)
3. In the following, we assume a POOL is conducted right after a CONV

Backpropagation

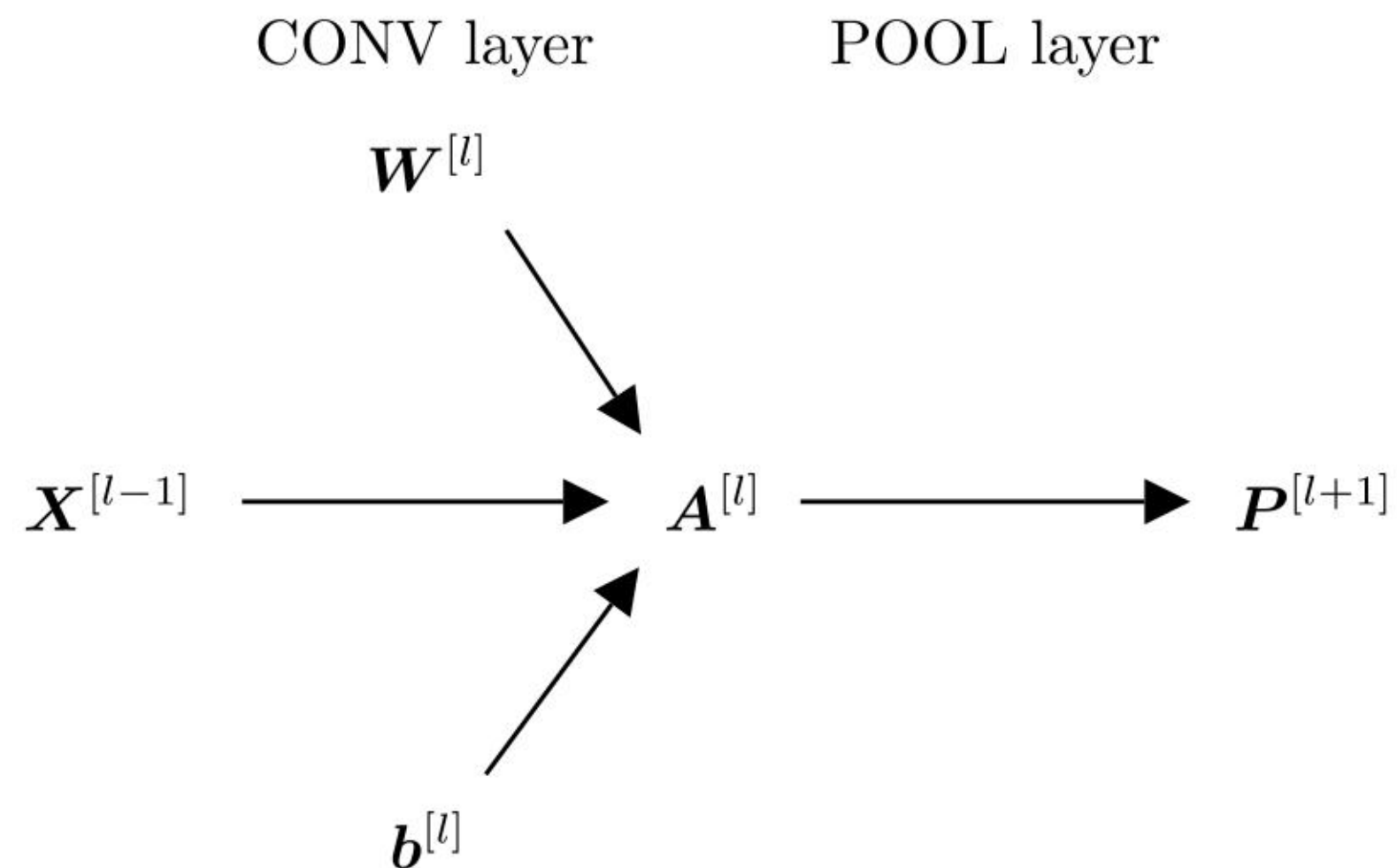
1. Consider the following case:

- The l th layer is a CONV layer
- The $(l + 1)$ th layer is a POOL layer

2. Denote

- $\mathbf{W}^{[l]} \in \mathbb{R}^{d_C^{[l]} \times d_C^{[l-1]} \times f^{[l]} \times f^{[l]}}$: Kernel for the l th layer
- $\mathbf{b}^{[l]} \in \mathbb{R}^{d_C^{[l]} \times 1 \times 1}$: Bias for the l th layer, taking $\mathbf{P}^{[l-1]}$ as input
- $\mathbf{A}^{[l]} \in \mathbb{R}^{d_C^{[l]} \times d_H^{[l]} \times d_W^{[l]}}$: Output of the l th layer (CONV)
- $\mathbf{P}^{[l+1]} \in \mathbb{R}^{d_C^{[l]} \times d_H^{[l+1]} \times d_W^{[l+1]}}$: Output of the $(l + 1)$ th layer (POOL)

