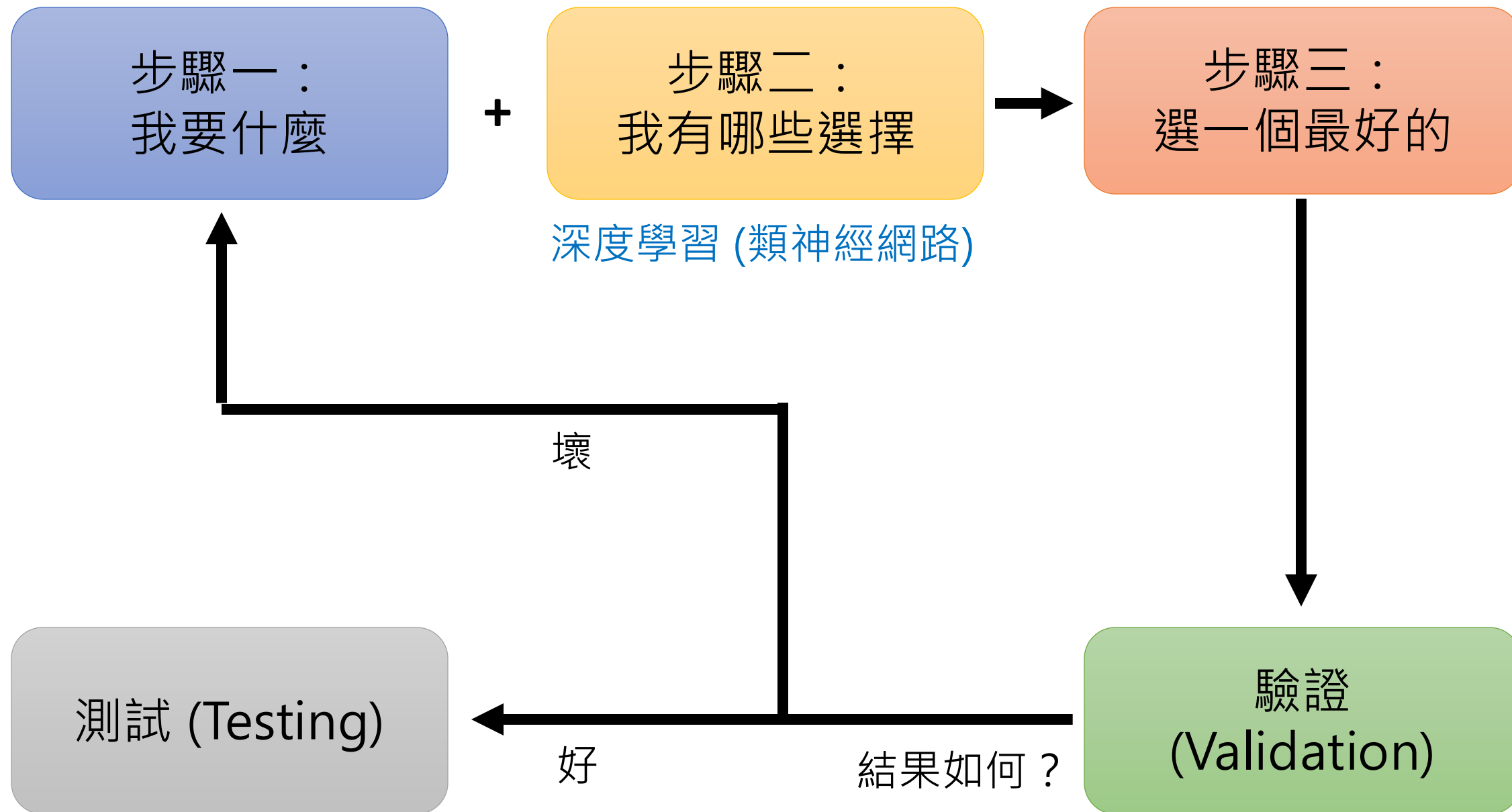




一堂課看懂 訓練類神經網路 的各種訣竅

李宏毅

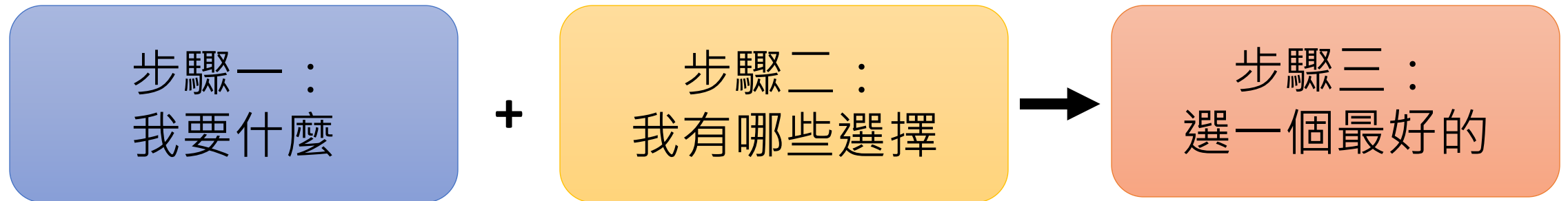




介紹各種訓練類神經網路常用技巧

- 聽到一個跟訓練類神經網路有關的方法時，你要這樣問自己

方法名	改了那一個步驟	帶來什麼好處
.....		



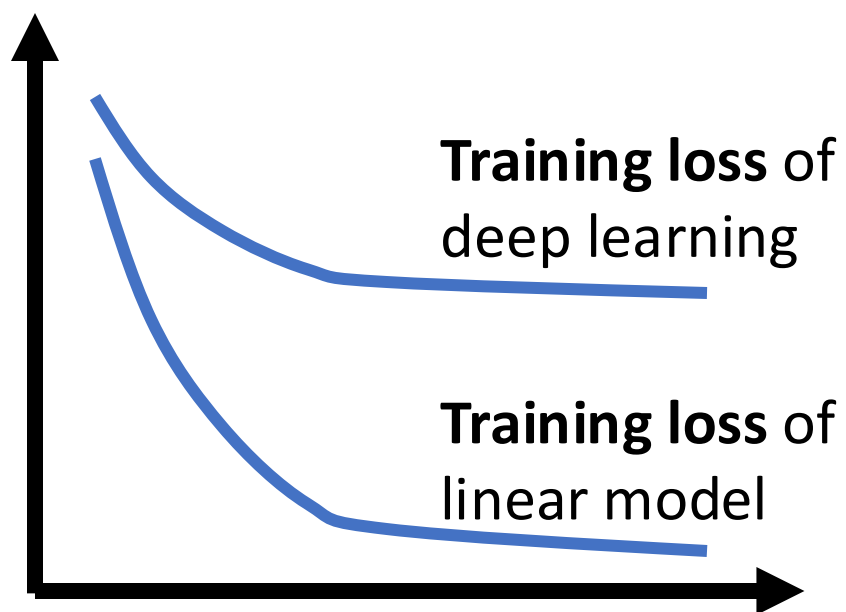
介紹各種訓練類神經網路常用技巧

- 聽到一個跟訓練類神經網路有關的方法時，你要這樣問自己

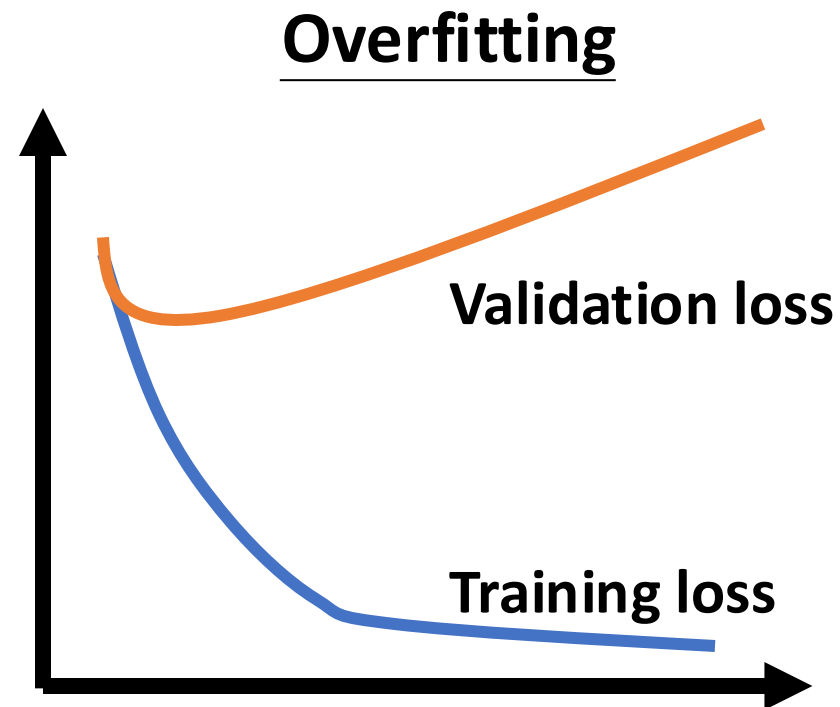
方法名	改了那一個步驟	帶來什麼好處
.....		

- Better Optimization：更低的 Training Loss
- Better Generalization：更低的 Validation Loss
-

選擇合適的技巧

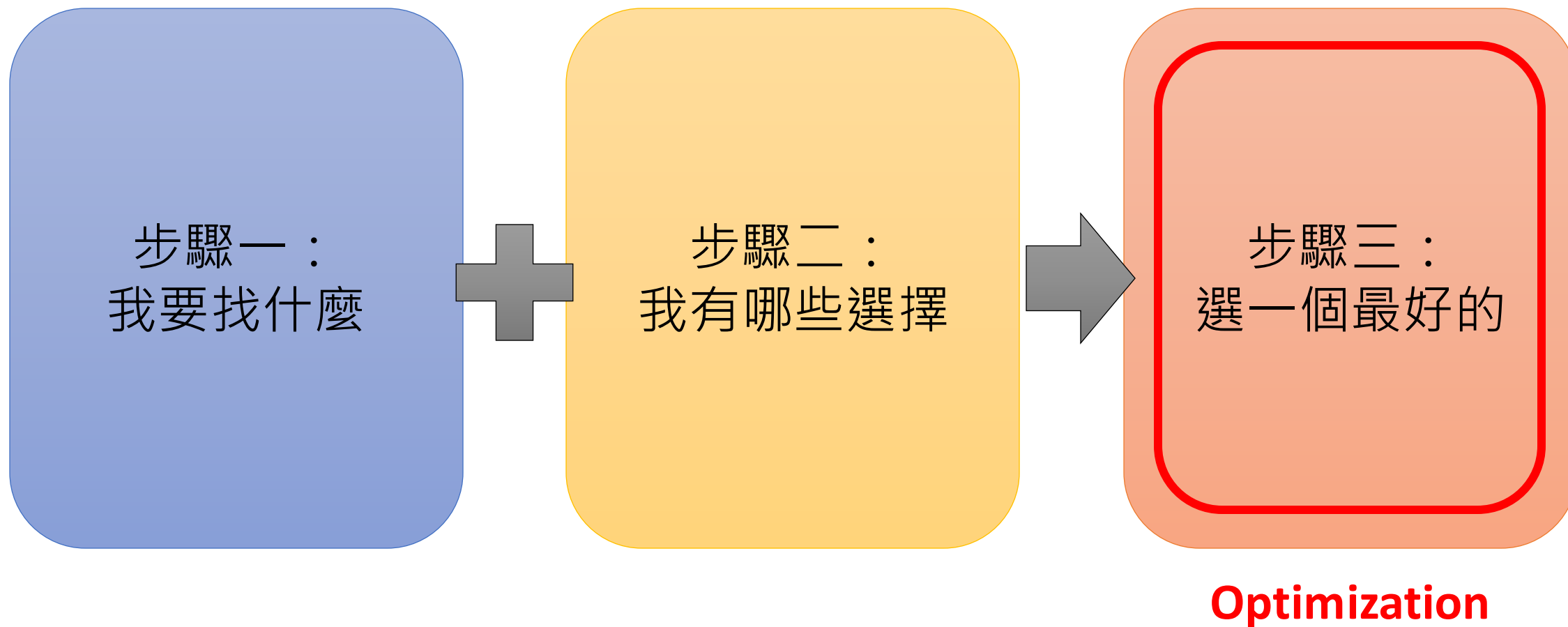


Better Optimization!



Better Generalization!

找函式步驟 3 + 1



Vanilla Gradient Descent

$$\boldsymbol{\theta}^* = \arg \min_{\boldsymbol{\theta}} L(\boldsymbol{\theta})$$

➤ (Randomly) Pick initial values $\boldsymbol{\theta}^0$

➤ Compute gradient $\boldsymbol{g}^0 = \nabla L(\boldsymbol{\theta}^0)$

$$\boldsymbol{\theta}^1 \leftarrow \boldsymbol{\theta}^0 - \eta \boldsymbol{g}^0$$

➤ Compute gradient $\boldsymbol{g}^1 = \nabla L(\boldsymbol{\theta}^1)$

$$\boldsymbol{\theta}^2 \leftarrow \boldsymbol{\theta}^1 - \eta \boldsymbol{g}^1$$

➤ Compute gradient $\boldsymbol{g}^2 = \nabla L(\boldsymbol{\theta}^2)$

$$\boldsymbol{\theta}^3 \leftarrow \boldsymbol{\theta}^2 - \eta \boldsymbol{g}^2$$

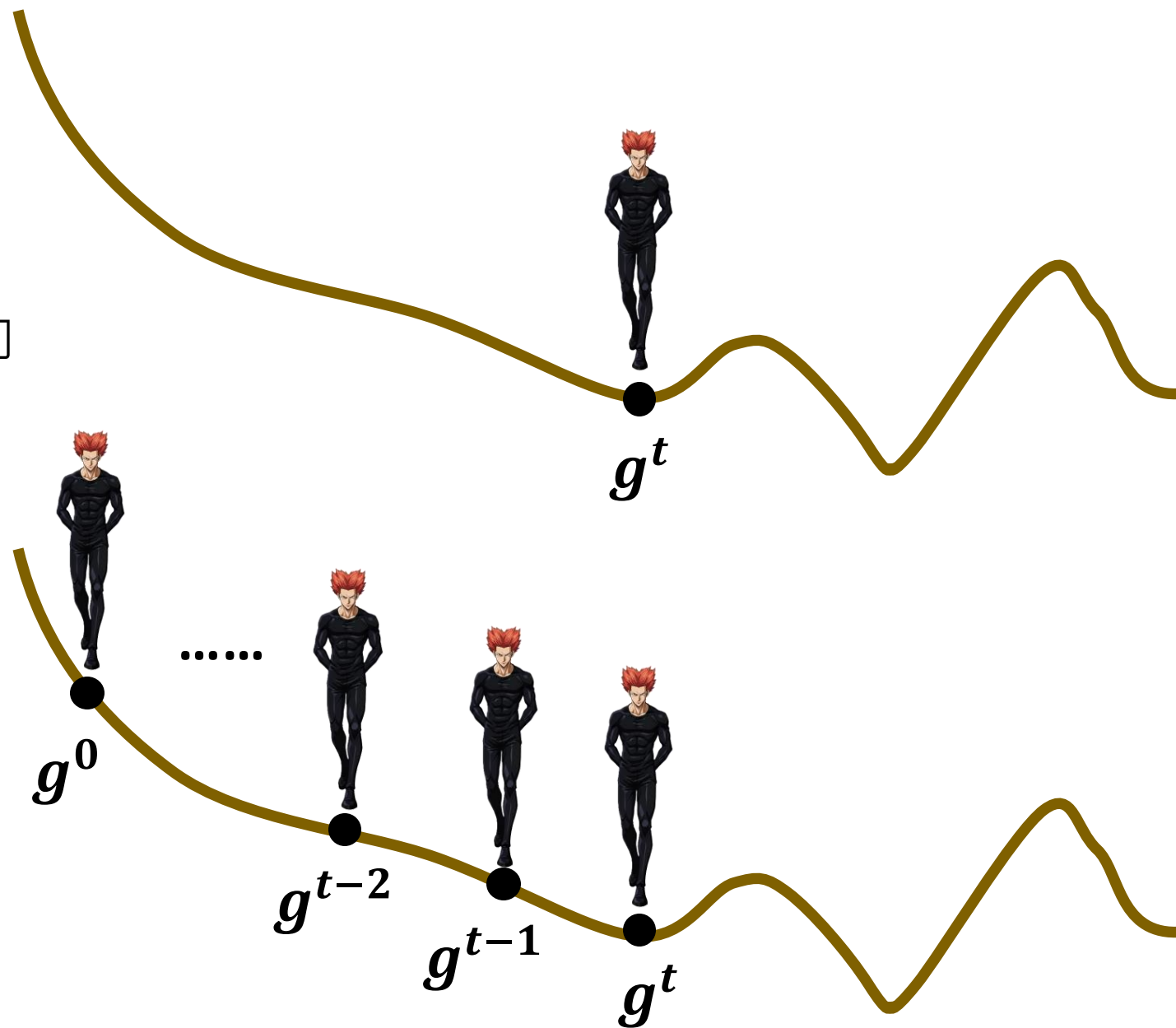
Optimizer

Vanilla Gradient Descent

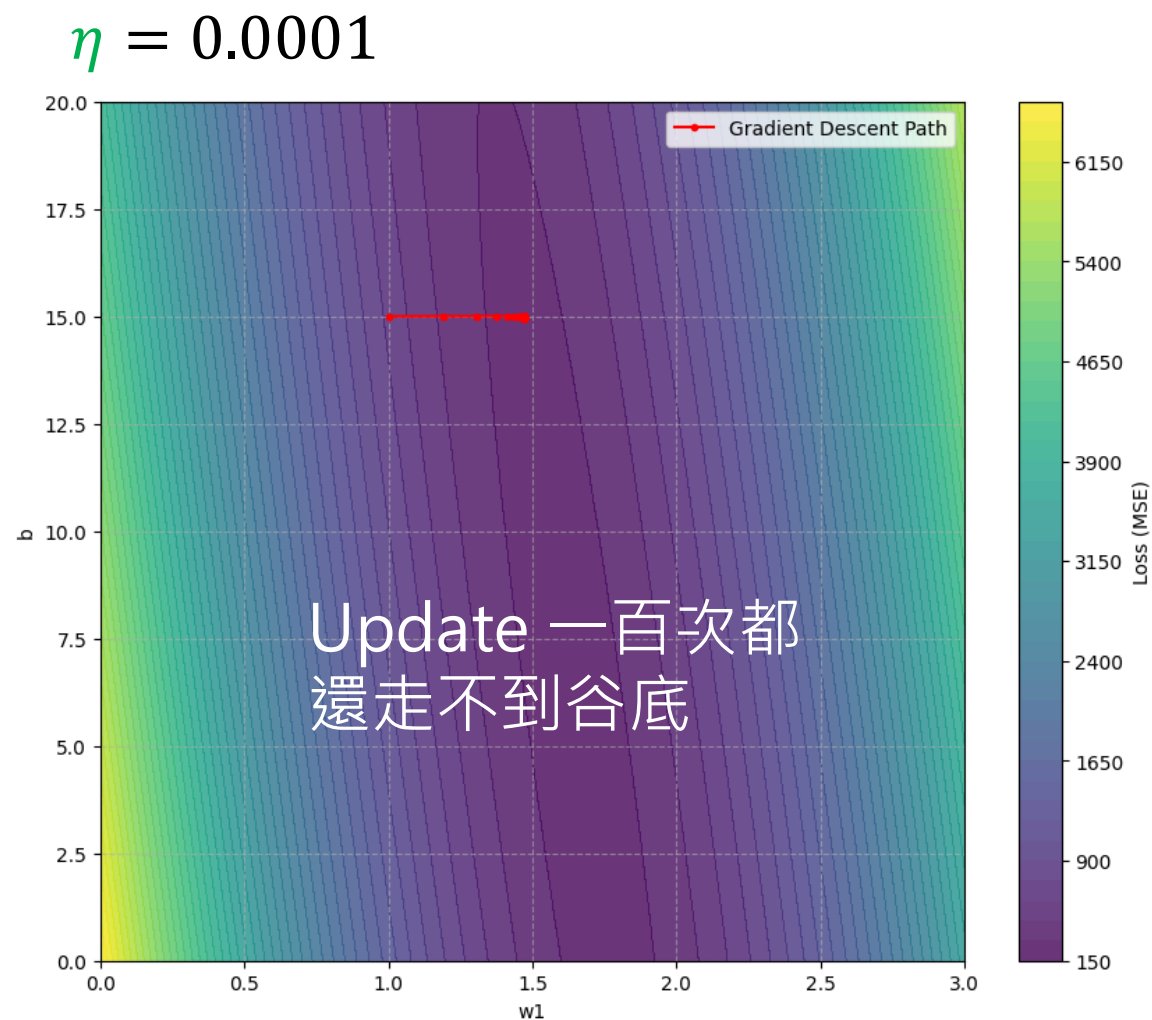
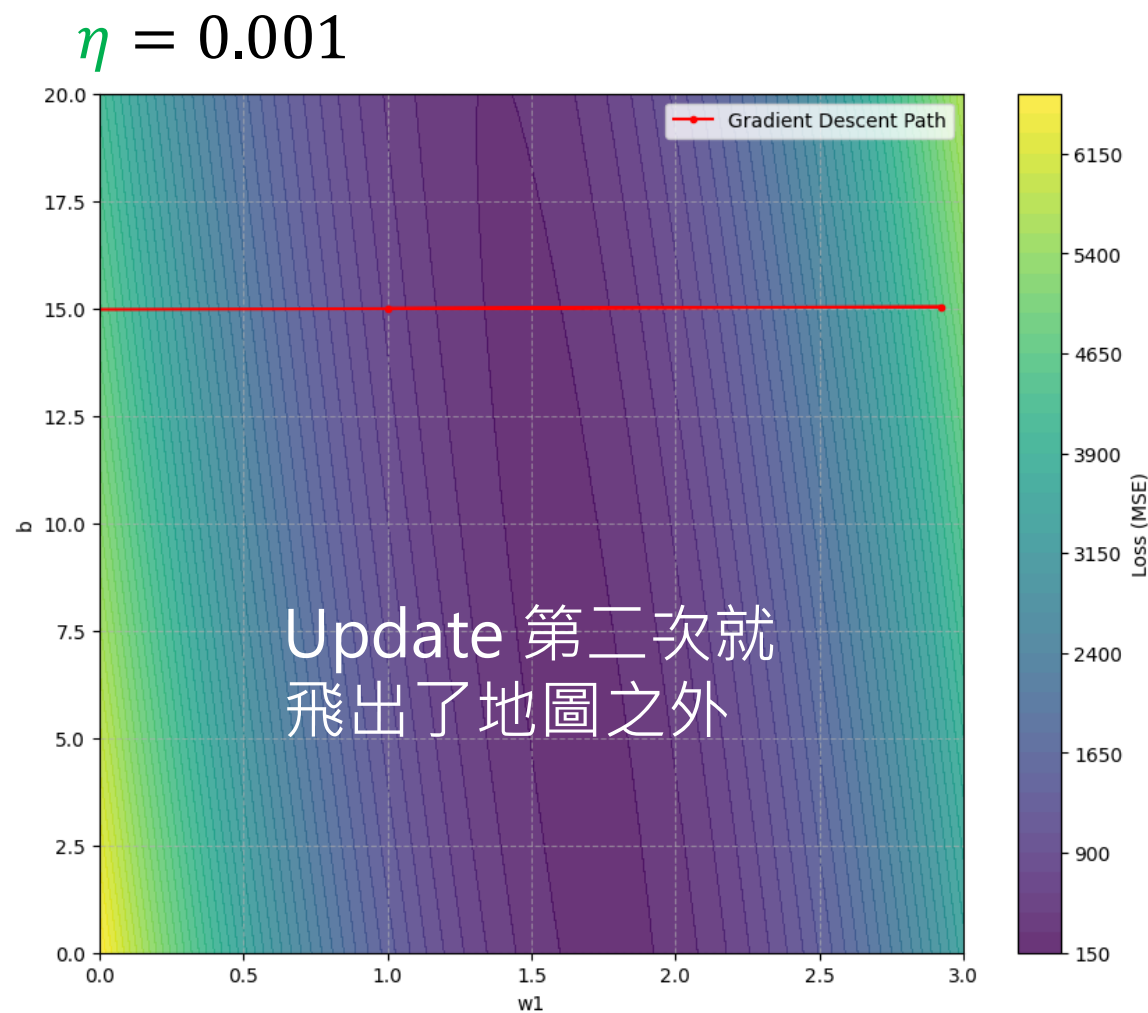
根據當下算出來的 g^t 來決定方向

