

Introduction to CNN & RNN

2020310634 수학과 이호준

Contents

0. Review

1. CNN

1-1. Convolution Layer

1-2. Pooling Layer

1-3. Structure of CNN

2. RNN

2-1. Time Serise Data

2-2. RNN Layer

2-3. Structure of RNN

3. Comparison

0. Review

0. Review

Data

$$\begin{aligned}x^{(1)} &= [x_1^{(1)} \quad \dots \quad x_n^{(1)}] \\x^{(2)} &= [x_1^{(2)} \quad \dots \quad x_n^{(2)}] \\&\vdots \\x^{(d)} &= [x_1^{(d)} \quad \dots \quad x_n^{(d)}]\end{aligned}$$

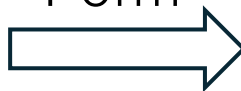
Bias

$$b = [b^{(1)} \quad \dots \quad b^{(m)}]$$

Weights

$$\begin{aligned}w^{(1)} &= \begin{bmatrix} w_1^{(1)} \\ \vdots \\ w_n^{(1)} \end{bmatrix} \\w^{(2)} &= \begin{bmatrix} w_1^{(2)} \\ \vdots \\ w_n^{(2)} \end{bmatrix} \\&\vdots \\w^{(m)} &= \begin{bmatrix} w_1^{(m)} \\ \vdots \\ w_n^{(m)} \end{bmatrix}\end{aligned}$$

Matrix
Form



$$X = \begin{bmatrix} x^{(1)} \\ \vdots \\ x^{(d)} \end{bmatrix} = \begin{bmatrix} x_1^{(1)} & \dots & x_n^{(1)} \\ \vdots & \ddots & \vdots \\ x_1^{(d)} & \dots & x_n^{(d)} \end{bmatrix}$$

$$W = \begin{bmatrix} w_1^{(1)} & \dots & w_1^{(m)} \\ \vdots & \ddots & \vdots \\ w_n^{(1)} & \dots & w_n^{(m)} \end{bmatrix}$$

$$b = [b^{(1)} \quad \dots \quad b^{(m)}]$$

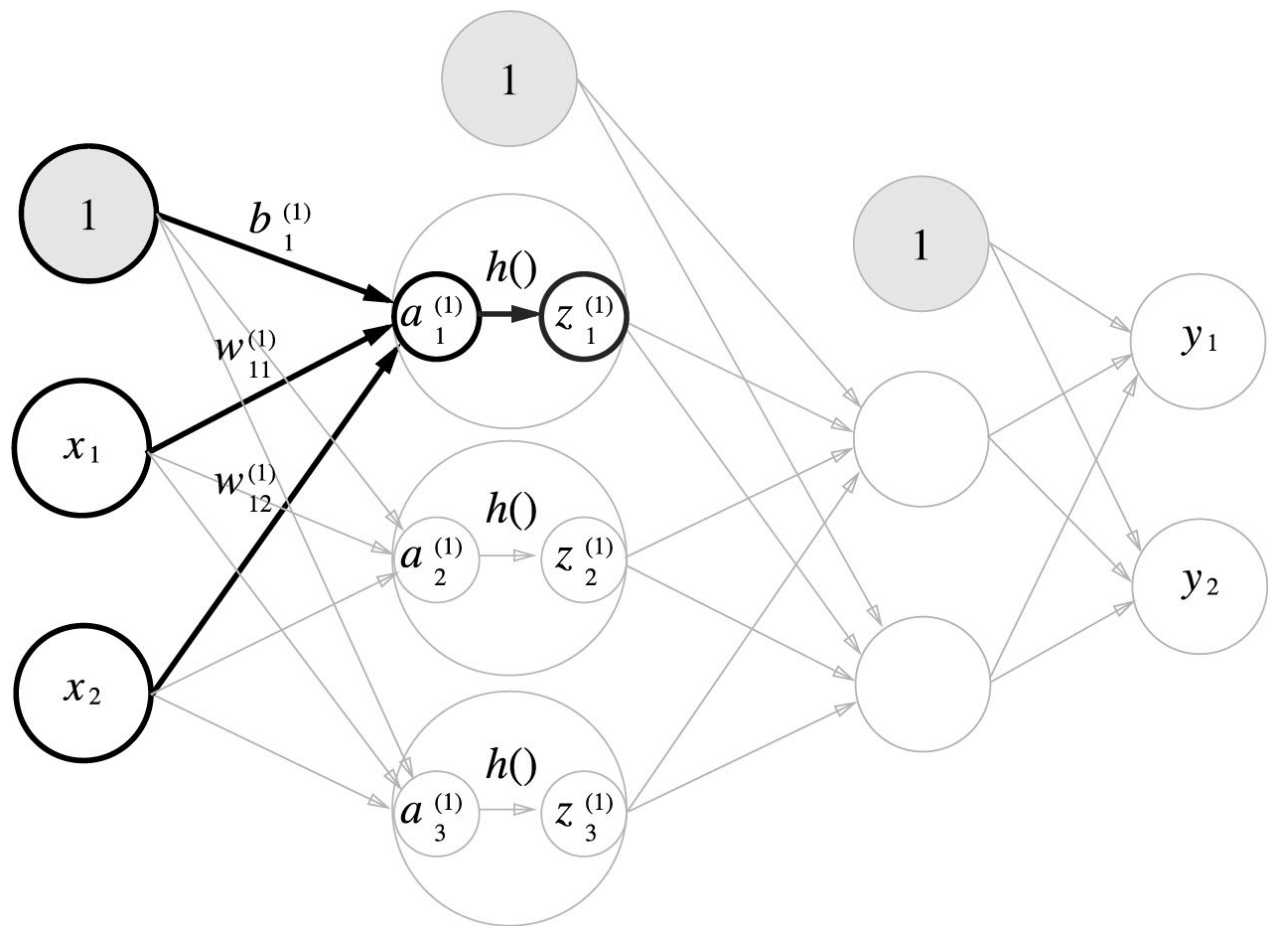
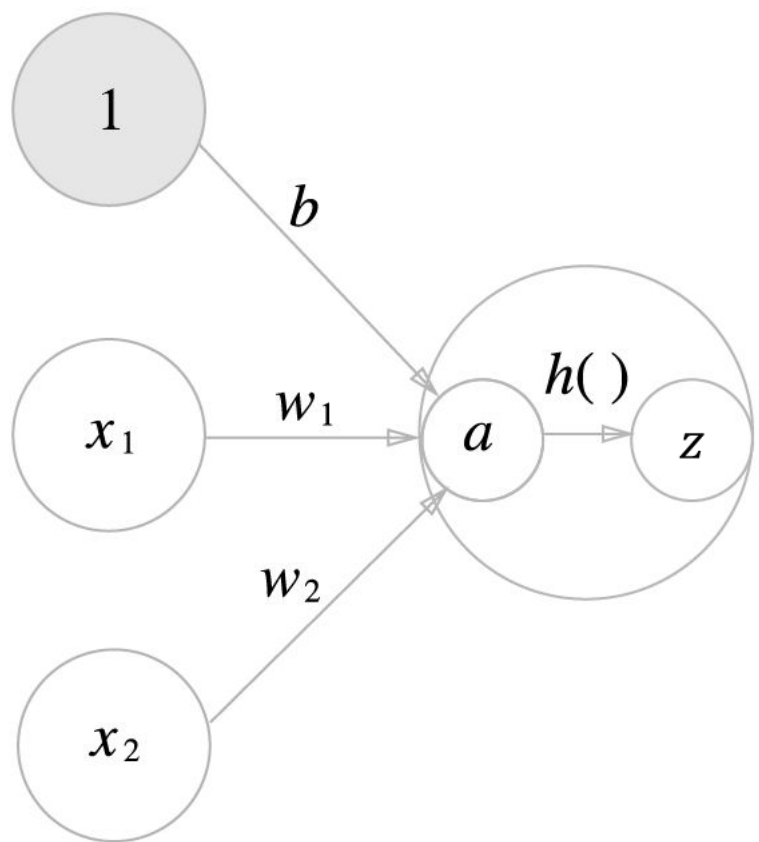
0. Review