Structure



1. Assume that $d\mathbf{P}^{[l+1]}$ is available
 - Backpropagation for vectorization is trivial

Backpropagation for average pooling

$\mathbf{A}^{[l]}$

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

Kernel $\mathbf{M}$

0.25	0.25
0.25	0.25

$\mathbf{P}^{[l+1]}$

4.25	6.0	6.25
3.0	6.0	2.25
3.25	3.75	3.0

$$\mathbf{P}_{11}^{[l+1]} = 0.25 \times (\mathbf{A}_{11}^{[l]} + \mathbf{A}_{12}^{[l]} + \mathbf{A}_{21}^{[l]} + \mathbf{A}_{22}^{[l]})$$

$$d\mathbf{A}_{11}^{[l]} = \frac{\partial \mathcal{J}}{\partial \mathbf{A}_{11}^{[l]}} = \sum_{i=1}^3 \sum_{j=1}^3 \frac{\partial \mathcal{J}}{\partial \mathbf{P}_{ij}^{[l+1]}} \times \frac{\partial \mathbf{P}_{ij}^{[l+1]}}{\partial \mathbf{A}_{11}^{[l]}} = \frac{\partial \mathcal{J}}{\partial \mathbf{P}_{11}^{[l+1]}} \times \frac{\partial \mathbf{P}_{11}^{[l+1]}}{\partial \mathbf{A}_{11}^{[l]}} = d\mathbf{P}_{11}^{[l+1]} \times 0.25$$

$$d\mathbf{A}_{12}^{[l]} = d\mathbf{A}_{21}^{[l]} = d\mathbf{A}_{22}^{[l]} = 0.25 \times d\mathbf{P}_{11}^{[l+1]}$$

Backpropagation for average pooling

1. For average pooling, we have

$$d\mathbf{A}^{[l]} = d\mathbf{P}^{[l+1]} \otimes \mathbf{M}$$

- $\mathbf{A} \otimes \mathbf{B}$: Kronecker production of two matrices $\mathbf{A}$ and $\mathbf{B}$
- Valid for a special case:
 - ▷ Kernel size “equals to” stride
 - ▷ Size of the original image can be “divided by” the kernel size

2. Chain rule can be used to derive the result for general cases (check by yourself)

Backpropagation for max pooling

$\mathbf{A}^{[l]}$

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

Kernel

$\mathbf{P}^{[l+1]}$

7	9	9
5	9	7
9	7	7

$$\mathbf{P}_{11}^{[l+1]} = \max\{\mathbf{A}_{11}^{[l]} + \mathbf{A}_{12}^{[l]} + \mathbf{A}_{21}^{[l]} + \mathbf{A}_{22}^{[l]}\} = \mathbf{A}_{22}^{[l]}$$

$$d\mathbf{A}_{11}^{[l]} = \frac{\partial \mathcal{J}}{\partial \mathbf{A}_{11}^{[l]}} = \sum_{i=1}^3 \sum_{j=1}^3 \frac{\partial \mathcal{J}}{\partial \mathbf{P}_{ij}^{[l+1]}} \times \frac{\partial \mathbf{P}_{ij}^{[l+1]}}{\partial \mathbf{A}_{11}^{[l]}} = \frac{\partial \mathcal{J}}{\partial \mathbf{P}_{11}^{[l+1]}} \times \frac{\partial \mathbf{P}_{11}^{[l+1]}}{\partial \mathbf{A}_{11}^{[l]}} = d\mathbf{P}_{11}^{[l+1]} \times 0 = 0$$

$$d\mathbf{A}_{12}^{[l]} = d\mathbf{A}_{21}^{[l]} = 0; \quad d\mathbf{A}_{22}^{[l]} = d\mathbf{P}_{11}^{[l+1]}$$

Backpropagation for max pooling

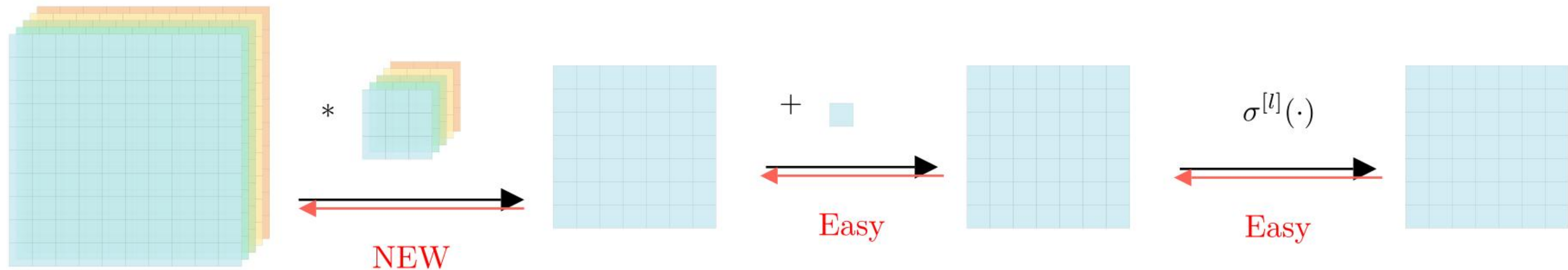
1. For max pooling, $d\mathbf{A}^{[l]}$ is obtained by
 - Stacking small matrices associated with each elements in $\mathbf{P}^{[l+1]}$ times a 2×2 mask matrix
 - Valid for a special case:
 - ▷ Kernel size “equals to” stride
 - ▷ Size of the original image can be “divided by” the kernel size
2. Chain rule can be used to derive the result for general cases (check by yourself)

Remarks

1. We have finished the backpropagation from $d\mathbf{P}^{[l+1]}$ to $d\mathbf{A}^{[l]}$
 - For average pooling, under specific conditions, gradients are **evenly distributed**
 - For max pooling, under specific conditions, gradient is propagated back **only at the position of the forward maximum**
 - Pooling layer does not involve any model parameters
2. Assume the availability of $d\mathbf{A}^{[l]}$ for the following analysis