Gradient Descent + Optimizer

根據 $g^0, g^1, g^2, \dots, g^t$
一起來決定方向

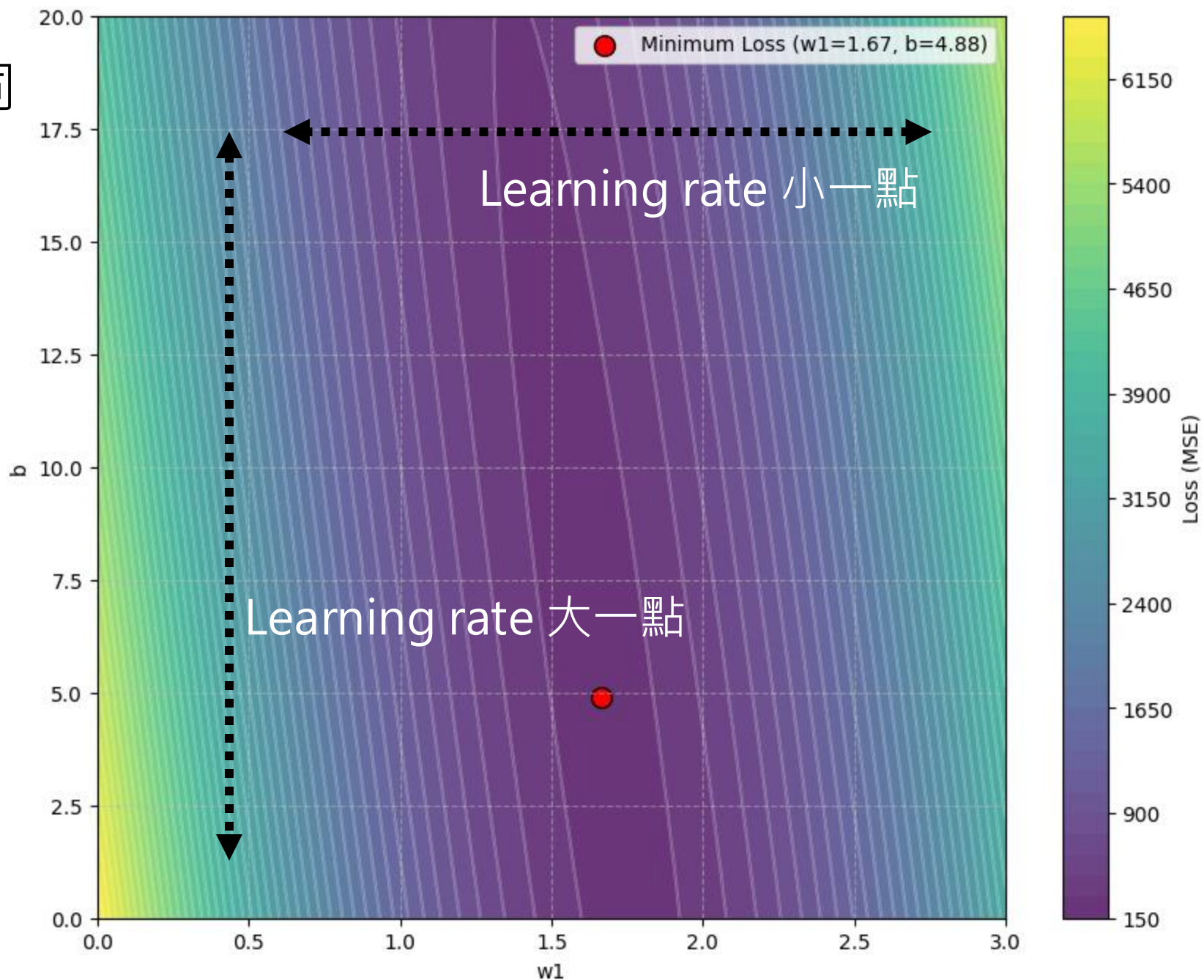


Learning Rate 實在是很難調 ...

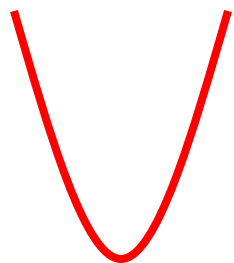


不同參數應該要有不同的 Learning Rate

怎麼知道那一個方向的 gradient 大、那一個小？



用過去的 Gradient 來決定 Learning Rate



Learning rate 小一點

$$g^0 = \begin{bmatrix} 500 \\ 0.4 \end{bmatrix}$$

$$g^1 = \begin{bmatrix} 432 \\ -0.3 \end{bmatrix}$$

$$g^2 = \begin{bmatrix} -211 \\ 0.2 \end{bmatrix}$$

$$g^3 = \begin{bmatrix} 139 \\ 0.1 \end{bmatrix}$$



Learning rate 大一點

Adagrad

- Compute gradient $\mathbf{g}^0 = \nabla L(\boldsymbol{\theta}^0)$

For each dimension i : $\sigma_i^0 = \sqrt{(g_i^0)^2} = |g_i^0|$

$$\theta_i^1 \leftarrow \theta_i^0 - \frac{\eta}{\sigma_i^0} g_i^0$$

- Compute gradient $\mathbf{g}^1 = \nabla L(\boldsymbol{\theta}^1)$

Average?

For each dimension i : $\sigma_i^1 = \sqrt{[(g_i^0)^2 + (g_i^1)^2]}$

$$\theta_i^2 \leftarrow \theta_i^1 - \frac{\eta}{\sigma_i^1} g_i^1$$

- Compute gradient $\mathbf{g}^2 = \nabla L(\boldsymbol{\theta}^2)$

For each dimension i : $\sigma_i^2 = \sqrt{[(g_i^0)^2 + (g_i^1)^2 + (g_i^2)^2]}$

$$\theta_i^3 \leftarrow \theta_i^2 - \frac{\eta}{\sigma_i^2} g_i^2$$

⋮

- Compute gradient $\mathbf{g}^t = \nabla L(\boldsymbol{\theta}^t)$

For each dimension i : $\sigma_i^t = \sqrt{\sum_{i=0}^t (g_i^t)^2}$

$$\theta_i^{t+1} \leftarrow \theta_i^t - \frac{\eta}{\sigma_i^t} g_i^t$$

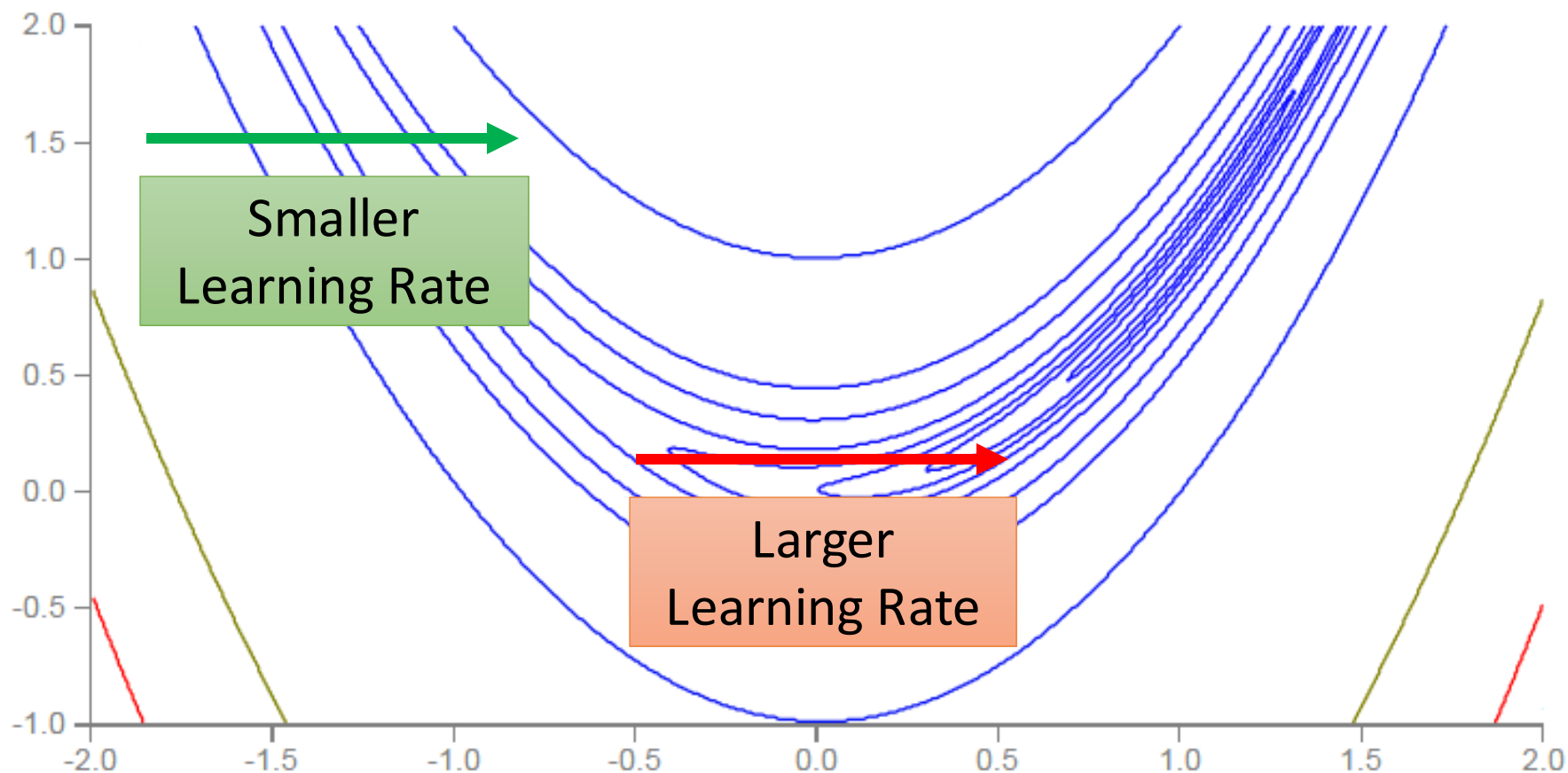
範例程式

連結：

<https://colab.research.google.com/drive/1XPIU-I77dXL9W74jnevmfb8K8XoEPRso?usp=sharing>



同一個參數的Gradient 大小不會一成不變



同一個參數的Gradient 大小不會一成不變

$$g^0 = \begin{bmatrix} 500 \\ 0.4 \end{bmatrix} \quad g^1 = \begin{bmatrix} 432 \\ -0.3 \end{bmatrix} \quad \dots \quad g^{t-1} = \begin{bmatrix} -0.1 \\ 229 \end{bmatrix} \quad g^t = \begin{bmatrix} 0.1 \\ 100 \end{bmatrix}$$

變大 變大



Adagrad: 全部平方加起來

RMSProp: 最近算出來的 gradient 給比較大的影響

RMSProp

- Compute gradient $\mathbf{g}^0 = \nabla L(\boldsymbol{\theta}^0)$

For each dimension i : $\sigma_i^0 = \sqrt{(g_i^0)^2}$

$$\theta_i^1 \leftarrow \theta_i^0 - \frac{\eta}{\sigma_i^0} g_i^0$$

- Compute gradient $\mathbf{g}^1 = \nabla L(\boldsymbol{\theta}^1)$

$$0 < \alpha < 1$$

For each dimension i : $\sigma_i^1 = \sqrt{\alpha(\sigma_i^0)^2 + (1 - \alpha)(g_i^1)^2}$

$$\theta_i^2 \leftarrow \theta_i^1 - \frac{\eta}{\sigma_i^1} g_i^1$$

- Compute gradient $\mathbf{g}^2 = \nabla L(\boldsymbol{\theta}^2)$

For each dimension i : $\sigma_i^2 = \sqrt{\alpha(\sigma_i^1)^2 + (1 - \alpha)(g_i^2)^2}$

$$\theta_i^3 \leftarrow \theta_i^2 - \frac{\eta}{\sigma_i^2} g_i^2$$

⋮

- Compute gradient $\mathbf{g}^t = \nabla L(\boldsymbol{\theta}^t)$

For each dimension i : $\sigma_i^t = \sqrt{\alpha(\sigma_i^{t-1})^2 + (1 - \alpha)(g_i^t)^2}$

$$\theta_i^{t+1} \leftarrow \theta_i^t - \frac{\eta}{\sigma_i^t} g_i^t$$

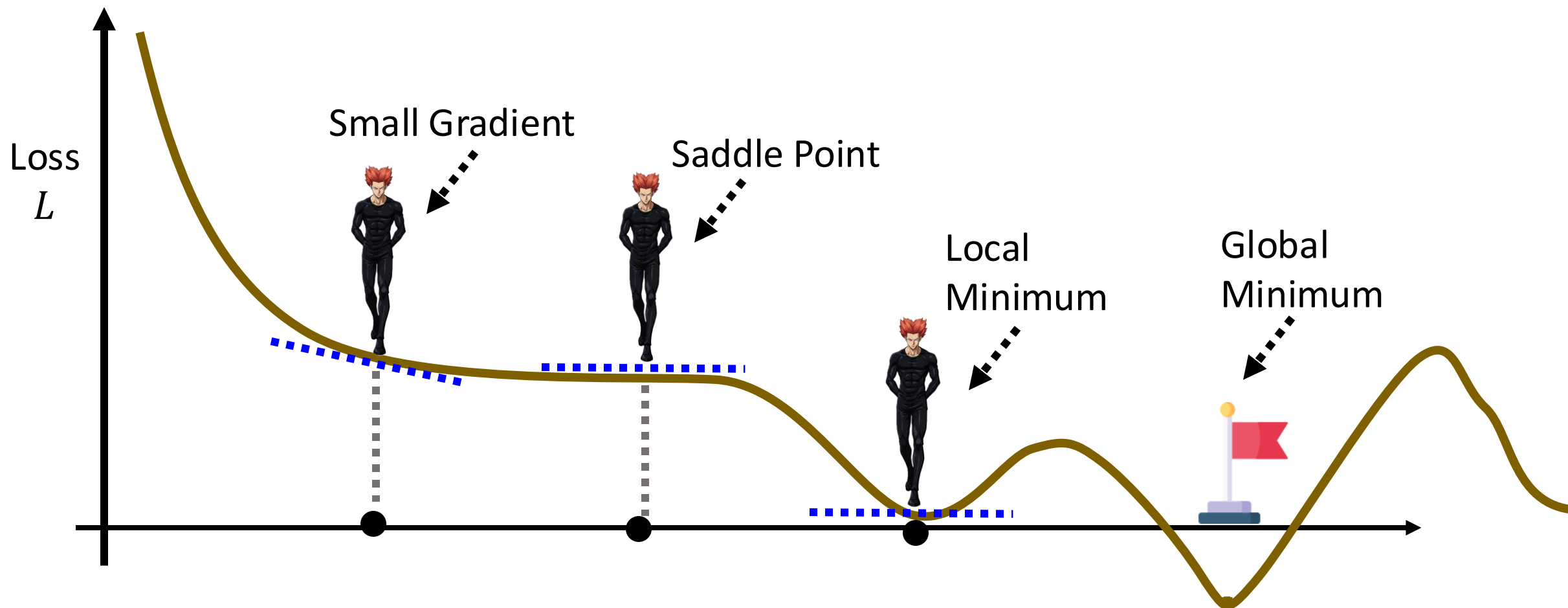
範例程式

連結：

<https://colab.research.google.com/drive/1XPIU-I77dXL9W74jnevmfb8K8XoEPRso?usp=sharing>

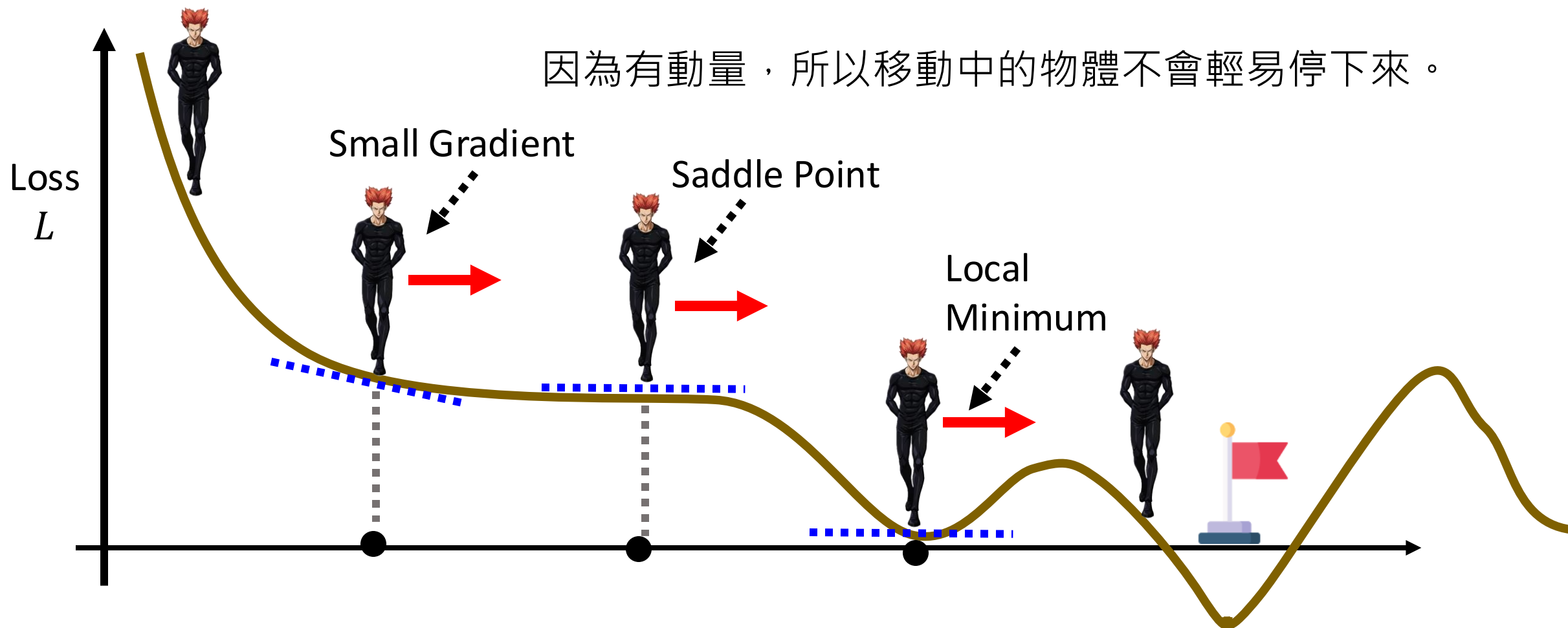


Optimization 會在 Gradient 很小時停止



考慮動量 (Momentum)