Affine Transform

$$A = XW + b$$

Activation Functions

- Sigmoid: $f(x) = \frac{1}{1 + e^{-x}}$
- ReLU: $f(x) = \max\{0, x\}$
- Softmax: $f(x_i) = \frac{e^{x_i}}{\sum_j e^{x_j}}$



0. Review

Loss function

$$L = \frac{1}{2} \sum_k (y_k - t_k)^2$$

Sum of squares for error (SSE)

$$L = - \sum_k t_k \log y_k$$

Cross entropy error (CEE)

Gradient Descent

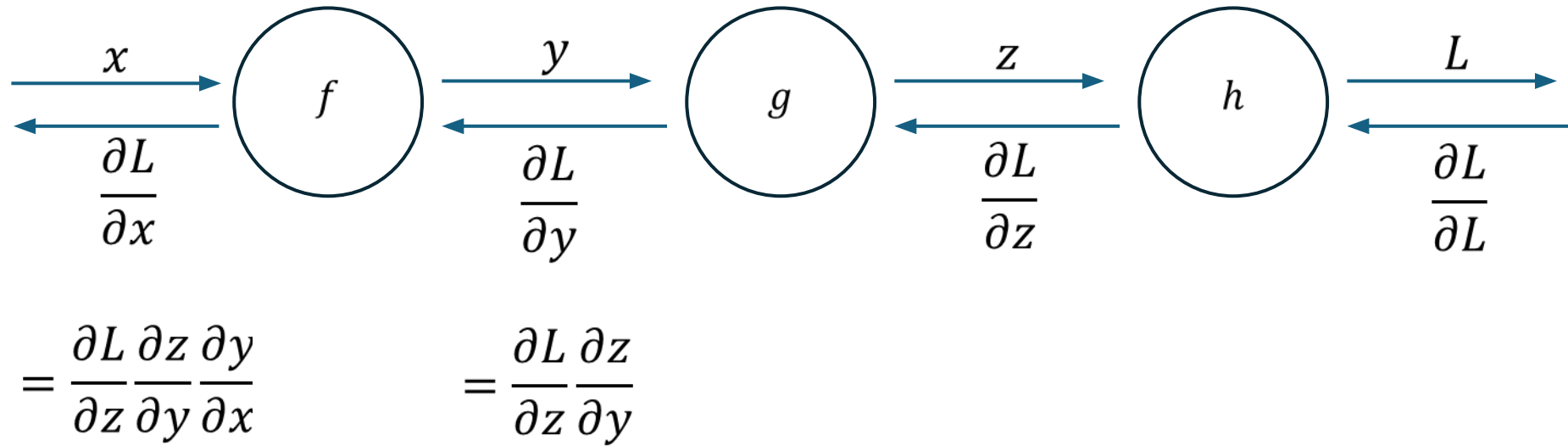
$$W = W - \alpha \frac{\partial L}{\partial W}$$