Backpropagation for CONV

1. For simplicity, let $d_C^{[l]} = 1$.



$$\mathbf{X}^{[l-1]} \in \mathbb{R}^{d_C^{[l-1]} \times d_H^{[l-1]} \times d_W^{[l-1]}}$$

$$\mathbf{Z}_0^{[l]} \in \mathbb{R}^{1 \times d_H^{[l]} \times d_W^{[l]}}$$

$$\mathbf{Z}_1^{[l]} \in \mathbb{R}^{1 \times d_H^{[l]} \times d_W^{[l]}}$$

$$\mathbf{W}^{[l]} \in \mathbb{R}^{1 \times d_C^{[l-1]} \times f^{[l]} \times f^{[l]}}$$

$$\mathbf{b}^{[l]} \in \mathbb{R}^{1 \times 1 \times 1}$$

$$\mathbf{A}^{[l]} \in \mathbb{R}^{1 \times d_H^{[l]} \times d_W^{[l]}}$$

$$d\mathbf{W}^{[l]} = ?$$

$$d\mathbf{b}^{[l]} = \sum_{i=1}^{d_H^{[l]}} \sum_{j=1}^{d_W^{[l]}} dZ_{1,ij}^{[l]}$$

$$d\mathbf{X}^{[l-1]} = ?$$

$$d\mathbf{Z}_0^{[l]} = d\mathbf{Z}_1^{[l]}$$

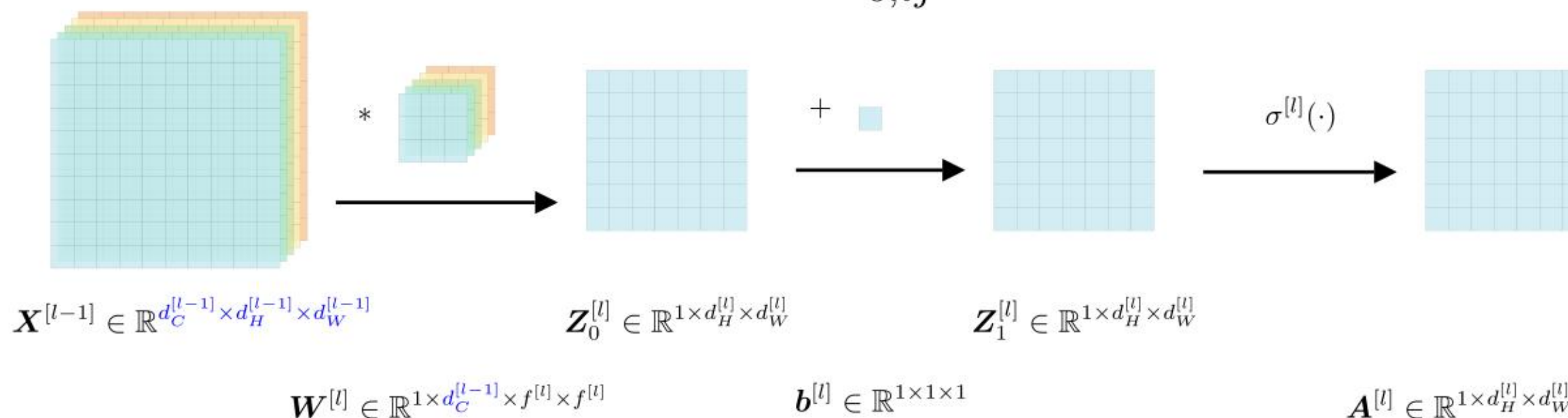
$$d\mathbf{Z}_1^{[l]} = \sigma^{[l]'}(\mathbf{Z}^{[l]}) \circ d\mathbf{A}^{[l]}$$

Backpropagation for CONV

1. To find derivative with respect to something, we need to find where its information is contained
2. Initialize $d\mathbf{X}^{[l-1]}$ as a zero matrix of the same dimension as the input image $\mathbf{X}^{[l-1]}$
3. Informally, obtain

$$d\mathbf{X}_{\text{slice},ij}^{[l-1]} + = d\mathbf{Z}_{0,ij}^{[l]} \times \mathbf{W}^{[l]}$$

- $\mathbf{X}_{\text{slice},ij}^{[l-1]}$: The part of $\mathbf{X}^{[l-1]}$ used to obtain $\mathbf{Z}_{0,ij}^{[l]}$