因為有動量，所以移動中的物體不會輕易停下來。



Momentum

- Compute gradient $\mathbf{g}^0 = \nabla L(\boldsymbol{\theta}^0)$

For each dimension i : $m_i^0 = g_i^0$

$$\theta_i^1 \leftarrow \theta_i^0 - \eta m_i^0$$

- Compute gradient $\mathbf{g}^1 = \nabla L(\boldsymbol{\theta}^1)$

For each dimension i : $m_i^1 = g_i^0 + g_i^1$

$$\theta_i^2 \leftarrow \theta_i^1 - \eta m_i^1$$

- Compute gradient $\mathbf{g}^2 = \nabla L(\boldsymbol{\theta}^2)$

For each dimension i : $m_i^2 = g_i^0 + g_i^1 + g_i^2$

$$\theta_i^3 \leftarrow \theta_i^2 - \eta m_i^2$$

⋮

- Compute gradient $\mathbf{g}^t = \nabla L(\boldsymbol{\theta}^t)$

For each dimension i : $m_i^t = g_i^0 + g_i^1 + g_i^2 + \cdots + g_i^t$

$$\theta_i^{t+1} \leftarrow \theta_i^t - \eta m_i^t$$

Momentum (這不是經典的 Momentum)

➤ Compute gradient $\mathbf{g}^0 = \nabla L(\boldsymbol{\theta}^0)$

For each dimension i : $m_i^0 = g_i^0$

$$\theta_i^1 \leftarrow \theta_i^0 - \eta m_i^0$$

➤ Compute gradient $\mathbf{g}^1 = \nabla L(\boldsymbol{\theta}^1)$

$$0 < \beta < 1$$

For each dimension i : $m_i^1 = \beta m_i^0 + (1 - \beta) g_i^1$

$$\theta_i^2 \leftarrow \theta_i^1 - \eta m_i^1$$

➤ Compute gradient $\mathbf{g}^2 = \nabla L(\boldsymbol{\theta}^2)$

For each dimension i : $m_i^2 = \beta m_i^1 + (1 - \beta) g_i^2$

$$\theta_i^3 \leftarrow \theta_i^2 - \eta m_i^2$$

⋮

➤ Compute gradient $\mathbf{g}^t = \nabla L(\boldsymbol{\theta}^t)$

For each dimension i : $m_i^t = \beta m_i^{t-1} + (1 - \beta) g_i^t$

$$\theta_i^{t+1} \leftarrow \theta_i^t - \eta m_i^t$$

範例程式

連結：

<https://colab.research.google.com/drive/1XPIU-I77dXL9W74jnevmfb8K8XoEPRso?usp=sharing>



Adam: RMSProp + Momentum

- Compute gradient $\mathbf{g}^t = \nabla L(\boldsymbol{\theta}^t)$

Momentum

For each dimension i : $\mathbf{m}_i^t = \beta \mathbf{m}_i^{t-1} + (1 - \beta) \mathbf{g}_i^t$ $\theta_i^{t+1} \leftarrow \theta_i^t - \eta \mathbf{m}_i^t$

RMSProp

For each dimension i : $\sigma_i^t = \sqrt{\alpha (\sigma_i^{t-1})^2 + (1 - \alpha) (\mathbf{g}_i^t)^2}$ $\theta_i^{t+1} \leftarrow \theta_i^t - \frac{\eta}{\sigma_i^t} \mathbf{g}_i^t$

Adam

$$\theta_i^{t+1} \leftarrow \theta_i^t - \frac{\eta}{\sigma_i^t} \mathbf{m}_i^t$$

Adam has bias-corrected terms,
which are omitted here for simplicity.

範例程式

連結：

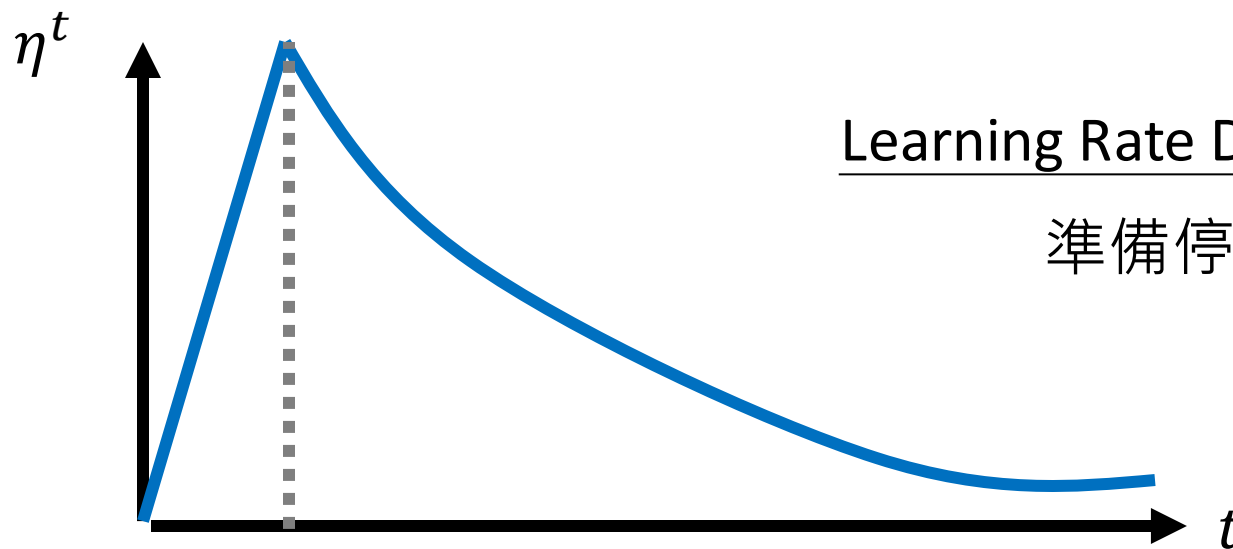
<https://colab.research.google.com/drive/1XPIU-I77dXL9W74jnevmfb8K8XoEPRso?usp=sharing>



Learning Rate Scheduling

$$\theta_i^{t+1} \leftarrow \theta_i^t - \frac{\eta}{\sigma_i^t} m_i^t \quad \longrightarrow \quad \theta_i^{t+1} \leftarrow \theta_i^t - \frac{\eta^t}{\sigma_i^t} m_i^t$$

Warm Up
探索地形



Learning Rate Decay
準備停下來

Summarization of Optimizer

$$\theta_i^{t+1} \leftarrow \theta_i^t - \frac{\eta^t}{\sigma_i^t} m_i^t$$

Learning rate scheduling

Momentum: sum of the previous gradients

Consider direction

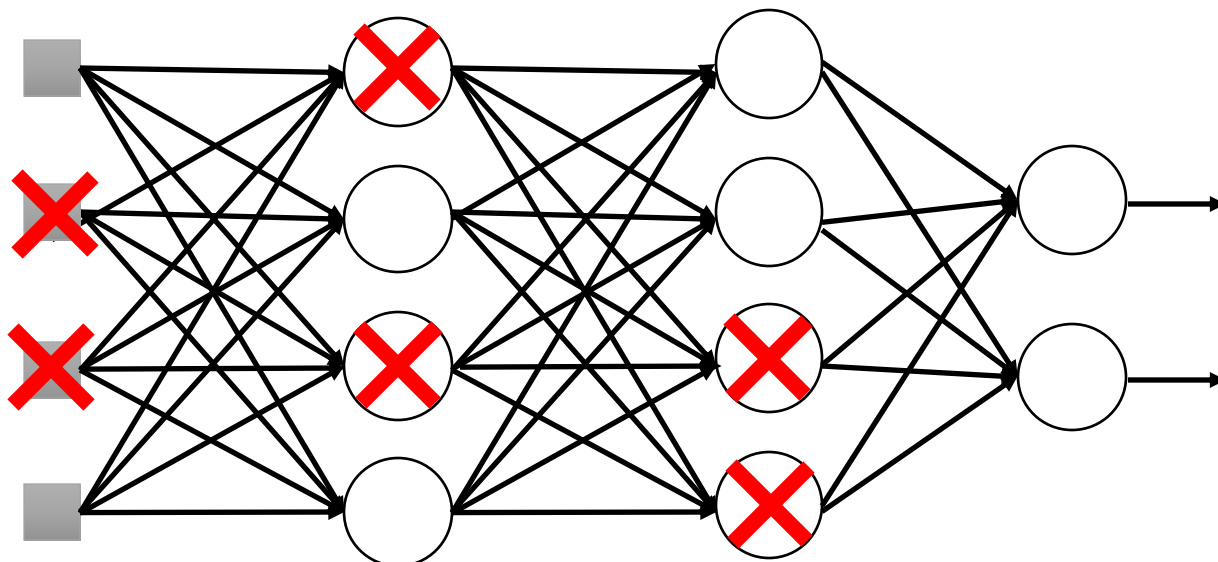
square of the gradients

only magnitude

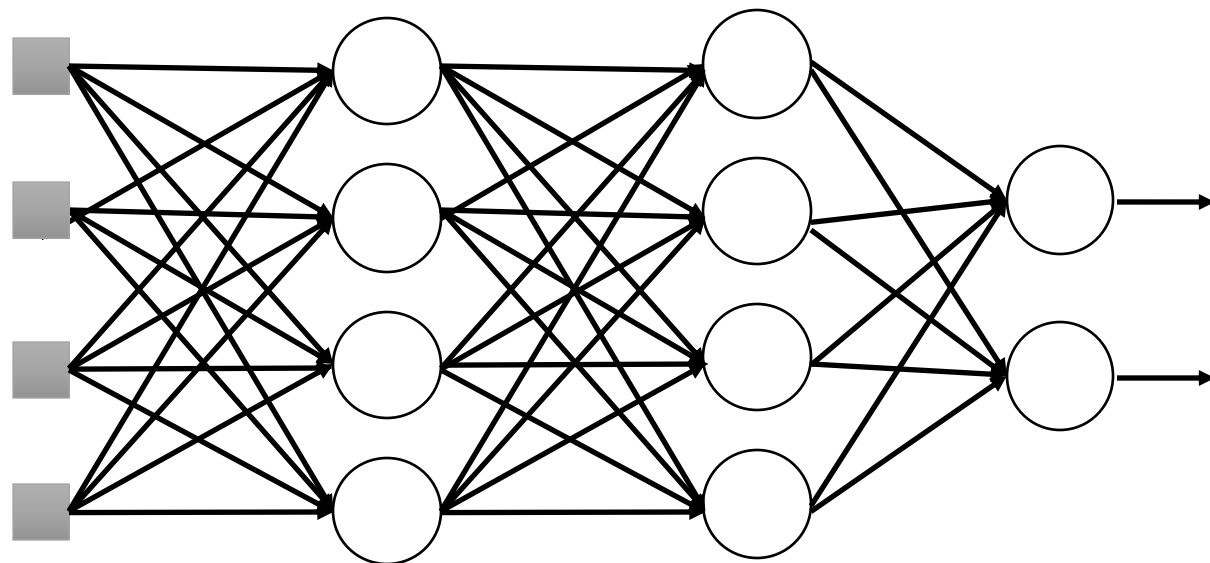
方法名	改了那一個步驟	帶來什麼好處
Adagrad, RMSProp, Momentum, Adam, etc.	找最好的函式	Better Optimization (not for Generalization)

Dropout

訓練時隨機丟掉
一些神經元



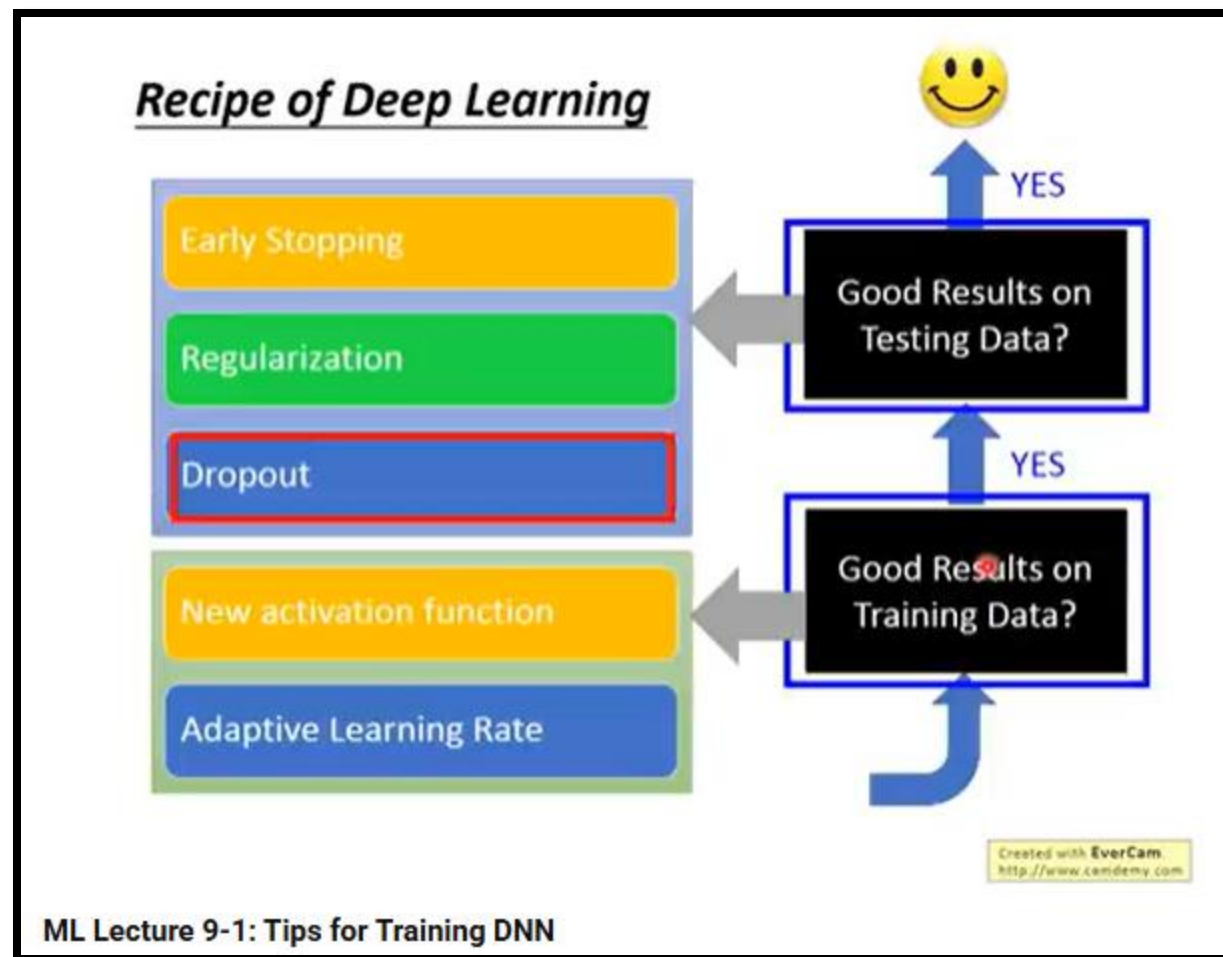
驗證和測試時
火力全開



方法名	改了那一個步驟	帶來什麼好處
Dropout	找最好的函式	Better Generalization (Worse Optimization)

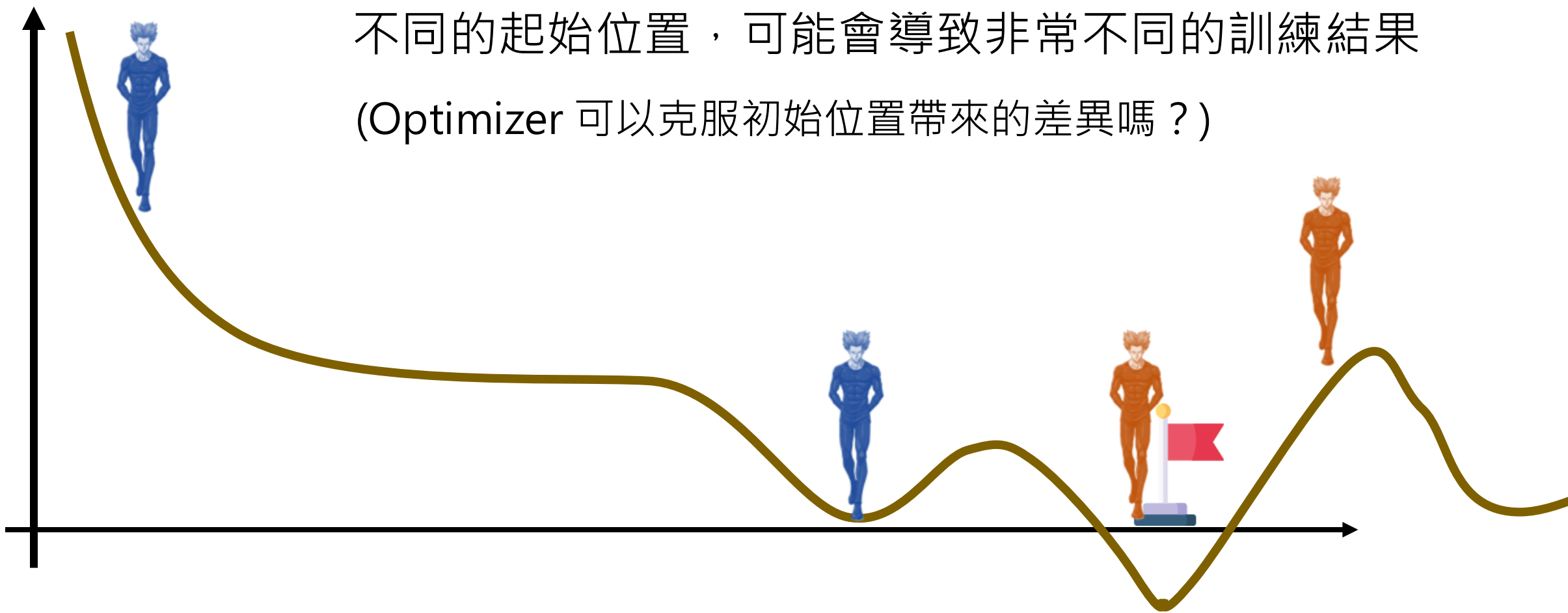
注意 Dropout
的選擇時機!

https://youtu.be/xki61j7z-30?si=EoF1VjcP3_L5UWI3&t=4227



Initialization

不同的起始位置，可能會導致非常不同的訓練結果
(Optimizer 可以克服初始位置帶來的差異嗎？)



範例程式

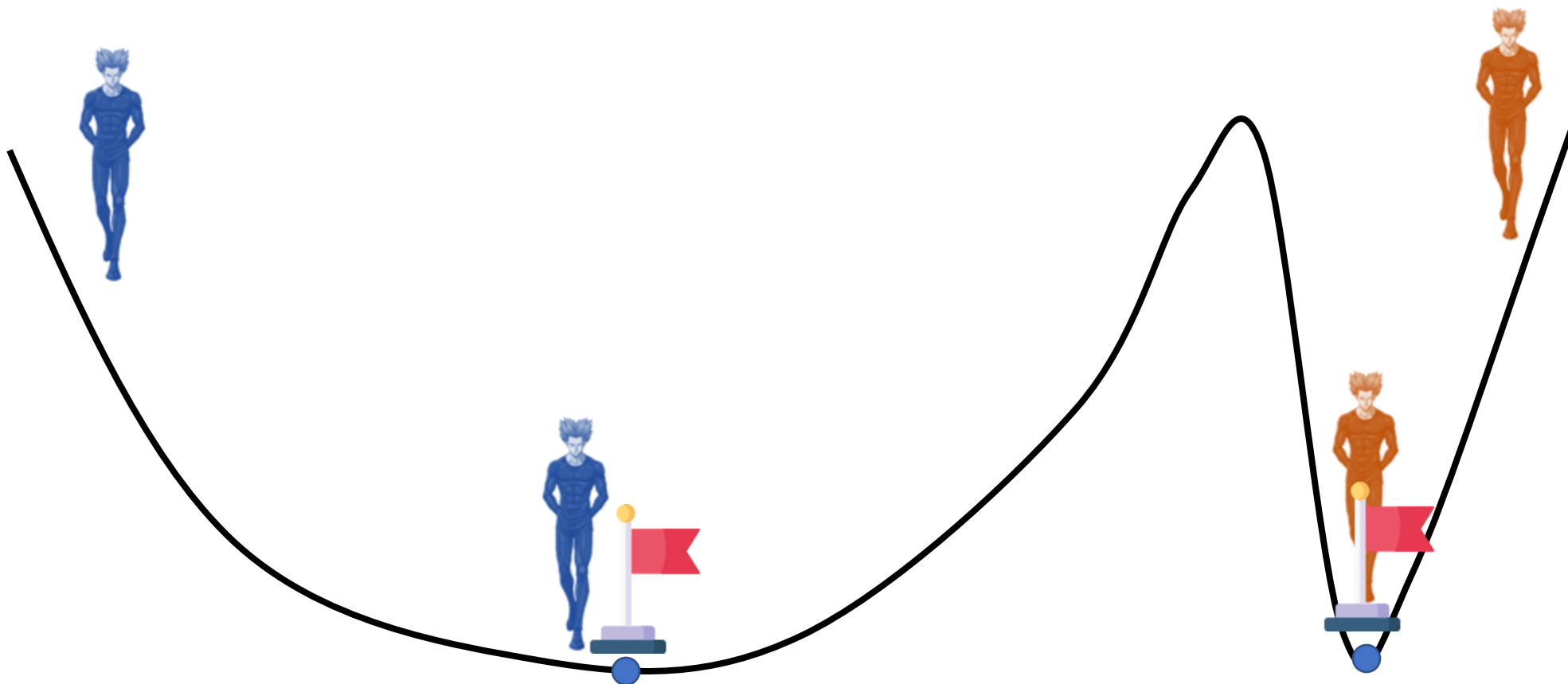
連結：

<https://colab.research.google.com/drive/1XPIU-I77dXL9W74jnevmfb8K8XoEPRso?usp=sharing>