α : learning rate

L : Loss function

0. Review

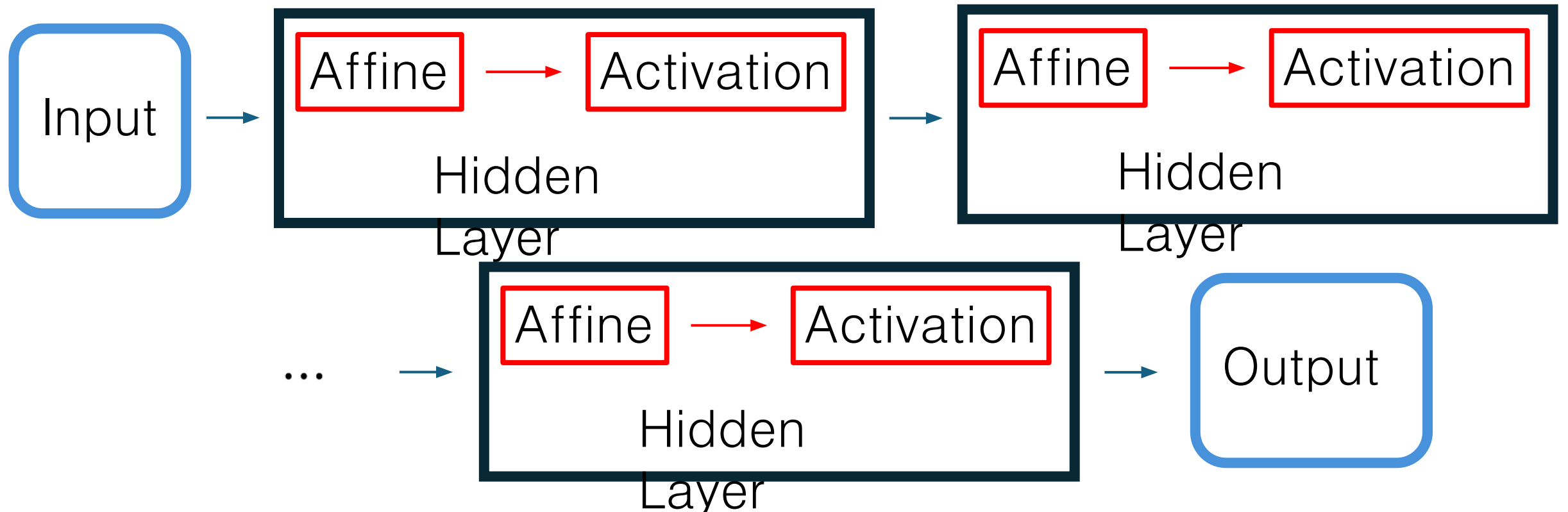
Backpropagation



0. Review

Fully Connected Neural Network (or MLP: Multi-Layer Perceptron)

Affine Layer → Activation Function



1. CNN

Convolutional Neural Network

1-1. Convolution Layer

Convolution (1-dim)

$$(y * w)(t) = \sum_{\tau=-\infty}^{\infty} y(\tau)w(t - \tau)$$

$$(y * w)(t) = \int_{-\infty}^{\infty} y(\tau)w(t - \tau)d\tau$$



We use (discrete) Convolution
(2-dim)