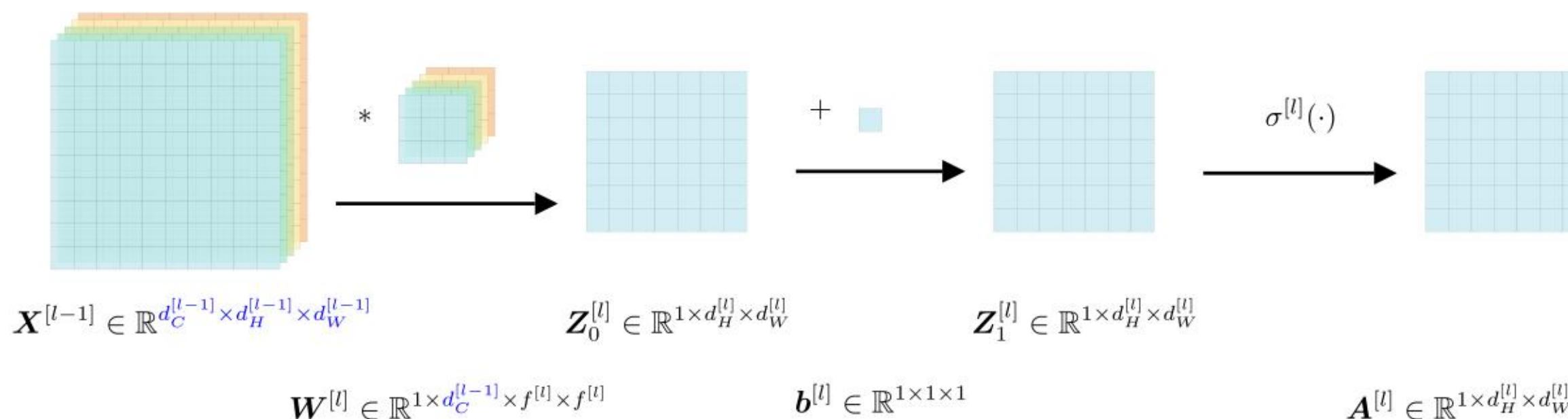


Backpropagation for CONV

1. It is obvious that every element in $\mathbf{Z}_0^{[l]}$ contains information of $\mathbf{W}^{[l]}$
2. Thus, we have

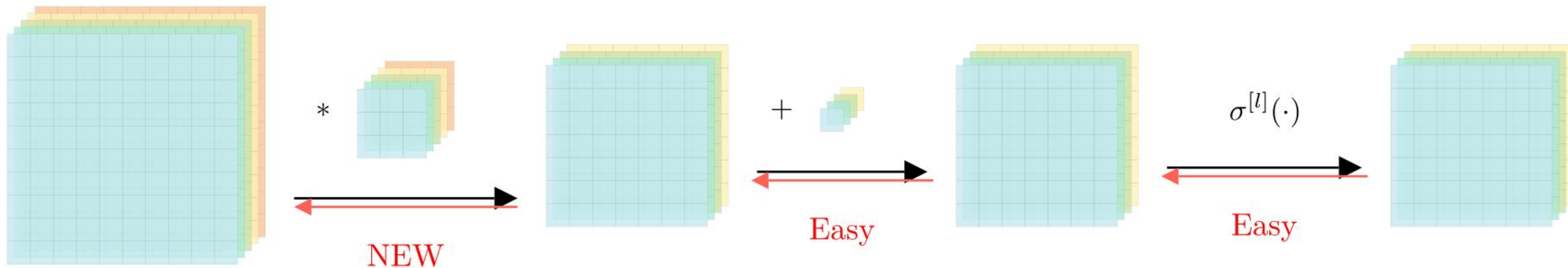
$$d\mathbf{W}^{[l]} = \sum_{i=1}^{d_H^{[l]}} \sum_{j=1}^{d_W^{[l]}} d\mathbf{Z}_{0,ij}^{[l]} \times \mathbf{B}_{ij}^{[l-1]}$$

- $\mathbf{B}_{ij}^{[l-1]}$: The part of $\mathbf{X}^{[l-1]}$ used to obtain $\mathbf{Z}_{0,ij}^{[l]}$



Backpropagation for CONV

1. Consider the general case where there exist $d_C^{[l]}$ channels in the l th layer



$$\mathbf{X}^{[l-1]} \in \mathbb{R}^{d_C^{[l-1]} \times d_H^{[l-1]} \times d_W^{[l-1]}}$$

$$\mathbf{Z}_0^{[l]} \in \mathbb{R}^{d_C^{[l]} \times d_H^{[l]} \times d_W^{[l]}}$$

$$\mathbf{Z}_1^{[l]} \in \mathbb{R}^{d_C^{[l]} \times d_H^{[l]} \times d_W^{[l]}}$$

$$\mathbf{W}^{[l]} \in \mathbb{R}^{d_C^{[l]} \times d_C^{[l-1]} \times f^{[l]} \times f^{[l]}}$$

$$\mathbf{b}^{[l]} \in \mathbb{R}^{d_C^{[l]} \times 1 \times 1}$$

$$\mathbf{A}^{[l]} \in \mathbb{R}^{d_C^{[l]} \times d_H^{[l]} \times d_W^{[l]}}$$

$$d\mathbf{W}^{[l]} = ?$$

$$d\mathbf{b}_c^{[l]} = \sum_{i=1}^{d_H^{[l]}} \sum_{j=1}^{d_W^{[l]}} dZ_{1,cij}^{[l]} \quad (c = 1, \dots, d_C^{[l]})$$

$$d\mathbf{X}^{[l-1]} = ?$$

$$d\mathbf{Z}_0^{[l]} = d\mathbf{Z}_1^{[l]}$$

$$d\mathbf{Z}_1^{[l]} = \sigma^{[l]'}(\mathbf{Z}^{[l]}) \circ d\mathbf{A}^{[l]}$$