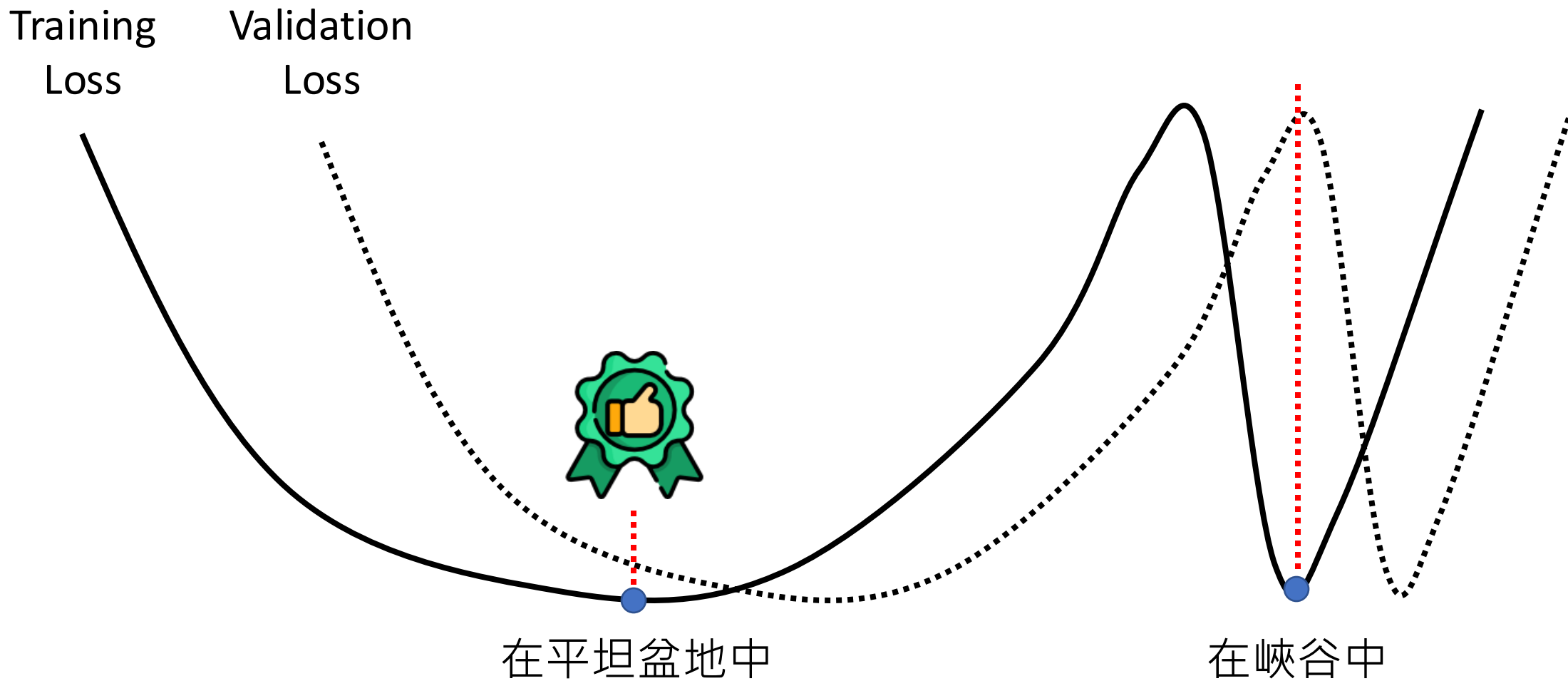
Initialization

當有多個 Training Loss 最低點的時候，
起始位置決定我們走到哪一個最低點



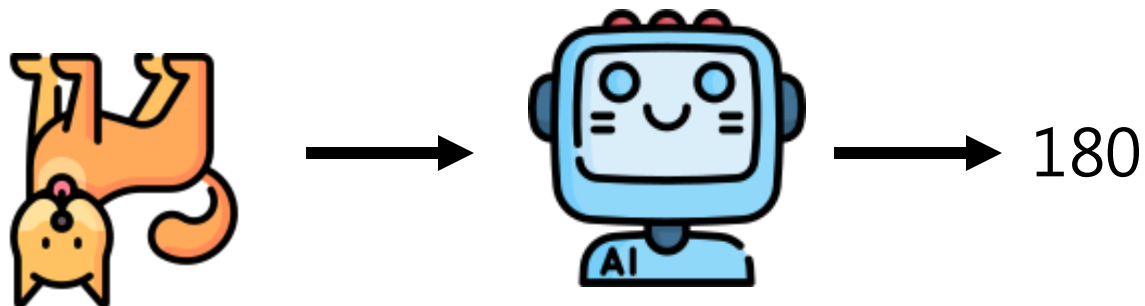
Initialization

當有多個 Training Loss 最低點的時候，
起始位置決定我們走到哪一個最低點



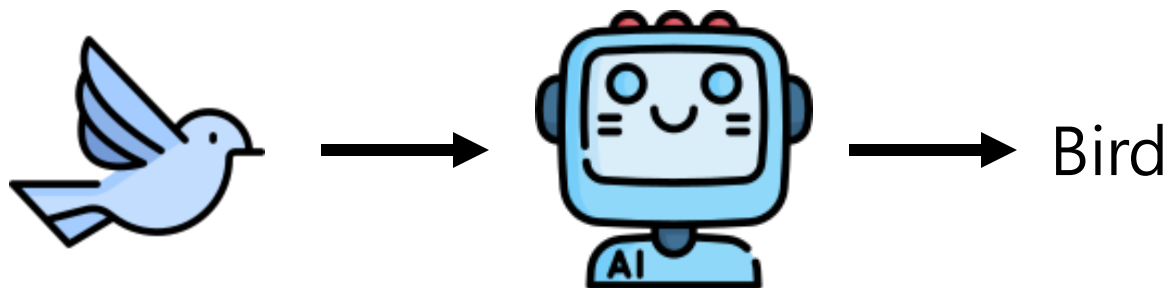
Initialization

Pretext Task ← ... 可以輕易蒐集大量
資料進行訓練



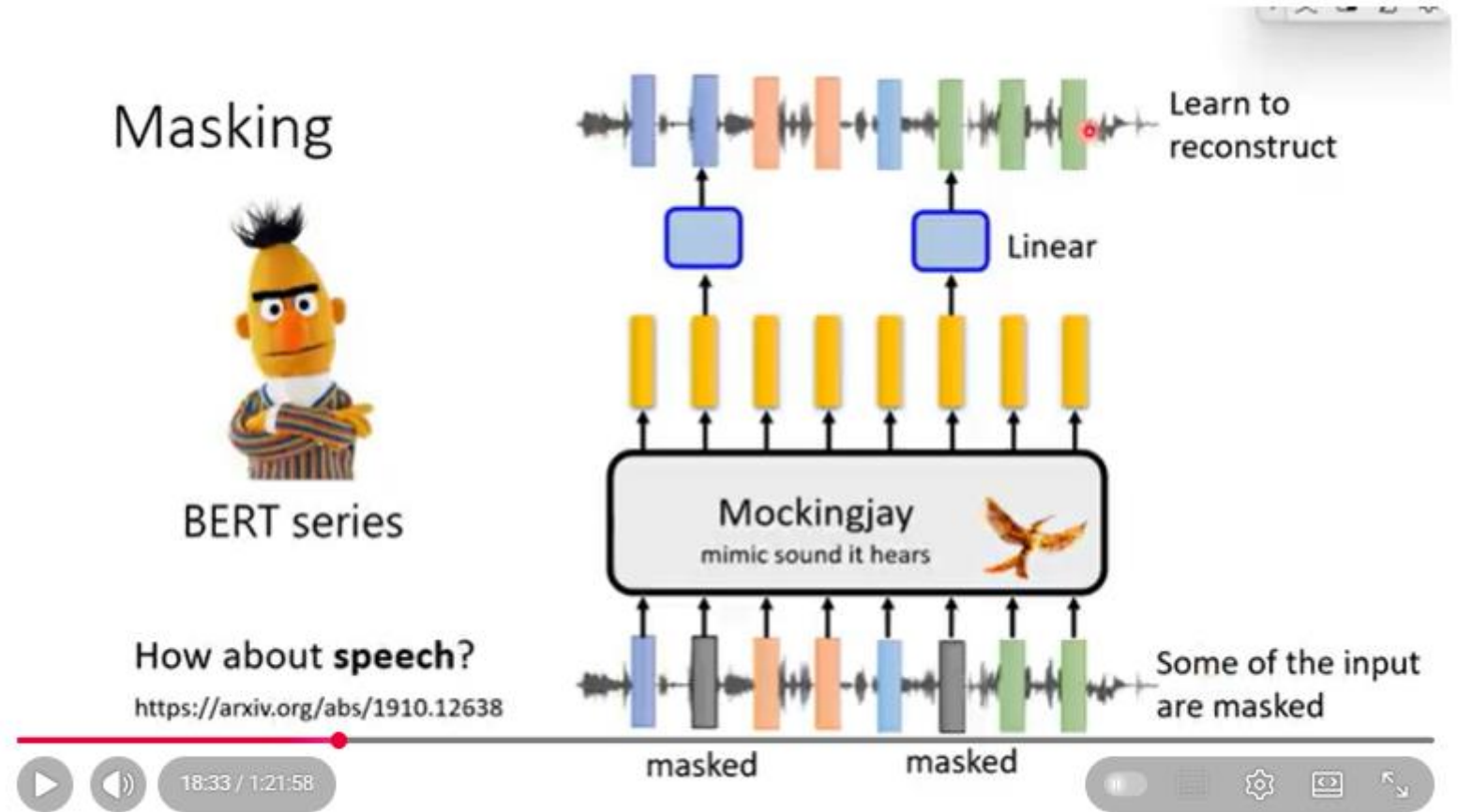
Pre-train

Initialization



Downstream Task

More about Pretext Task



【機器學習 2022】 語音與影像上的神奇自督導式學習 (Self-supervised Learning) 模型

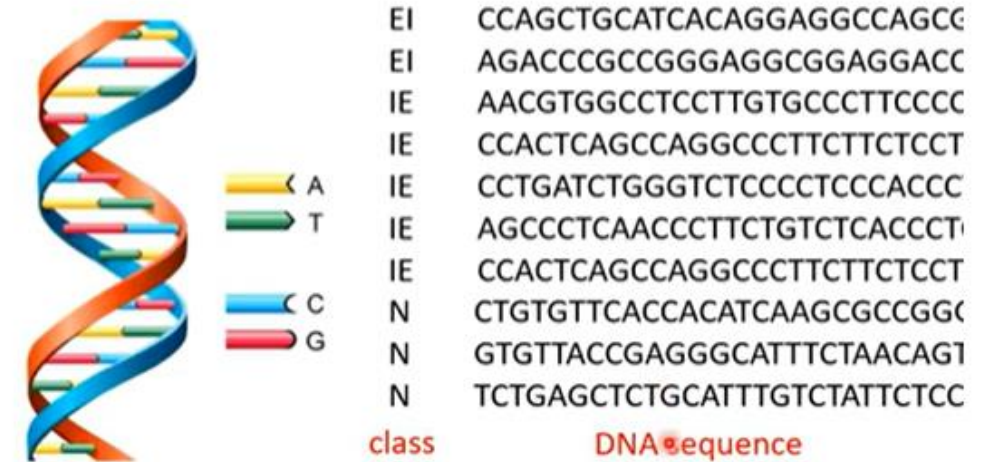
https://youtu.be/IMIN1iKYNmA?si=M6hL7pKbOU_ufefo

Initialization

Pre-train 對於 Optimization
和 Generalization 都有幫助

Why does BERT work?

- Applying BERT to **protein, DNA, music classification**

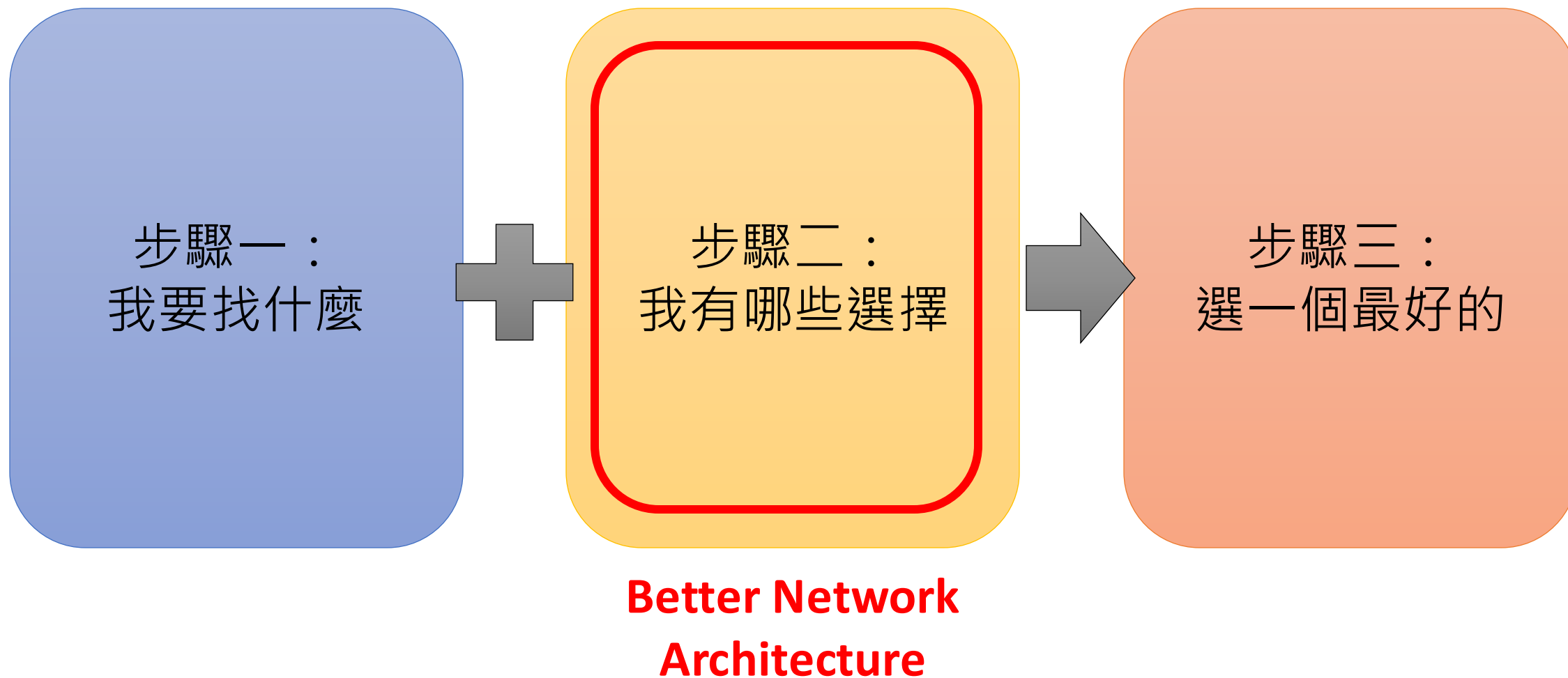


【機器學習2021】自督導式學習 (Self-supervised Learning) (三) – BERT的奇聞軼事

<https://youtu.be/ExXA05i8DEQ?si=bpml5yogGDjI3FUf&t=508>

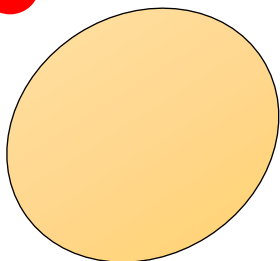
方法名	改了那一個步驟	帶來什麼好處
Initialization (e.g., pre-train)	找最好的函式	Better Optimization, Better Generalization

找函式步驟 3 + 1



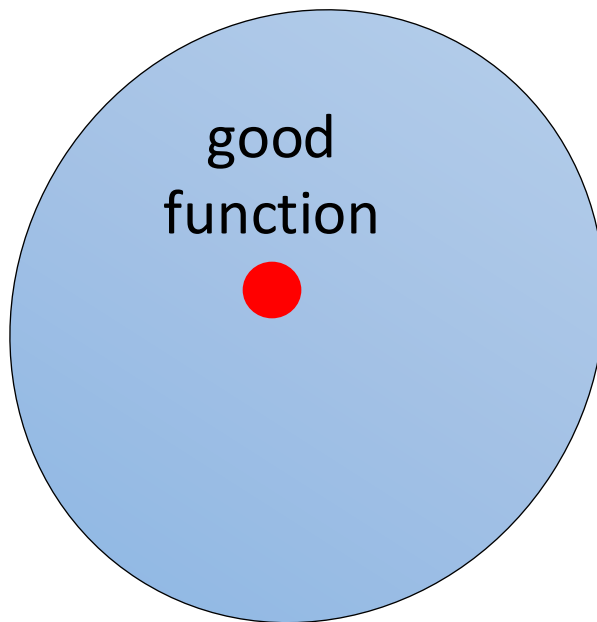
劃定一個剛剛好的函式範圍

good
function



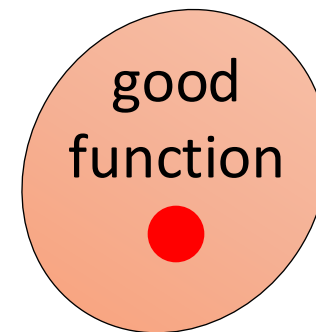
函式集合太小，可能
會沒有包含好的函式

good
function



函式集合大一點，比較
可能包到好的函式
(容易 overfitting)