$$(X * W)(i, j) = \sum_{u=0}^{FH-1} \sum_{v=0}^{FW-1} X(i + u, j + v)W(u, v)$$

1-1. Convolution Layer

Example: one matrix X and one W

1	2	3
0	3	5
2	1	0

X

*

2	1
0	1

W

=

1-1. Convolution Layer

Example: one matrix X and one W

1	2	3
0	3	5
2	1	0

X

*

2	1
0	1

W

=

1-1. Convolution Layer

Example: one matrix X and one W

1	2	3
0	3	5
2	1	0

X

*

2	1
0	1

W

=

1-1. Convolution Layer

Example: one matrix X and one W

1	2	3
0	3	5
2	1	0

X

*

2	1
0	1

W

=

1-1. Convolution Layer

Padding

	1	2		
	2	3		

X

*

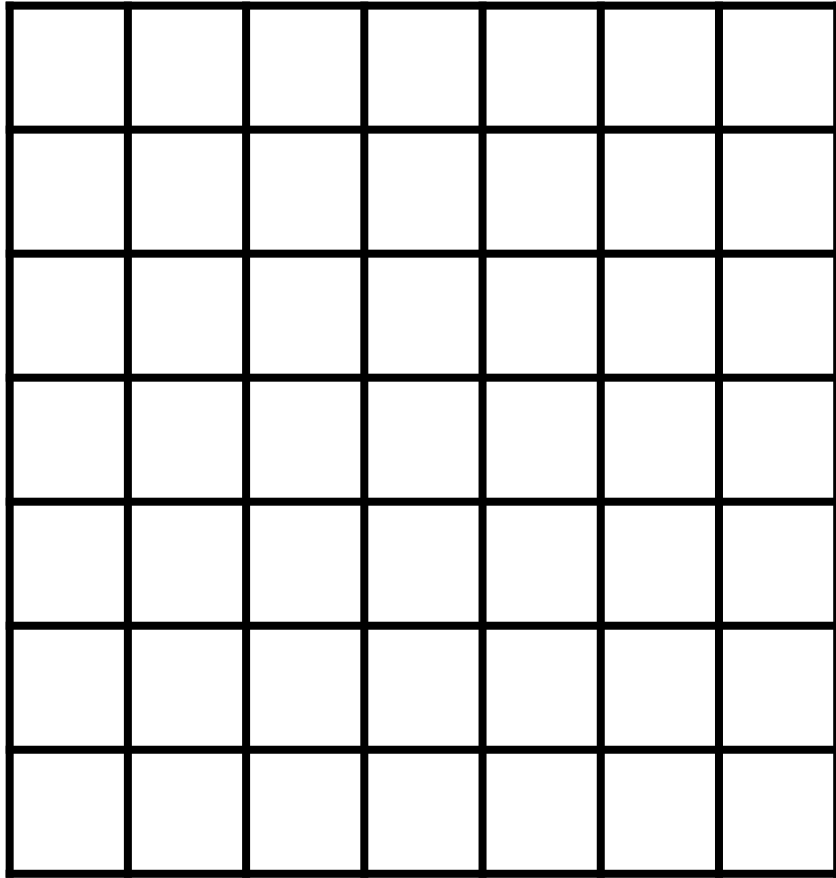
1	2	3
0	1	0
0	2	4

W

=

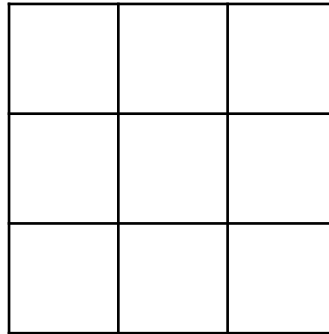
1-1. Convolution Layer

Stride



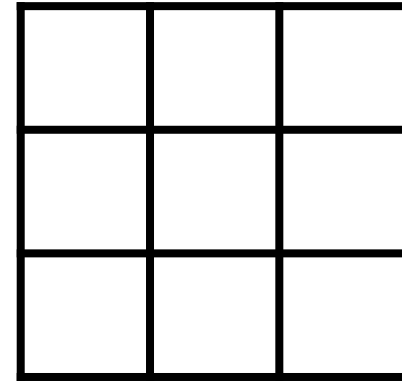
X

*



W

=



1-1. Convolution Layer

Size of the Output

$$(X * W)(i, j) = \sum_{u=0}^{FH-1} \sum_{v=0}^{FW-1} X(i+u, j+v)W(u, v)$$

$$OH = \frac{H + 2P - FH}{S} + 1$$

$$OW = \frac{W + 2P - FW}{S} + 1$$

X is $H \times W$ matrix

W is $FH \times FW$ matrix

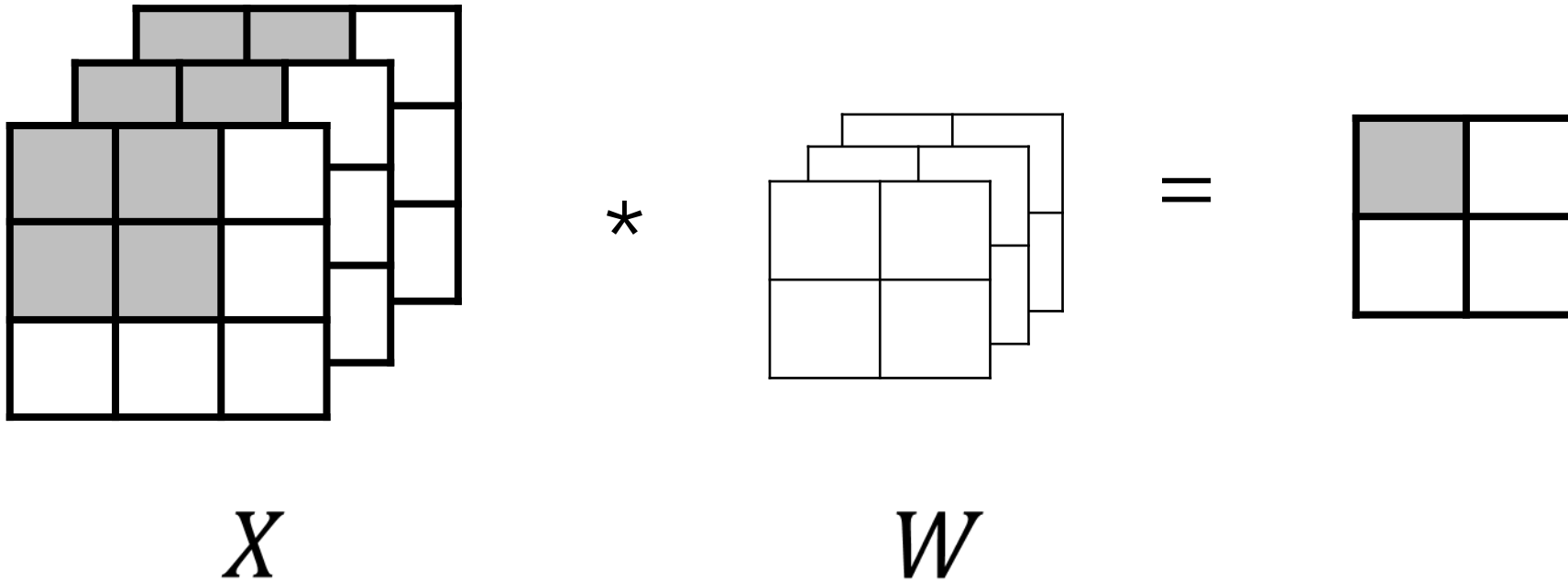
P is Padding

S is Stride.