Backpropagation for CONV

1. Think about where the information is contained when calculating derivatives
2. The information about $\mathbf{X}^{[l-1]}$ is contained in every channel of $\mathbf{Z}_0^{[l]}$
3. Initialize $d\mathbf{X}^{[l-1]}$ as a zero matrix of the same dimension as $\mathbf{X}^{[l-1]}$
4. Informally, obtain

$$d\mathbf{X}_{\text{slice},ij}^{[l-1]} + = \sum_{c=1}^{d_C^{[l]}} d\mathbf{Z}_{0,ijc}^{[l]} \times \mathbf{W}_c^{[l]}$$

- $\mathbf{X}_{\text{slice},ij}^{[l-1]}$: The part of $\mathbf{X}^{[l-1]}$ used to obtain $\mathbf{Z}_{0,ijc}^{[l]}$ for $c = 1, \dots, d_C^{[l]}$
- $\mathbf{W}_c^{[l]}$: The c th kernel in the l th layer

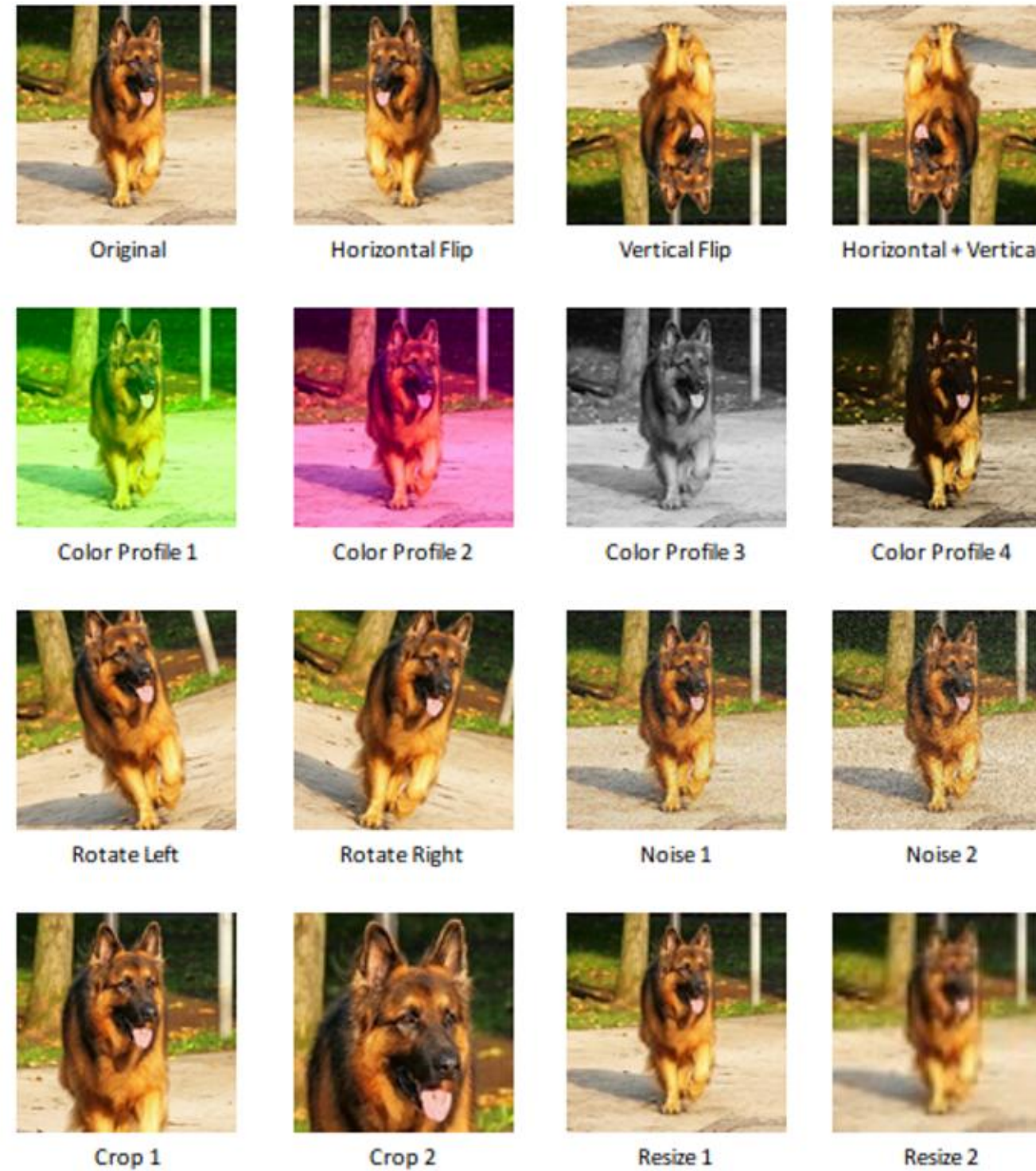
Backpropagation for CONV

1. It is obvious that every element in $\mathbf{Z}_{0,c}^{[l]}$ contains information of $\mathbf{W}_c^{[l]}$
for $c = 1, \dots, d_C^{[l]}$
2. However, $\mathbf{W}_c^{[l]}$ is not used for calculation of other channels
3. Thus, we have

$$d\mathbf{W}_{\textcolor{red}{c}}^{[l]} = \sum_{i=1}^{d_H^{[l]}} \sum_{j=1}^{d_W^{[l]}} d\mathbf{Z}_{0,\textcolor{red}{c}ij}^{[l]} \times \mathbf{B}_{ij}^{[l-1]} \quad (c = 1, \dots, d_C^{[l]})$$