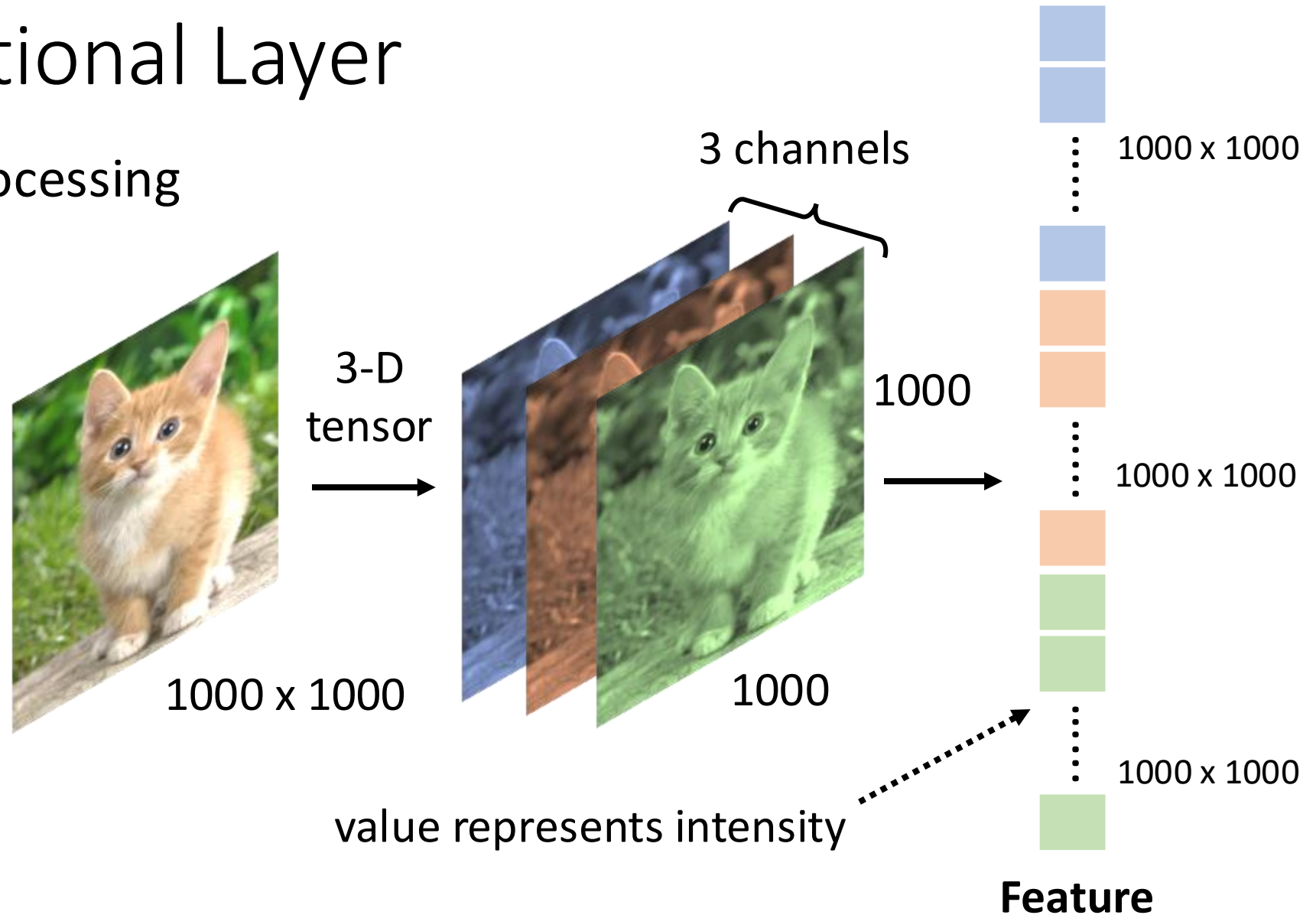
good
function

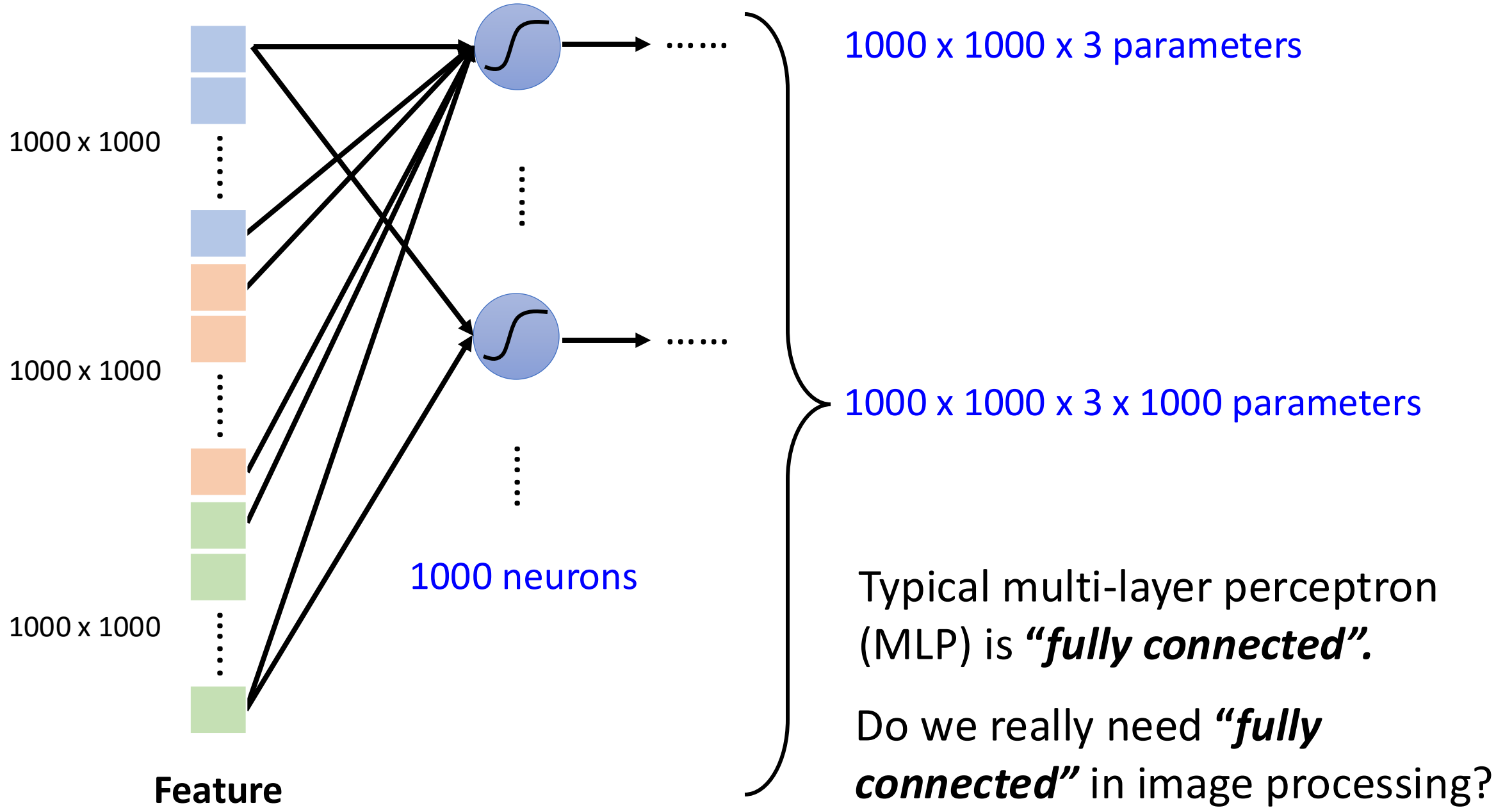


選一個小範圍但剛好
包含好的函式
(根據 domain knowledge)