1-1. Convolution Layer

Increasing the Number of Channels

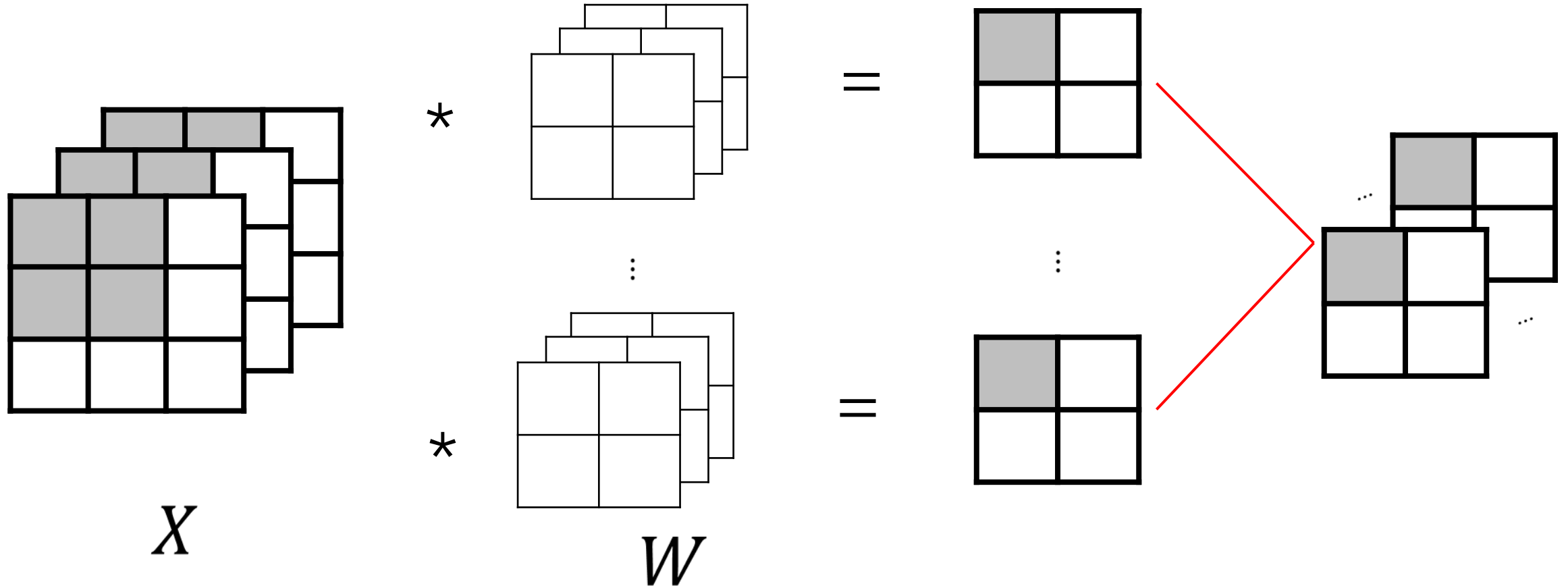
$$(X * W)(i, j) = \sum_{c=0}^{C-1} \sum_{u=0}^{FH-1} \sum_{v=0}^{FW-1} X(c, i+u, j+v) W(c, u, v)$$