- $\mathbf{B}_{ij}^{[l-1]}$: The part of $\mathbf{X}^{[l-1]}$ used to obtain $\mathbf{Z}_{0,\textcolor{red}{c}ij}^{[l]}$

Data augmentation



Understanding convolution networks

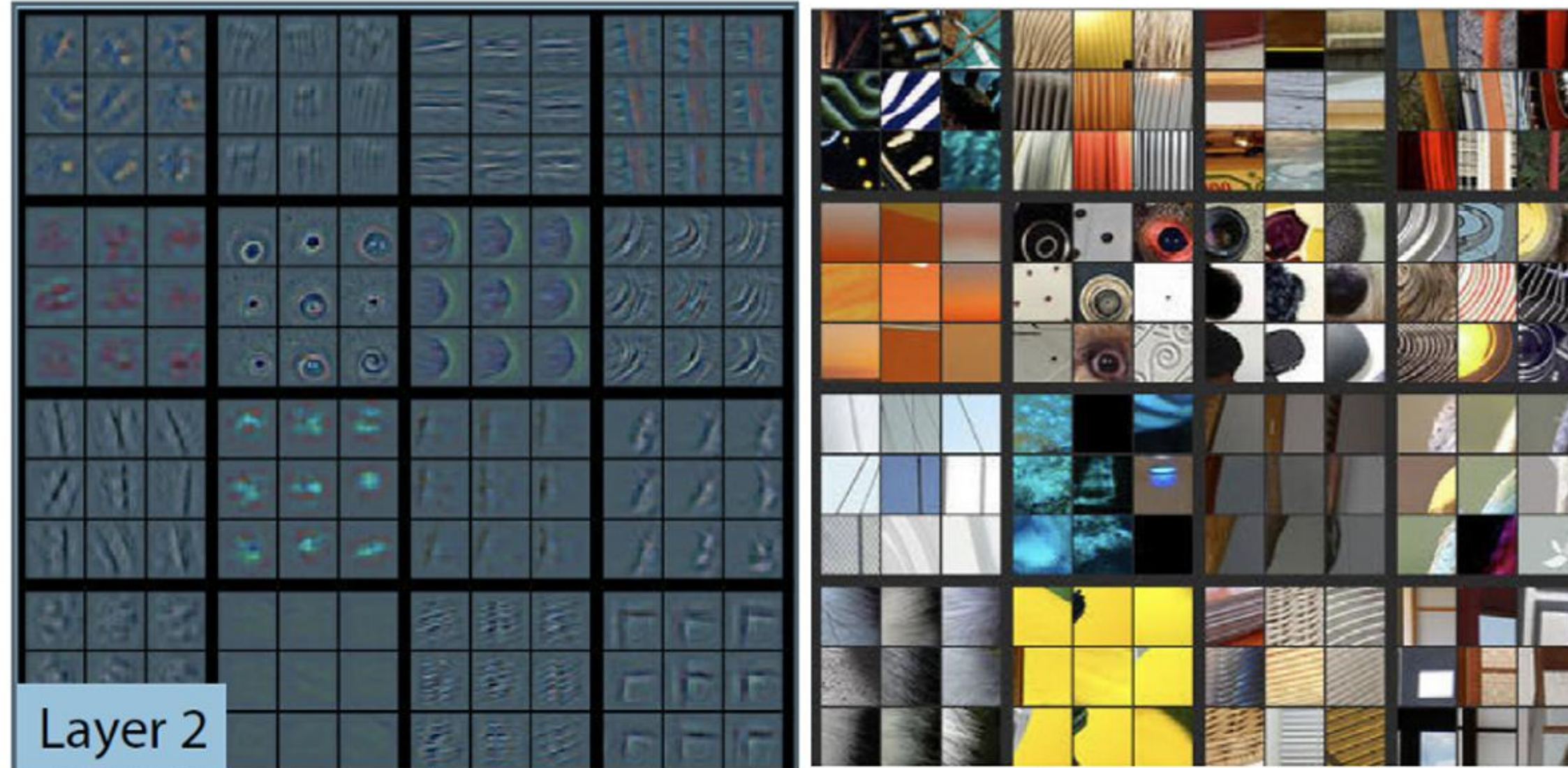


Layer 1



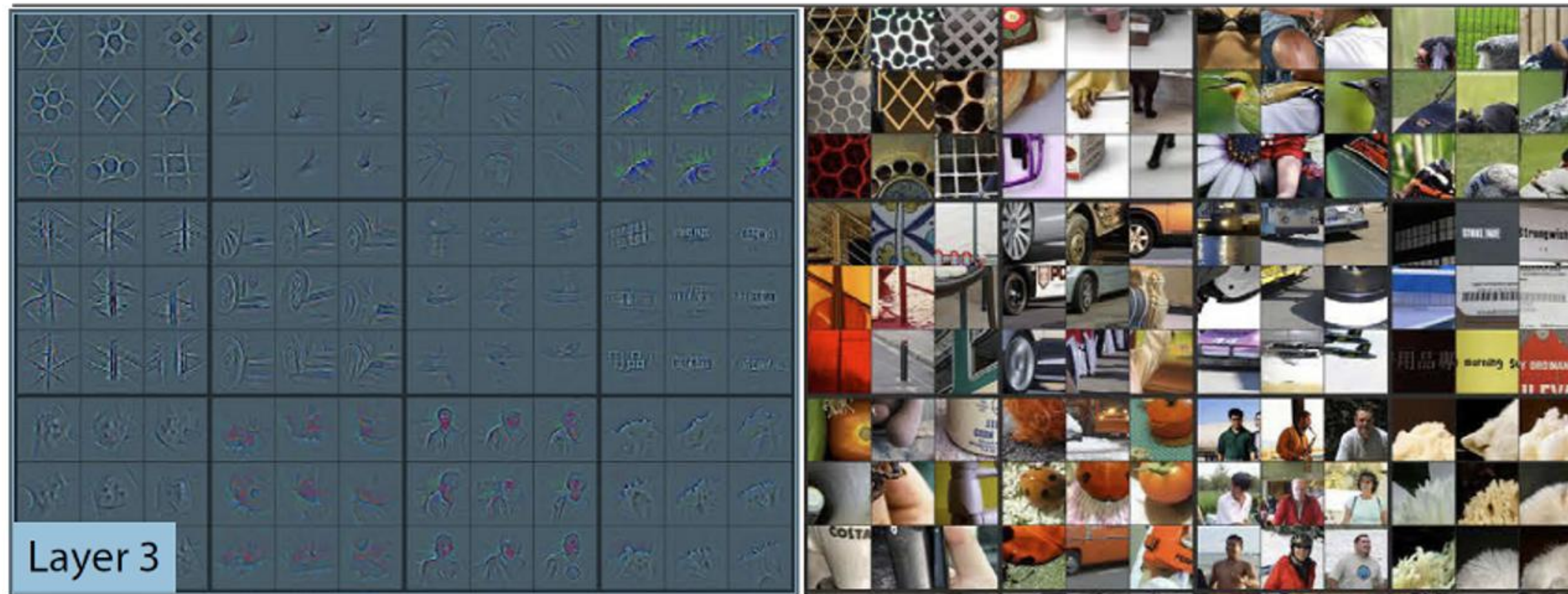
[Zeiler, M. D., & Fergus, R. (2014, September). Visualizing and understanding