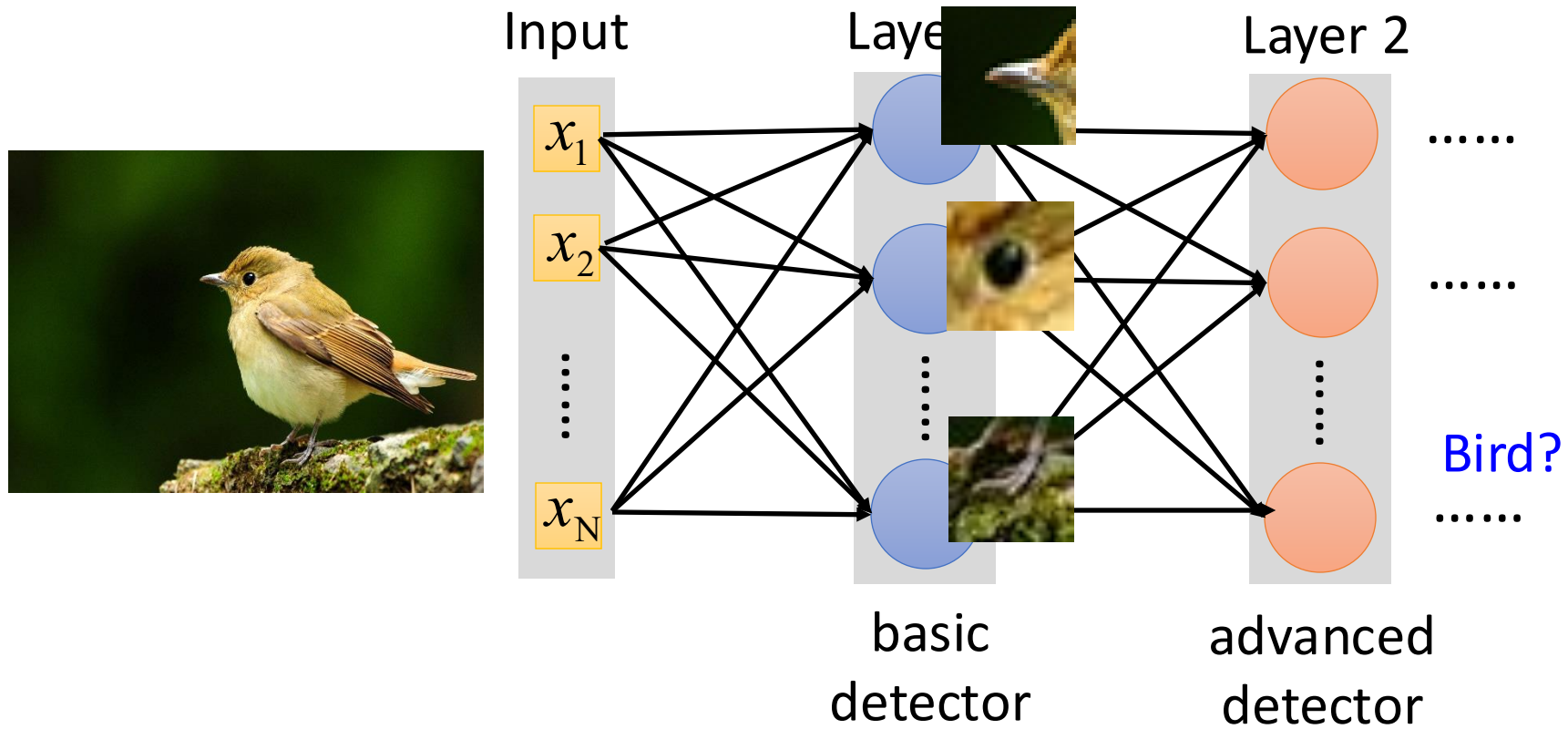
Convolutional Layer

For image processing

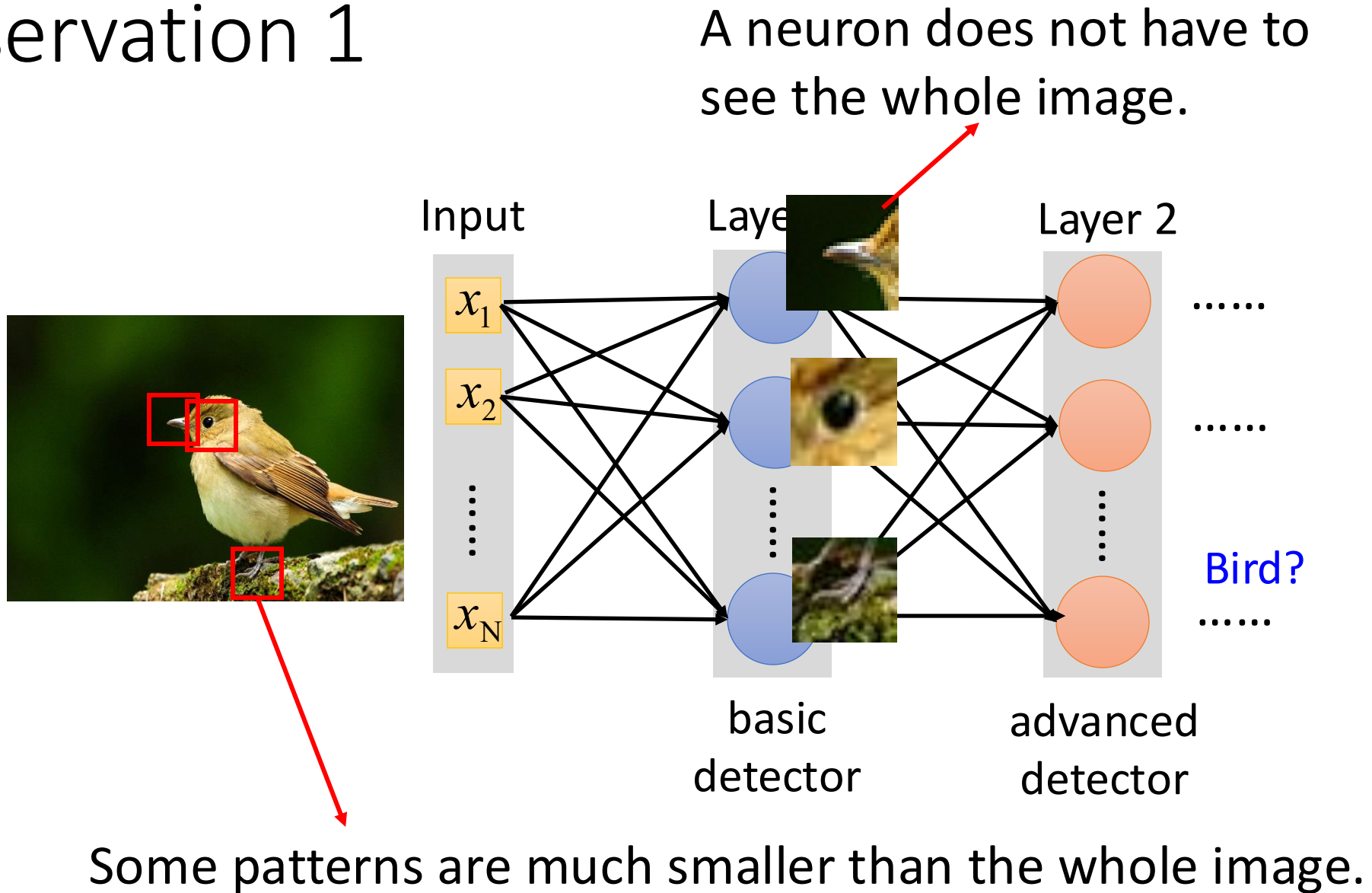




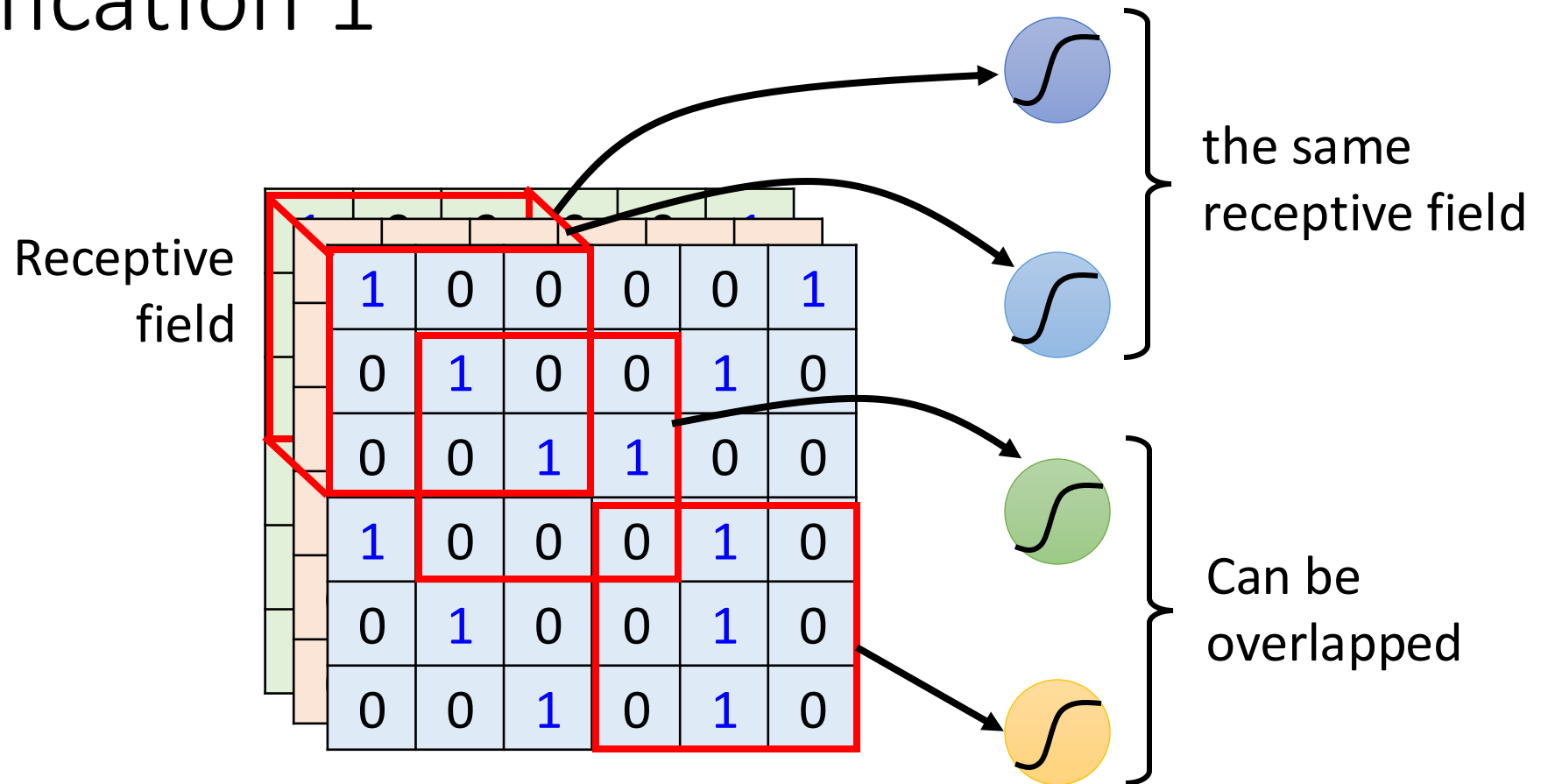
Observation 1



Observation 1

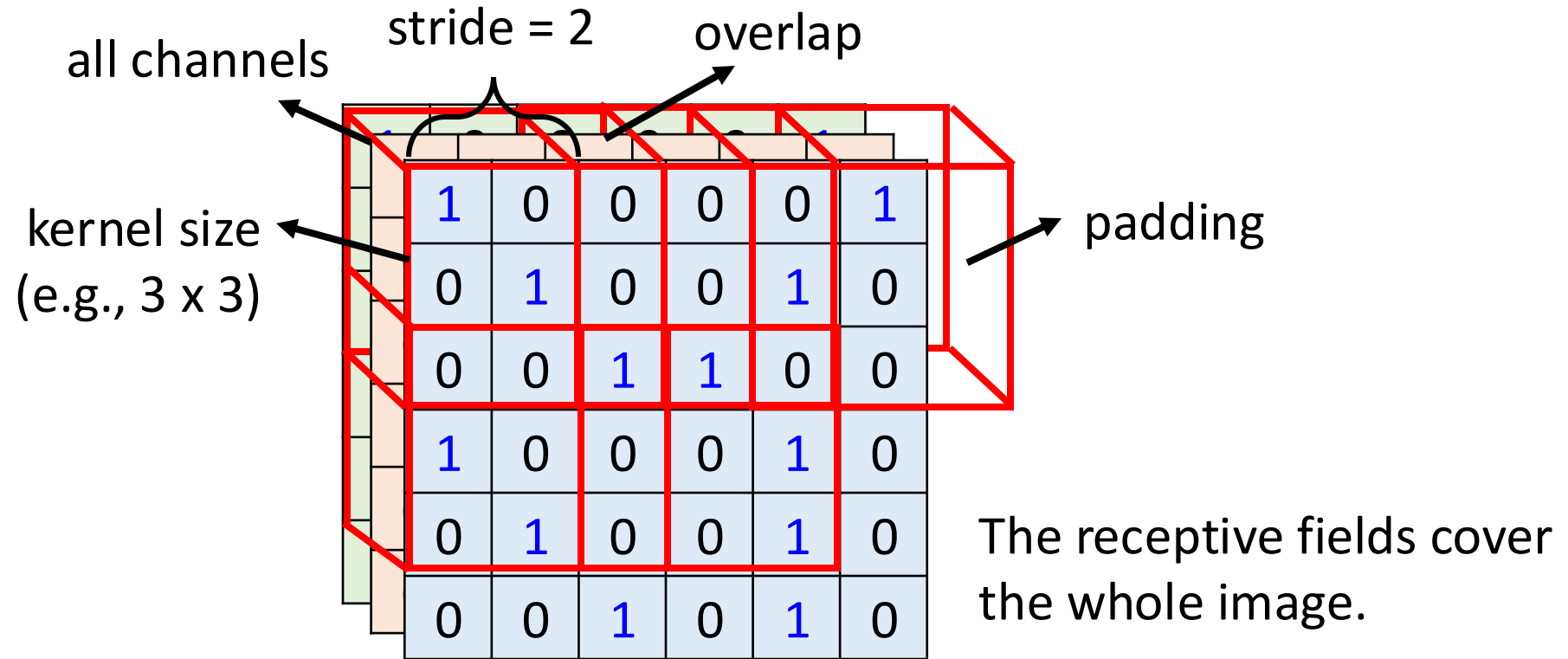


Simplification 1

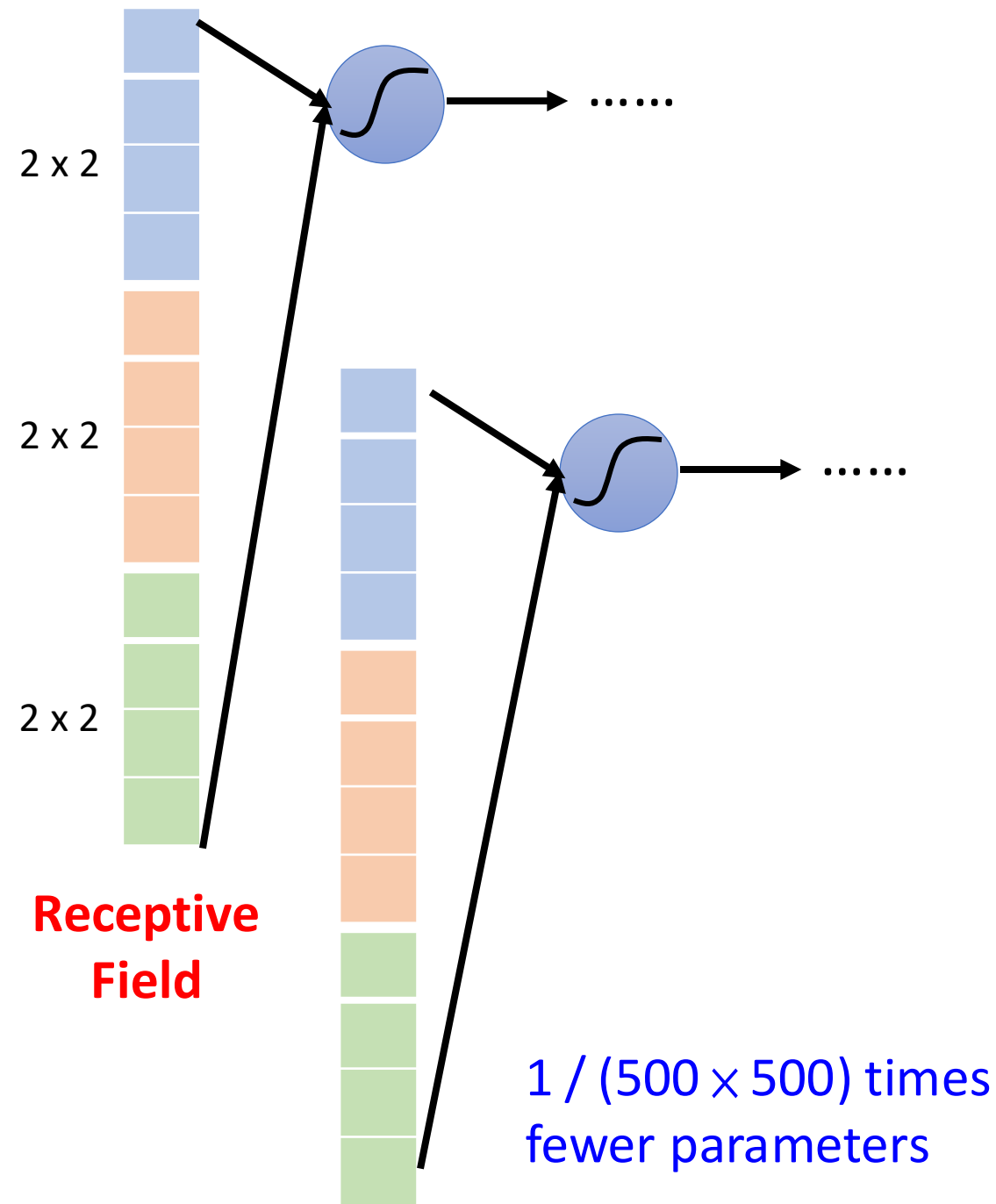
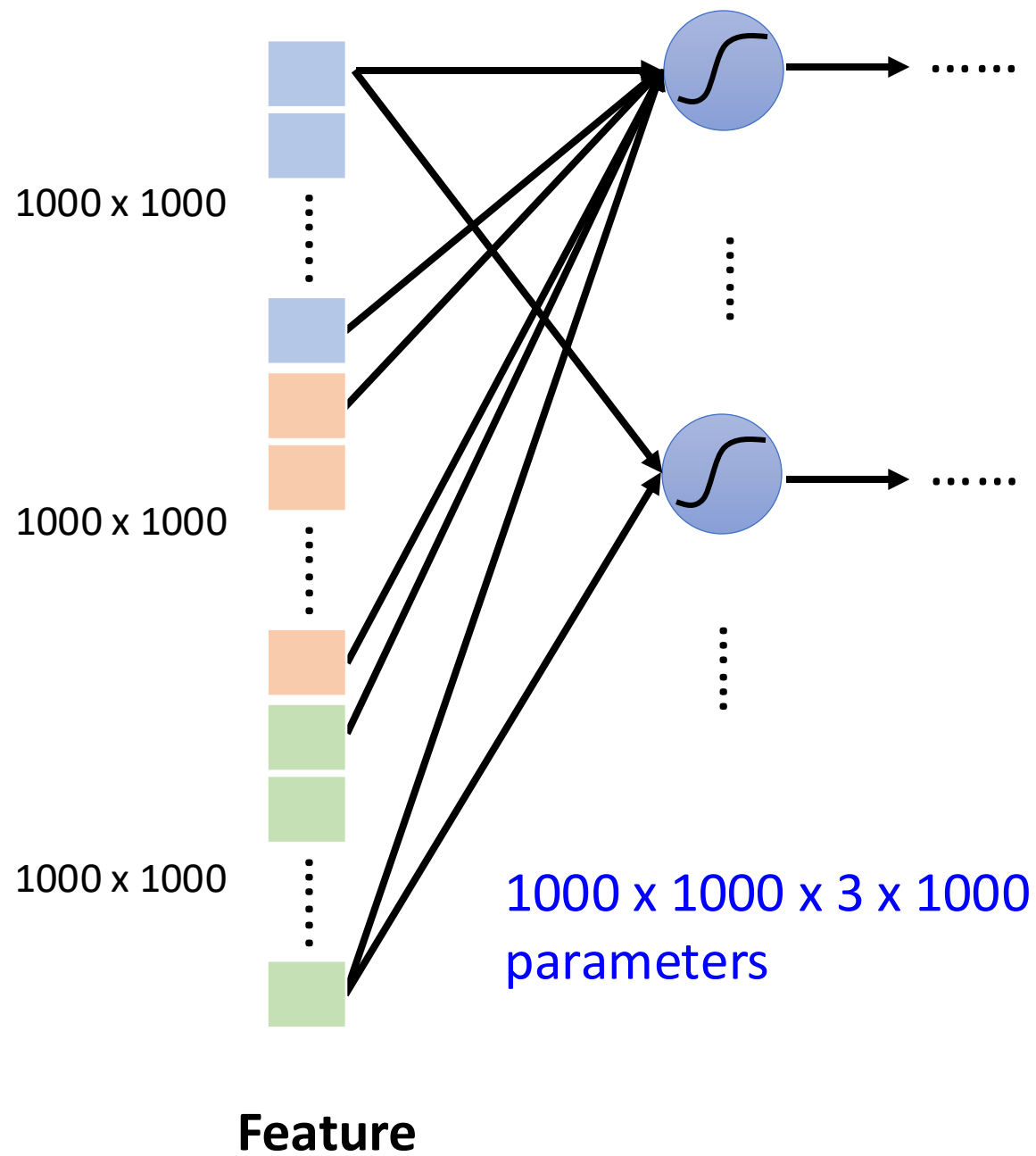


- Can different neurons have different sizes of receptive field?
- Not square receptive field?

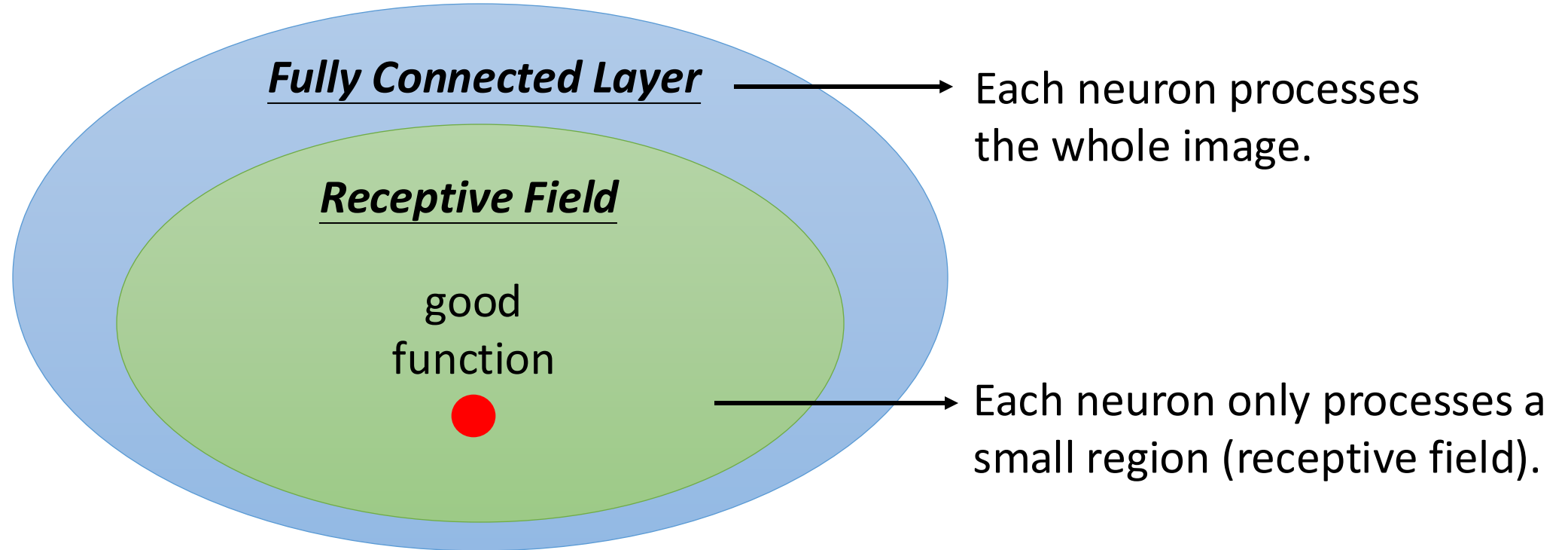
Simplification 1 – Typical Setting



Each receptive field has a set of neurons (e.g., 64 neurons).



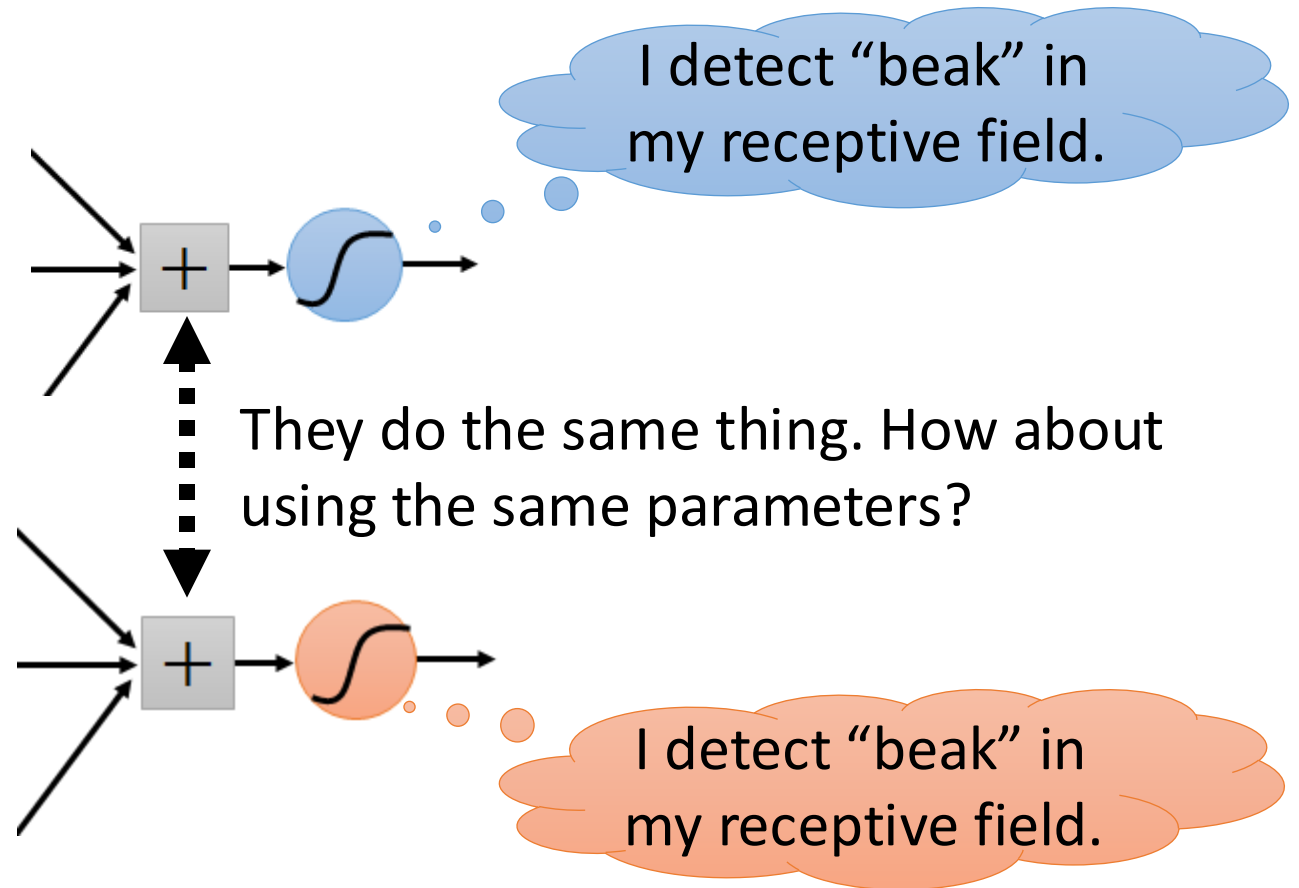
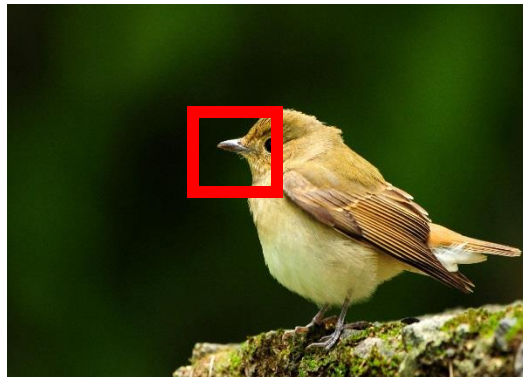
Benefit of Simplification 1



- Basic patterns are much smaller than the whole image. Each neuron only has to process a receptive field.

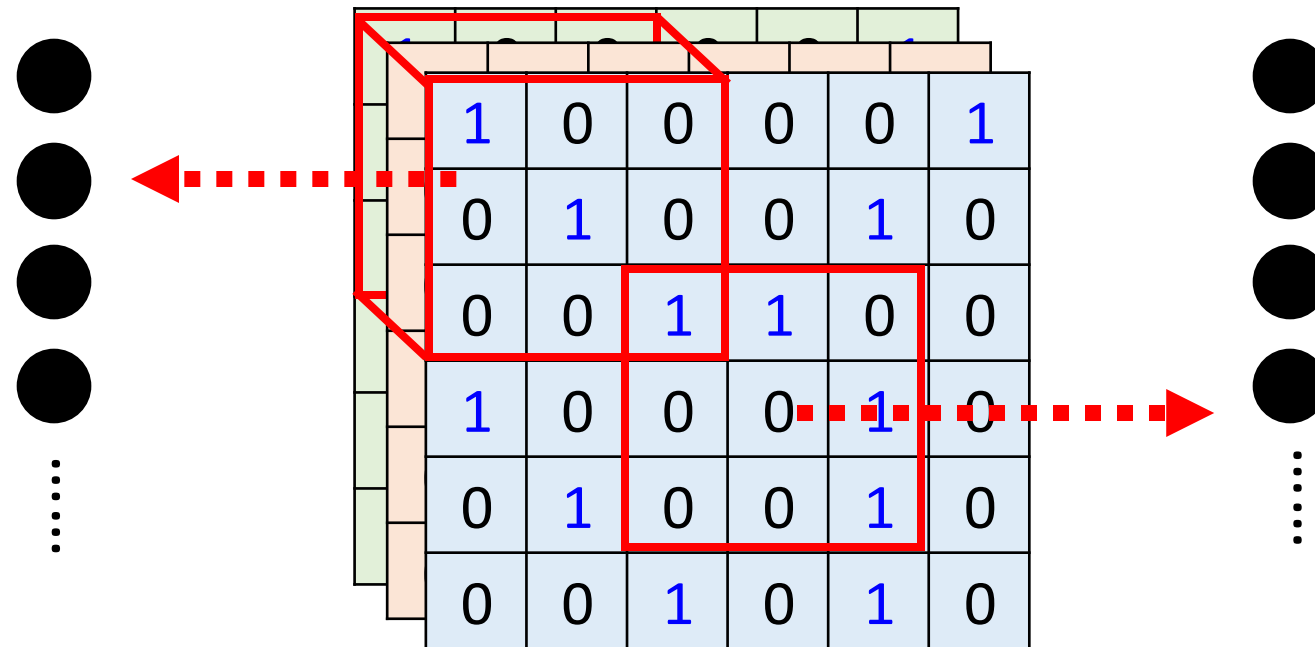
Observation 2

- The same patterns appear in different regions.



Simplification 2 – Sharing parameters

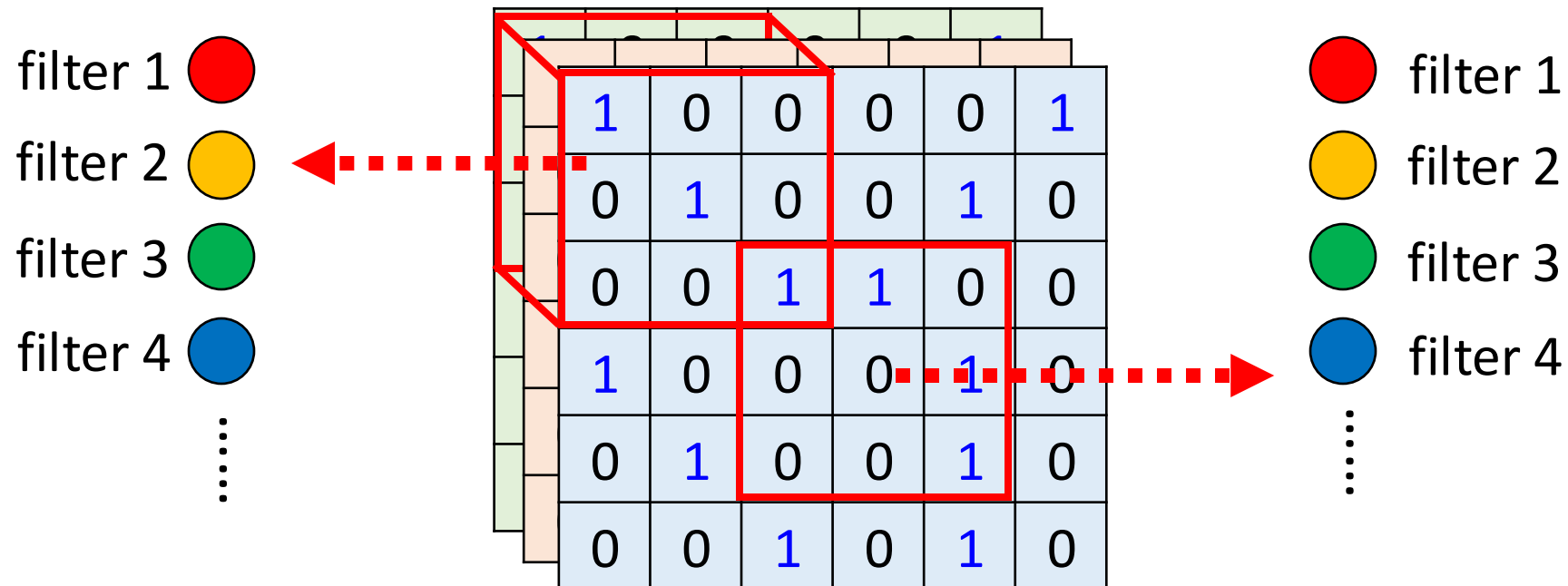
Each receptive field has a set of neurons (e.g., 64 neurons).



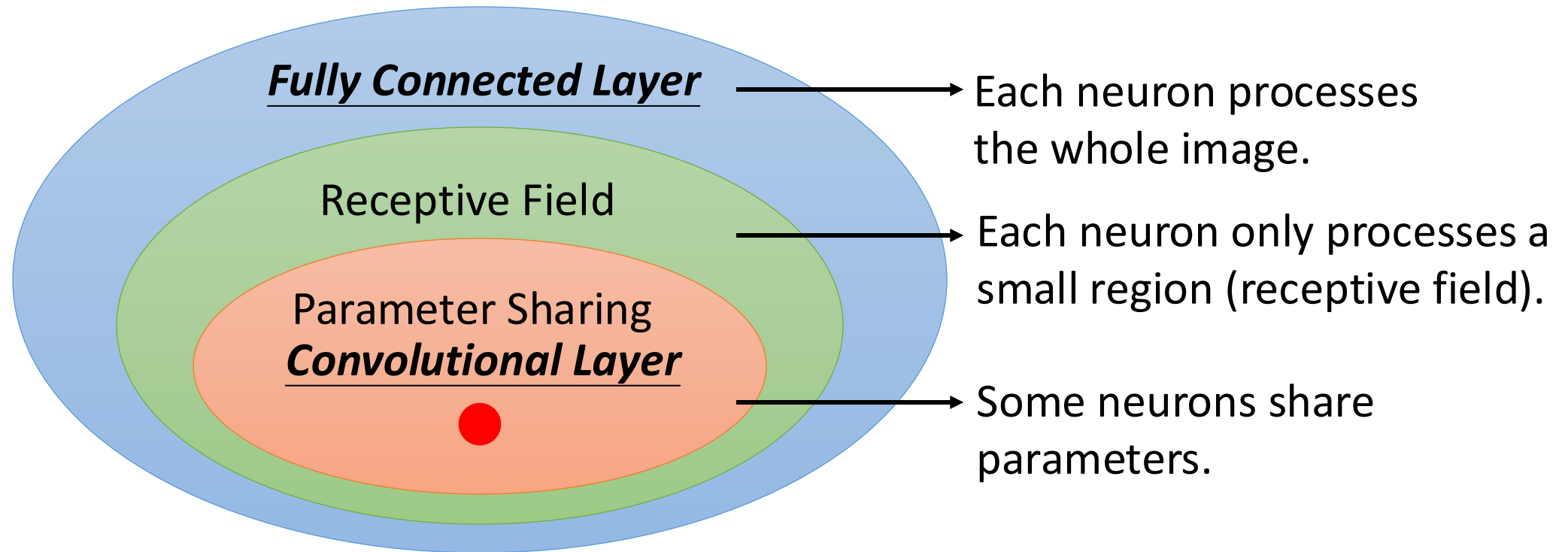
Simplification 2 – Sharing parameters

Each receptive field has a set of neurons (e.g., 64 neurons).

Each receptive field has the neurons with the same set of parameters.