1-1. Convolution Layer

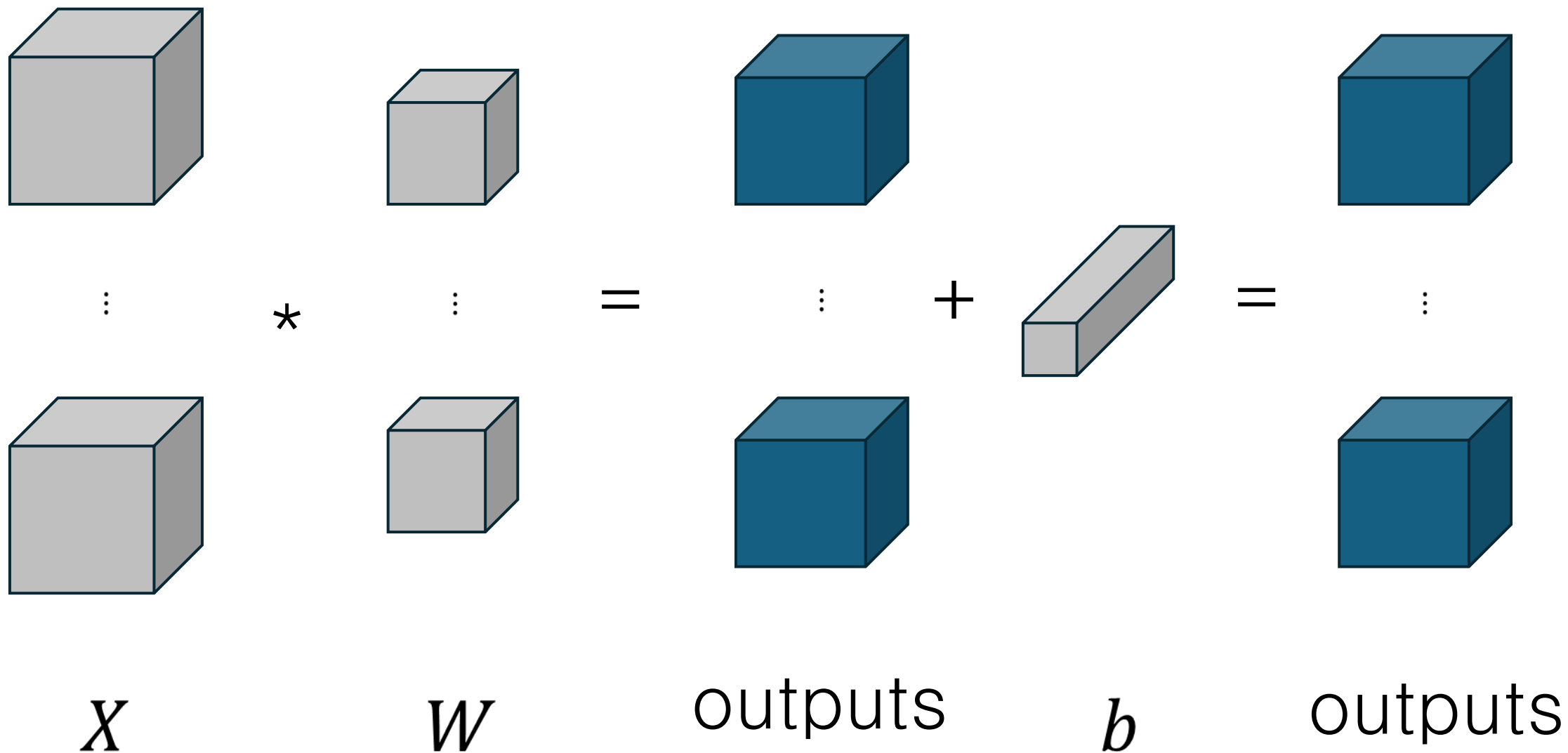
Increasing the Number of Filters

FN-times $(X * W)(i, j) = \sum_{c=0}^{C-1} \sum_{u=0}^{FH-1} \sum_{v=0}^{FW-1} X(c, i+u, j+v) W(c, u, v)$



1-1. Convolution Layer

Batch Processing



1-2. Pooling Layer

Pooling

- Max Pooling
- Average Pooling

Mathematical effect

- Downsampling
- Translation Invariance

Ex) Pooling size 2x2

1	2	3	4
0	5	1	2
0	1	2	1
2	0	5	4

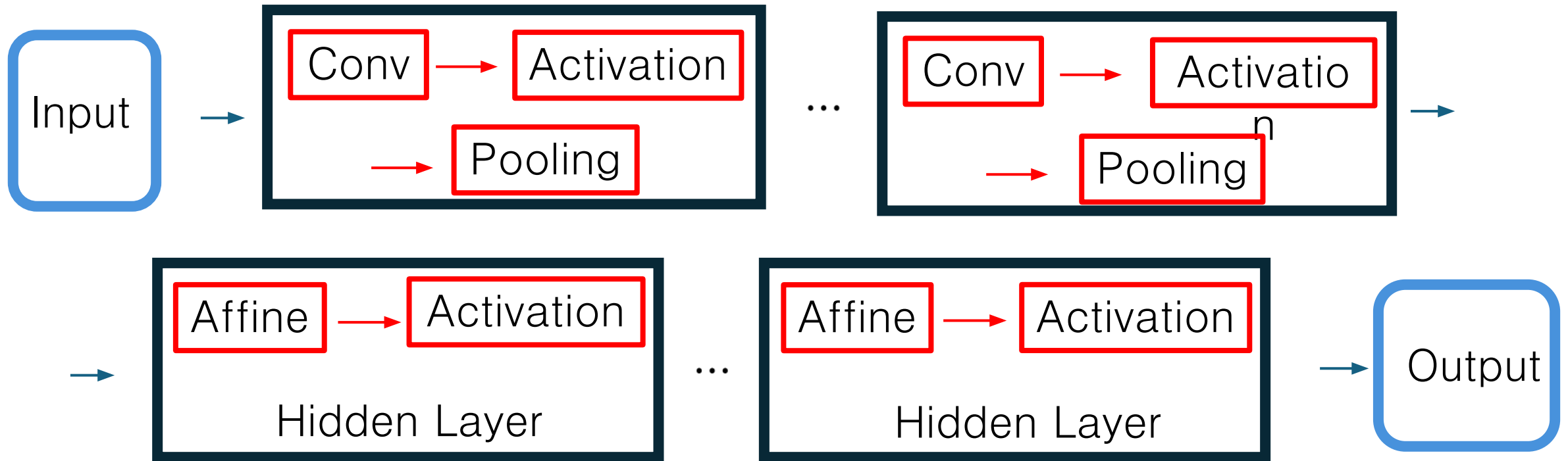
Max pooling
→

5	4
2	5

1-3. Structure of CNN

Convolutional Neural Network

Convolution Layer → Activation (→ Pooling): Feature Extraction
+ Affine → Activation: Classification



1–3. Structure of CNN

Gradient Descent

$$W = W - \alpha \frac{\partial L}{\partial W}$$

α : *learning rate*

L : *Loss function*

$$\frac{\partial L}{\partial W} = X * \frac{\partial L}{\partial Y}$$

1–3. Structure of CNN

Why do we use convolution?