convolutional networks. In European conference on computer vision (pp. 818-833). Springer, Cham]

Understanding convolution networks



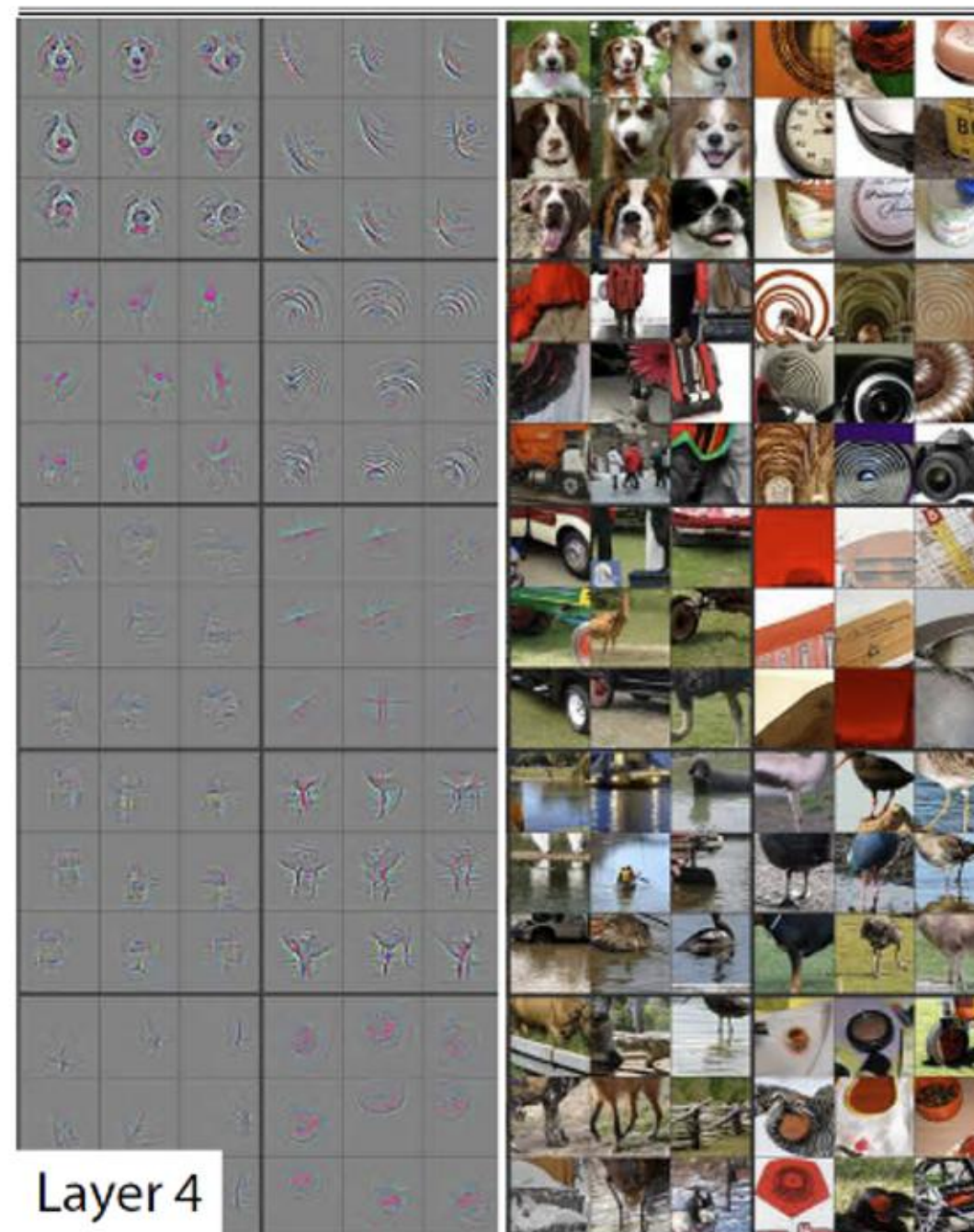
[Zeiler, M. D., & Fergus, R. (2014, September). Visualizing and understanding convolutional networks. In European conference on computer vision (pp. 818-833). Springer, Cham]