Benefit of Simplification 1 + 2

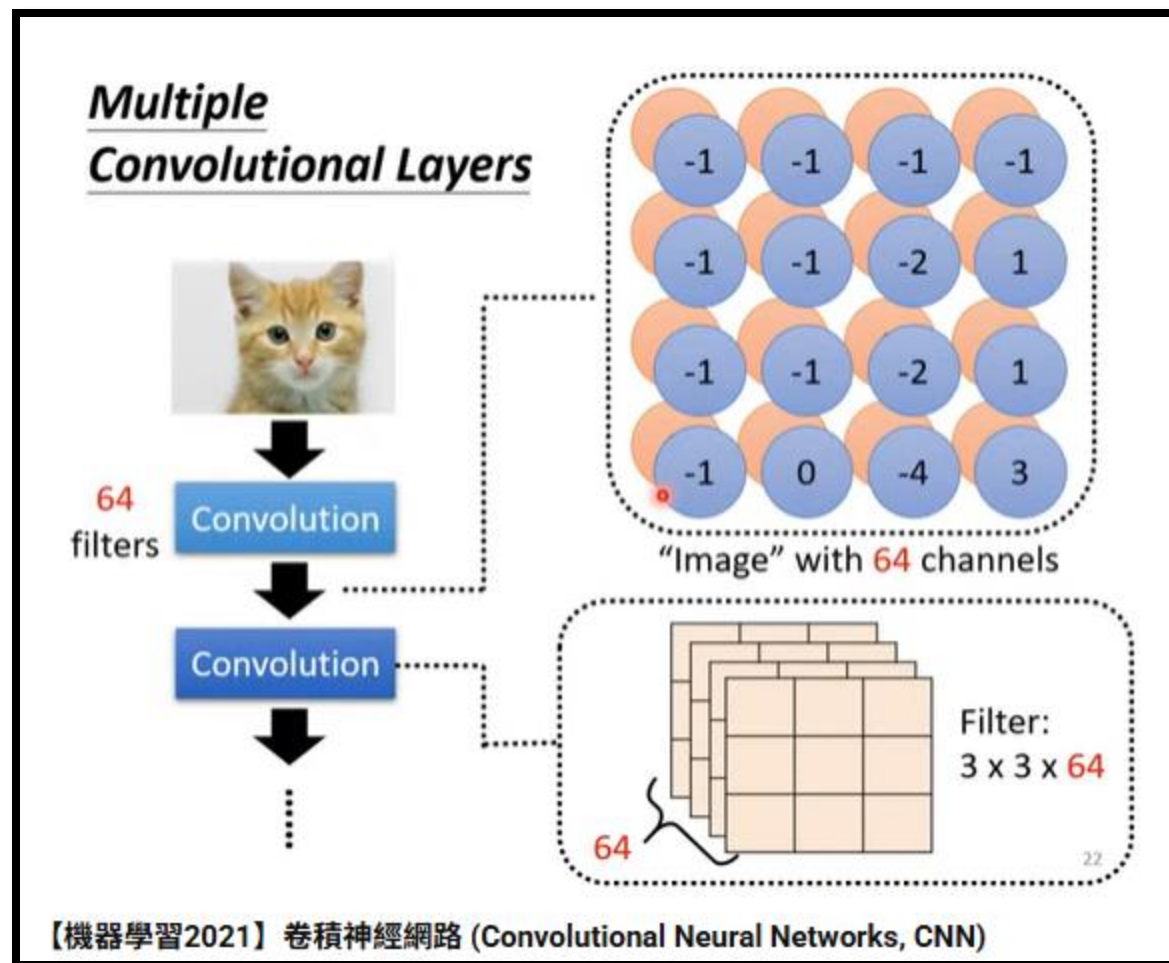


- Basic patterns are much smaller than the whole image. Each neuron only has to process a receptive field.
- Different receptive fields need the same detectors, so neurons for different receptive fields share parameters.

方法名	改了那一個步驟	帶來什麼好處
CNN (e.g., for image)	改變函式搜尋範圍	Better Generalization

To learn more about CNN ...

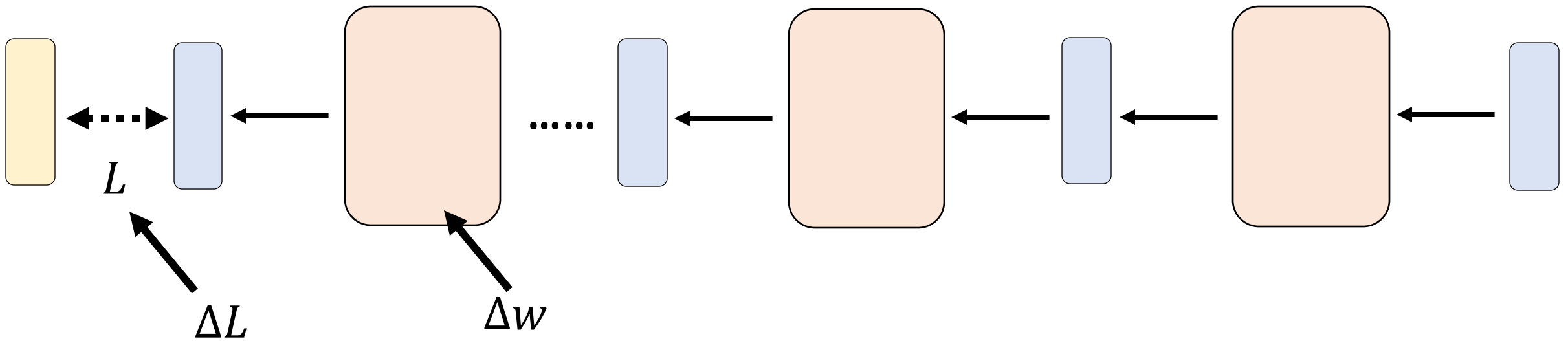
<https://youtu.be/OP5HcXJg2Aw?si=cNhXe9W-PgP45jN1&t=1686>



(為了簡化，此處沒有畫出 activation 和 bias)

Skip Connection (Residual Connection)

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

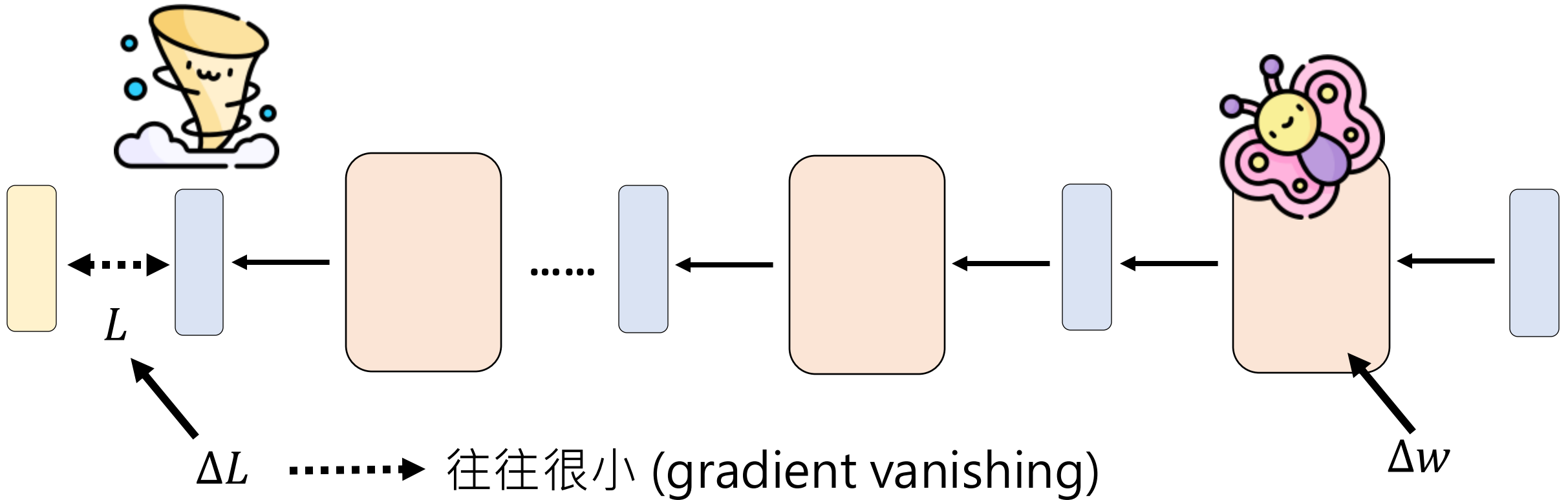


$$\frac{\partial L}{\partial w} \approx \frac{\Delta L}{\Delta w}$$

(為了簡化，此處沒有畫出 activation 和 bias)

Skip Connection (Residual Connection)

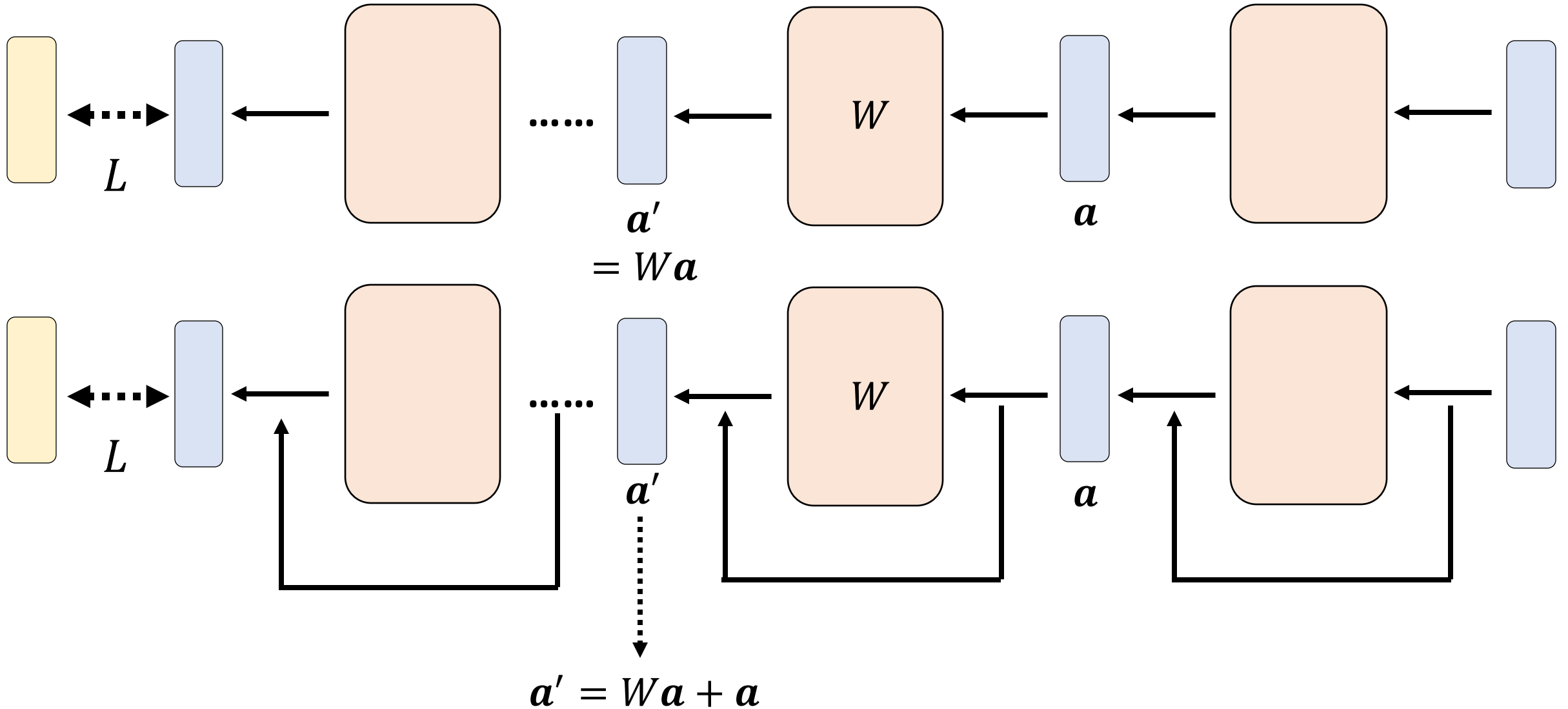
<https://arxiv.org/abs/1512.03385>

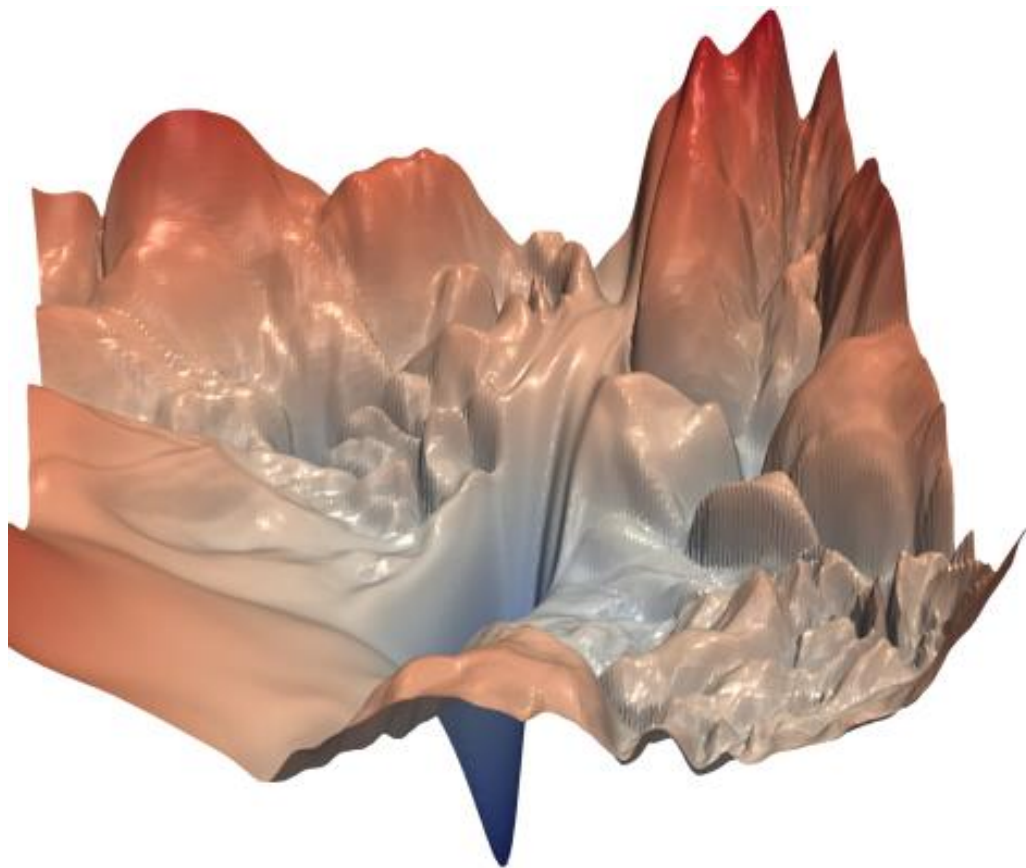


$$\frac{\partial L}{\partial w} \approx \frac{\Delta L}{\Delta w}$$

(為了簡化，此處沒有畫出 activation 和 bias)

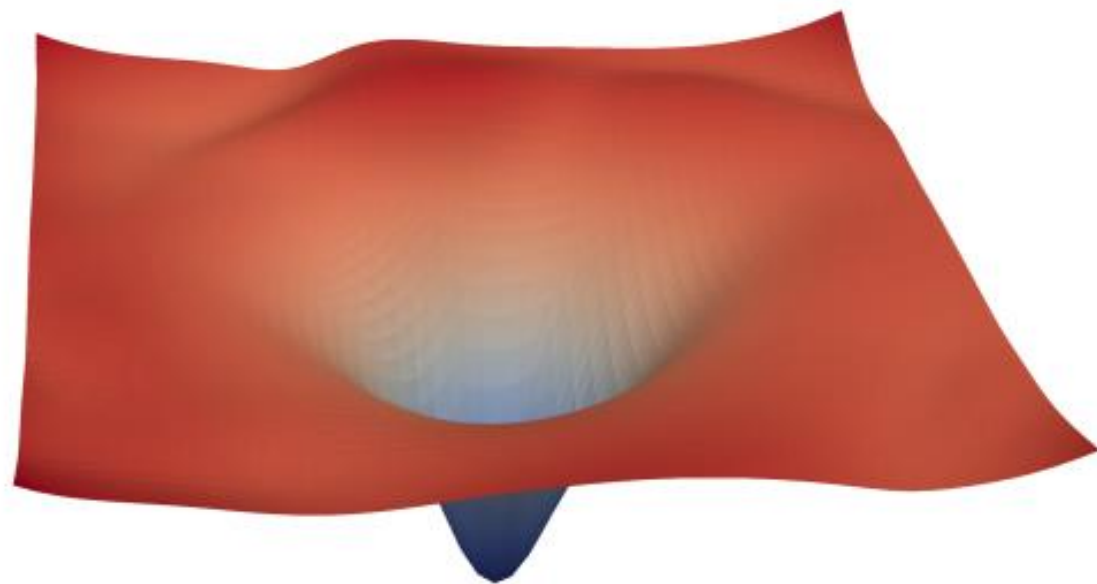
Skip Connection (Residual Connection)





Without skip connection

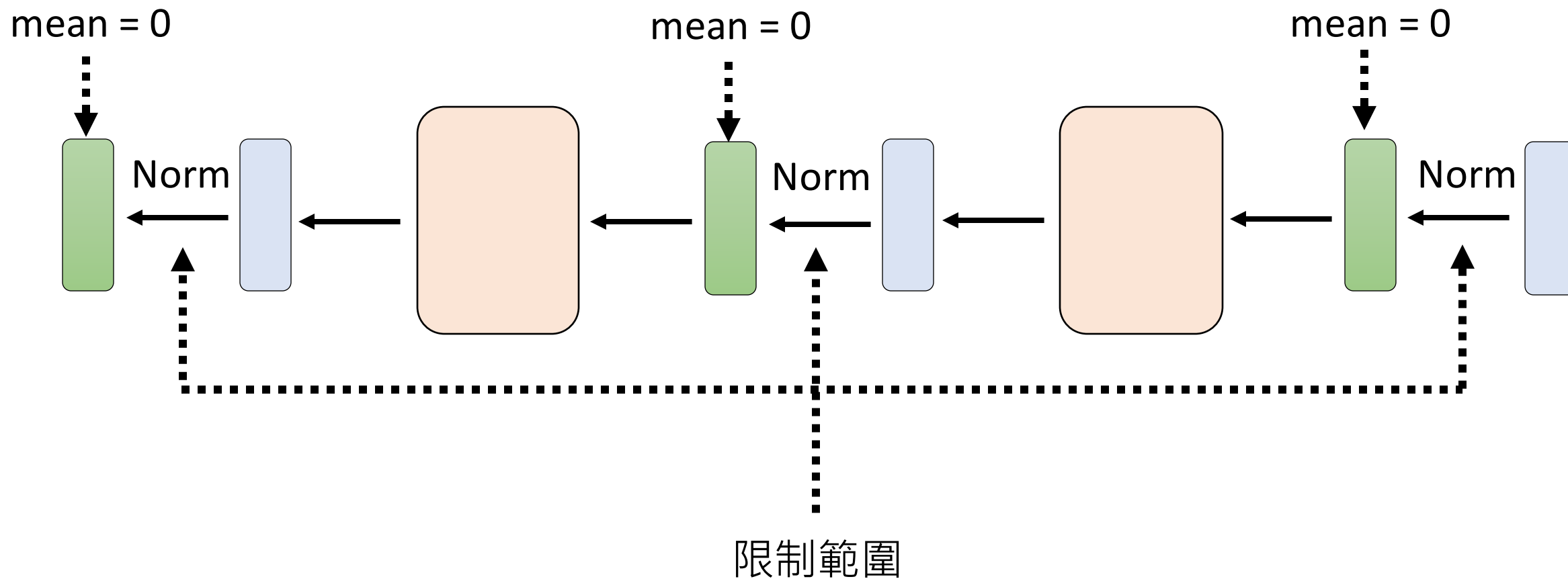
Source of image: <https://arxiv.org/abs/1712.09913>



With skip connection

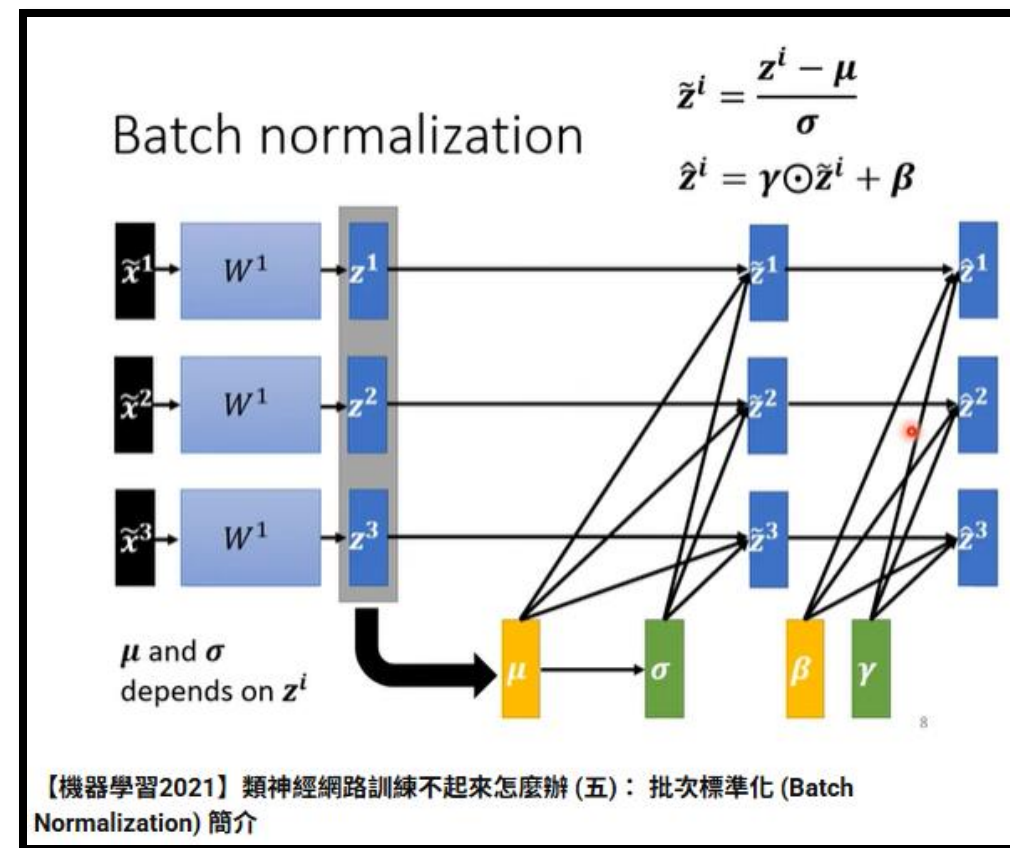
方法名	改了那一個步驟	帶來什麼好處
Skip Connection	改變函式搜尋範圍	Better Optimization