Ex) 28 x 28 image

– MLP

Input data: 28 x 28 = 784

Output: 100 neurons

→ $784 \times 100 = 78,400$ weights + 100 bias

– CNN

Filter: 3 x 3 = 9 weights

of filter : 100

→ $9 \times 100 = 900$ parameters + 100 bias

of Parameter is reduced



입력 이미지



필터 1



세로 에지에
반응

출력 이미지 1



필터 2



가로 에지에
반응

출력 이미지 2

2. RNN

Recurrent **N**eural **N**etwork

2-1. Time Series Data

Sentence data → time series data

“MIMIC is math interchange club.”

“Sungkyunkwan University (SKKU), located in Seoul, South Korea, is a prestigious institution with a history dating back over 600 years. It combines traditional values with modern research, offering a wide range of programs in sciences, humanities, engineering, and business. SKKU is known for its strong academic reputation, innovative research, and global partnerships.”

Dividing sentence into words (Tokenization)

→ Integer Encoding (Index)

→ Embedding Layer

→ RNN input data

What is next word?

2-1. Time Series Data

Word $\rightarrow$ vector (Embedding Layer)

MIMIC / is / math / interchange / club / .

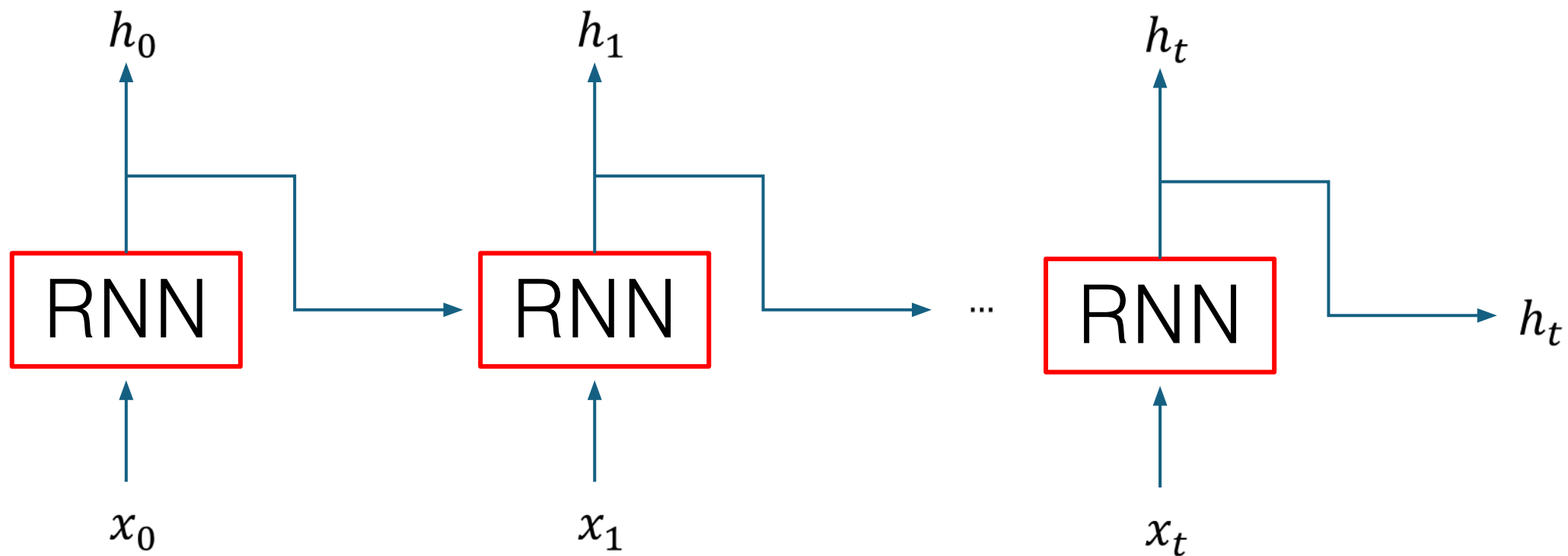
$\textcircled{0}$ 1 2 3 4 5 $\rightarrow$ indices



$$\text{Ex) } E = \begin{bmatrix} 0.61 & \dots & 0.52 \\ \vdots & \ddots & \vdots \\ 0.14 & \dots & 0.23 \end{bmatrix} \text{ where } E \text{ is } V \times d \text{ matrix}$$

We will learn matrix E with gradient descent.