Understanding convolution networks



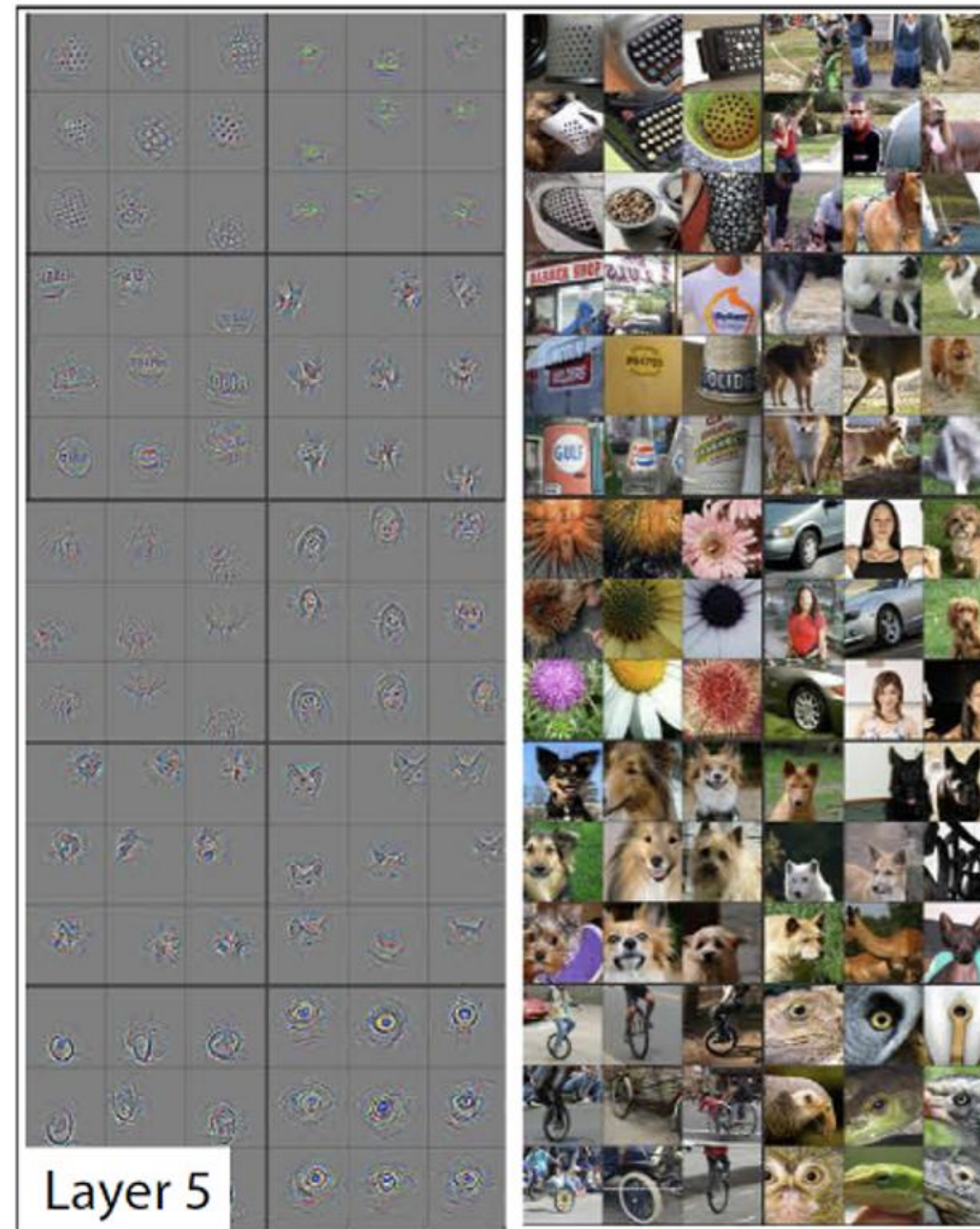
[Zeiler, M. D., & Fergus, R. (2014, September). Visualizing and understanding convolutional networks. In European conference on computer vision (pp. 818-833). Springer, Cham]

Understanding convolution networks



[Zeiler, M. D., & Fergus, R. (2014, September). Visualizing and understanding convolutional networks. In European conference on computer vision (pp. 818-833). Springer, Cham]

Understanding convolution networks



[Zeiler, M. D., & Fergus, R. (2014, September). Visualizing and understanding convolutional networks. In European conference on computer vision (pp. 818-833). Springer, Cham]

Summary

1. Motivation: Spatial structure, locality, and parameter sharing
2. Convolution and pooling layer: Extraction of local information
3. CNN exhibits a hierarchical organization to extract information from simple to complex