2-2. RNN Layer



2-2. RNN Layer

$$h_t = \varphi(h_{t-1}W_h + x_tW_x + b_h) \quad \varphi : \text{tanh or ReLU}$$



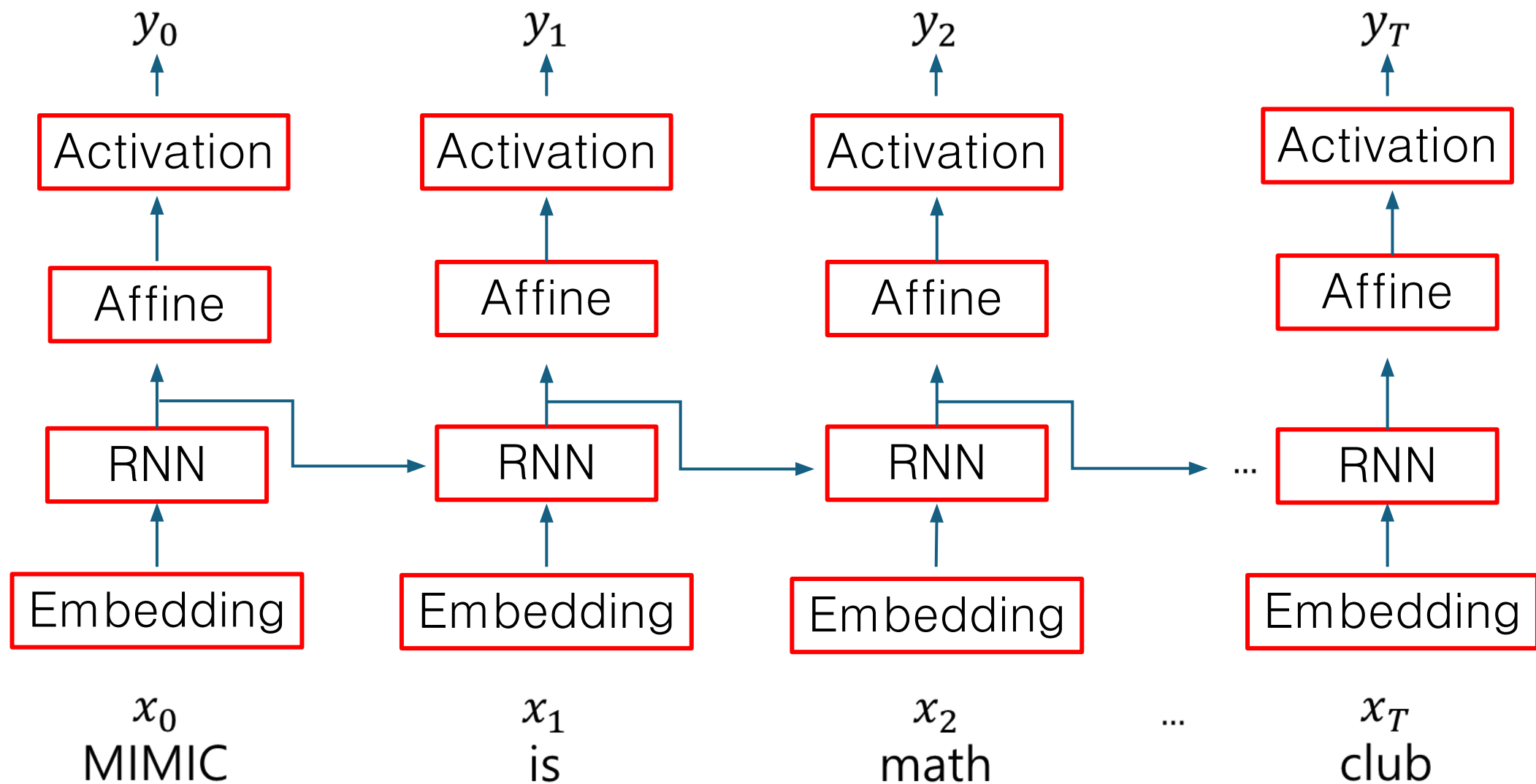
$$h_0 = \varphi(x_0W_x + b_h)$$

$$h_1 = \varphi(h_0W_h + x_1W_x + b_h) \quad \text{The same wights } W_x, W_h$$

$\vdots$

$$h_T = \varphi(h_{T-1}W_h + x_TW_x + b_h)$$

2-3. Structure of RNN



2-3. Structure of RNN

$$y_t = g(W_y h_t + b_y)$$

$g : \text{softmax or sigmoid}$



$$y_0 = g(W_y h_0 + b_y)$$

$$y_1 = g(W_y h_1 + b_y)$$

$\vdots$

$$y_T = g(W_y h_T + b_y)$$

The same weight W_y

2-3. Structure of RNN

Gradient Descent

$$W_h = W_h - \alpha \frac{\partial L}{\partial W_h}$$

$$W_x = W_x - \alpha \frac{\partial L}{\partial W_x}$$

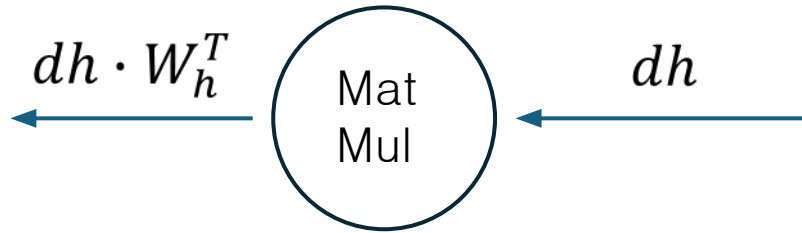
$$W_y = W_y - \alpha \frac{\partial L}{\partial W_y}$$

α : learning rate

L : Loss function

2-3. Structure of RNN

Exploding Gradients & Vanishing Gradients



$$W_h = W_h - \alpha \frac{\partial L}{\partial W_h}$$

 LSTM, Gradients clipping

3. Comparison

3. Comparison

	CNN	RNN
Data	Spatially structured data	Sequential or time-series data
Information Flow	One-directional	Recurrent
Structure	Repeated block cells [Conv - Act – (Pooling)]	Repeated RNN cell
Main Application	Image classification Object detection Video processing	Natural language processing Speech recognition Time-series forecasting
Parameter	The same filter -> all spatial inputs	The same weights -> all time steps

Q/A

Thank you