各種 Normalization



Normalization

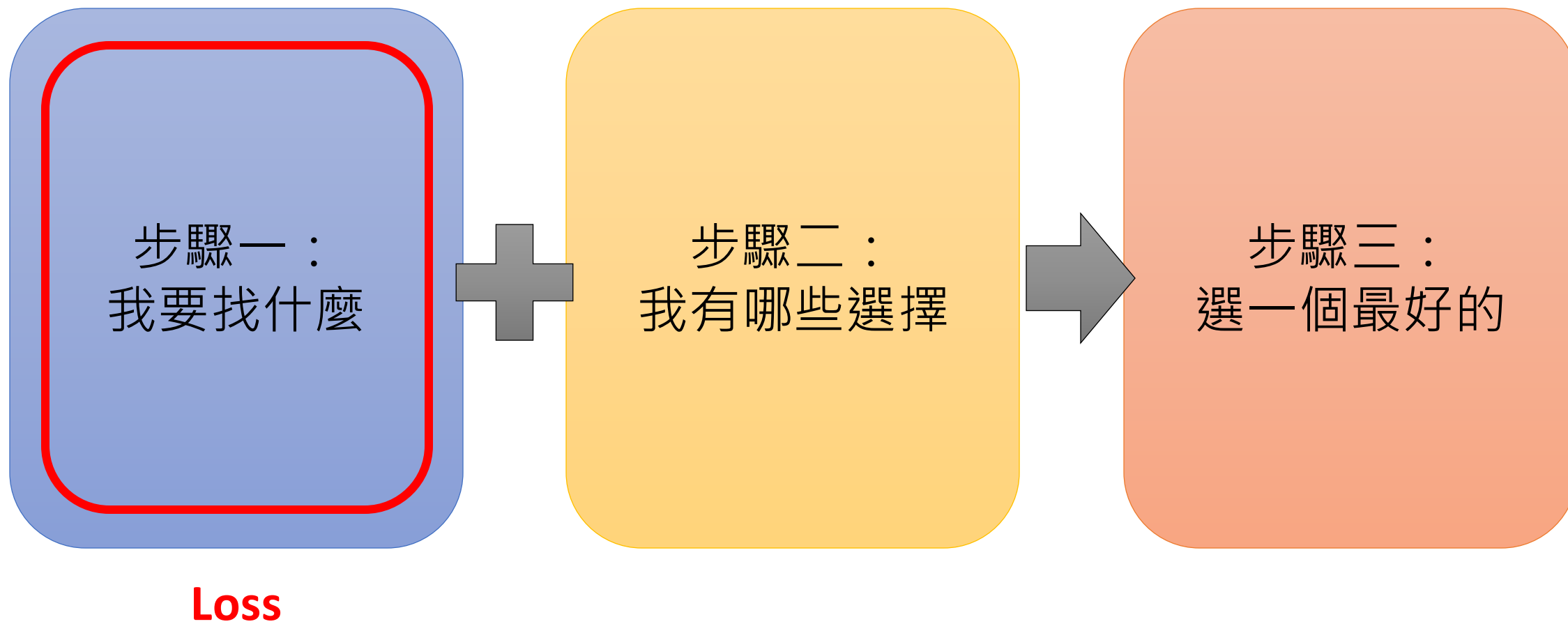
- Batch Renormalization
 - <https://arxiv.org/abs/1702.03275>
- Layer Normalization
 - <https://arxiv.org/abs/1607.06450>



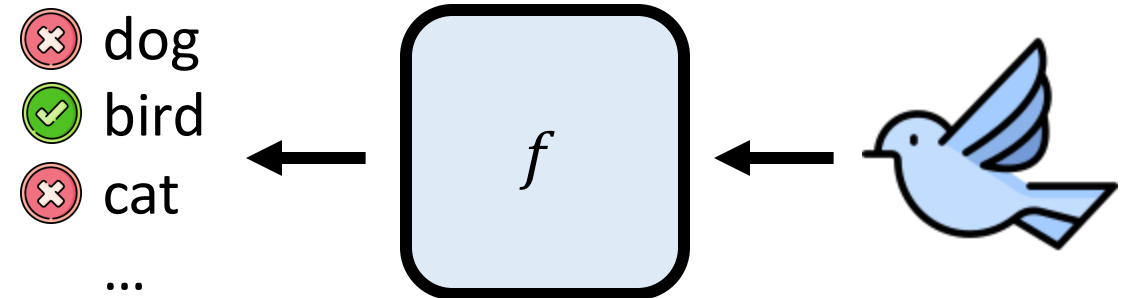
<https://youtu.be/BABPWOkSbLE?si=gzwtWMEBzbfg7sll>

方法名	改了那一個步驟	帶來什麼好處
Normalization	改變函式搜尋範圍	Better Optimization (Sometimes Better Generalization)

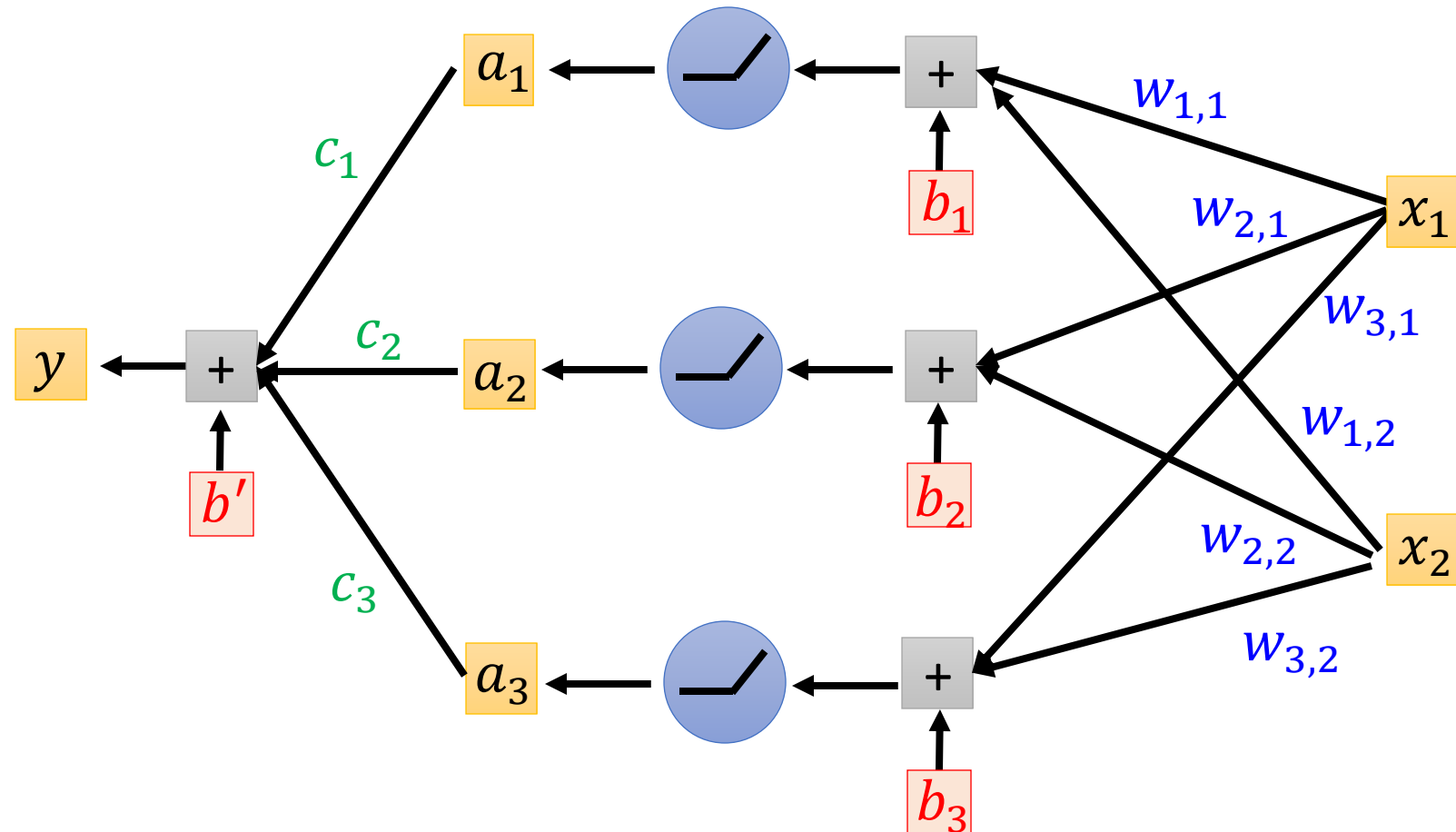
找函式步驟 3 + 1



Classification

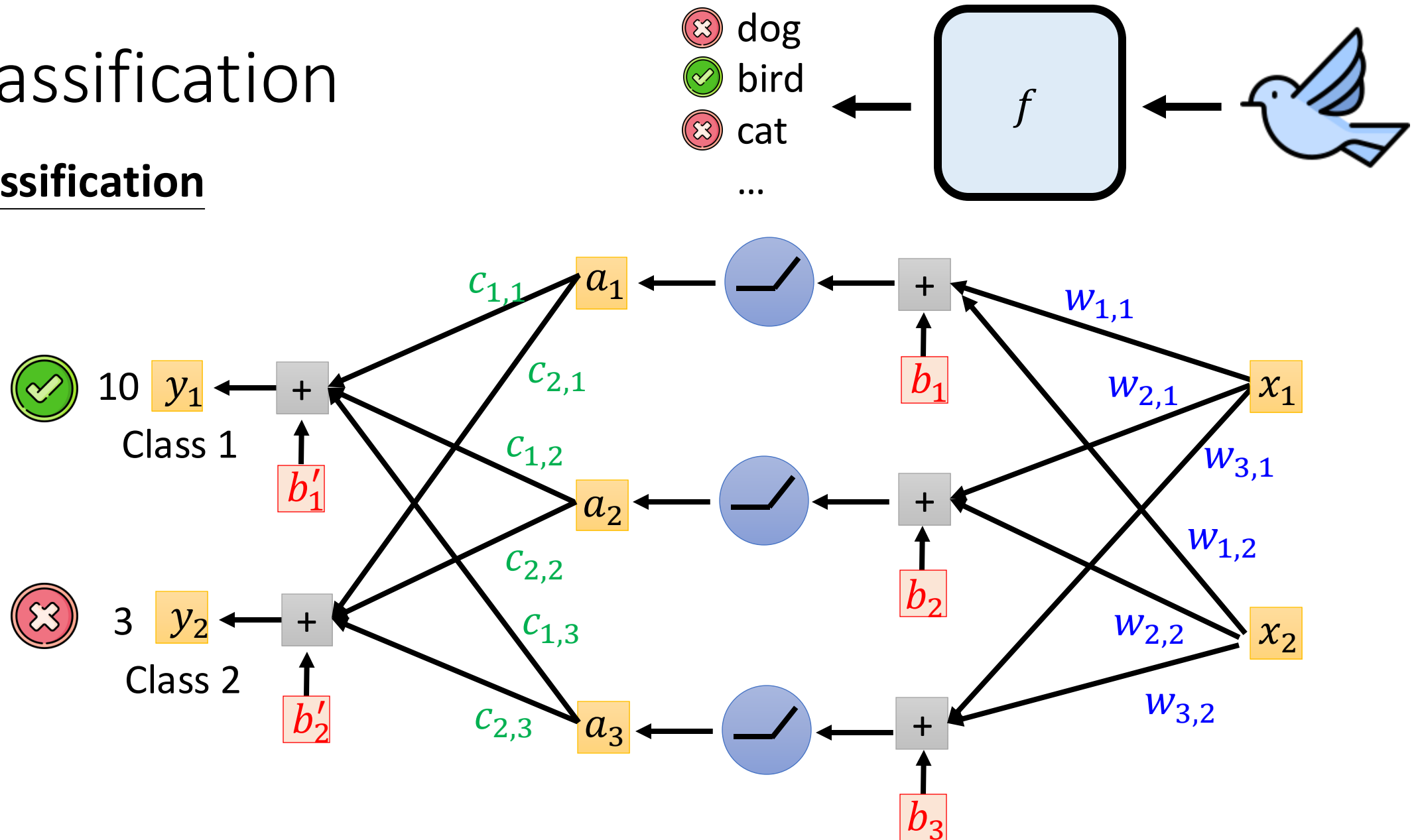


Regression



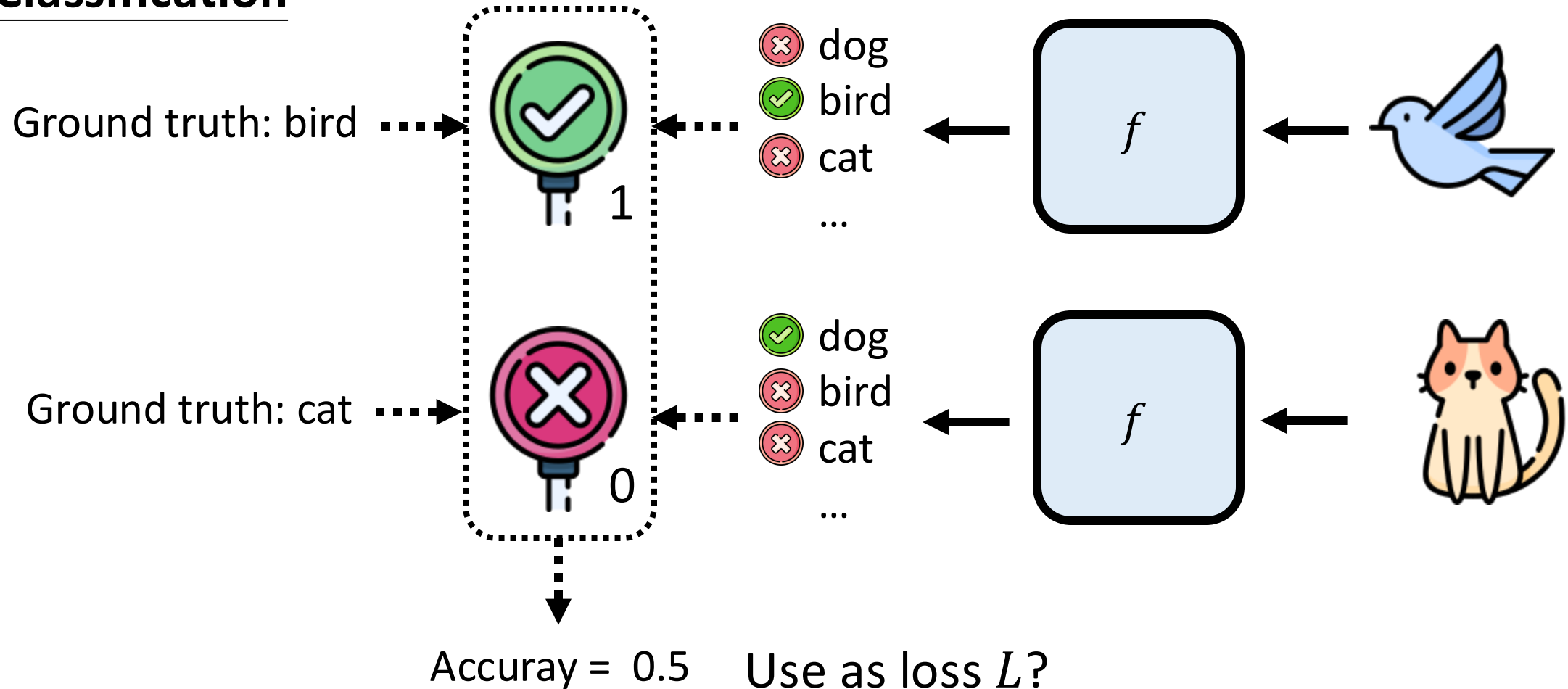
Classification

Classification



Loss = Evaluation Metrics?

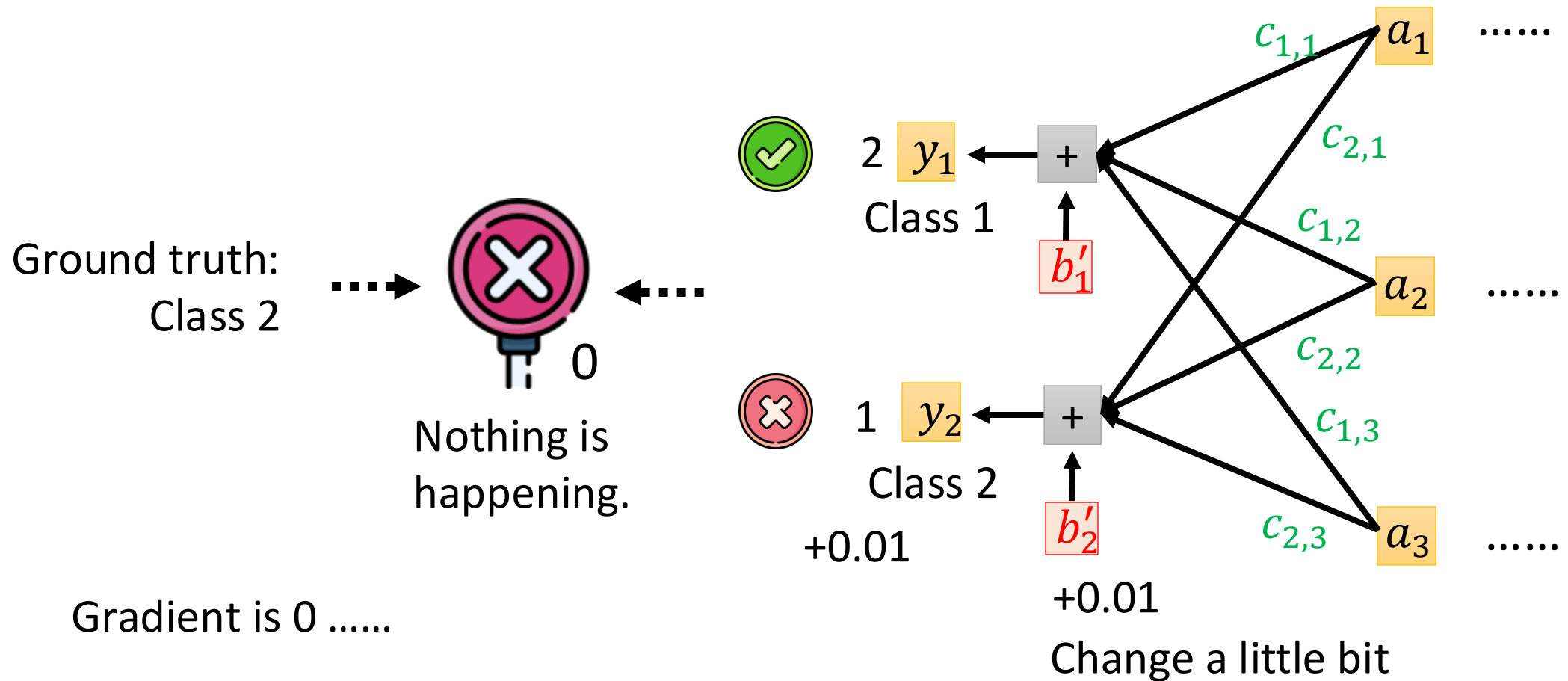
Classification



Loss = Evaluation Metrics?

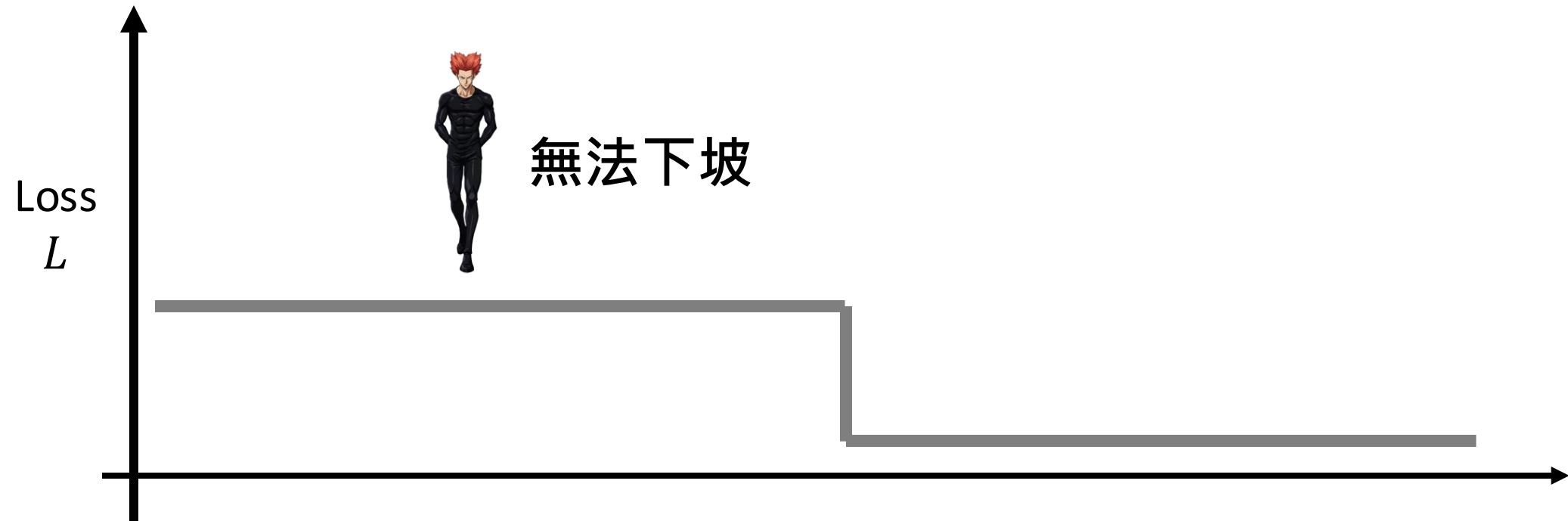
Classification

$$g = \nabla L(\theta)$$

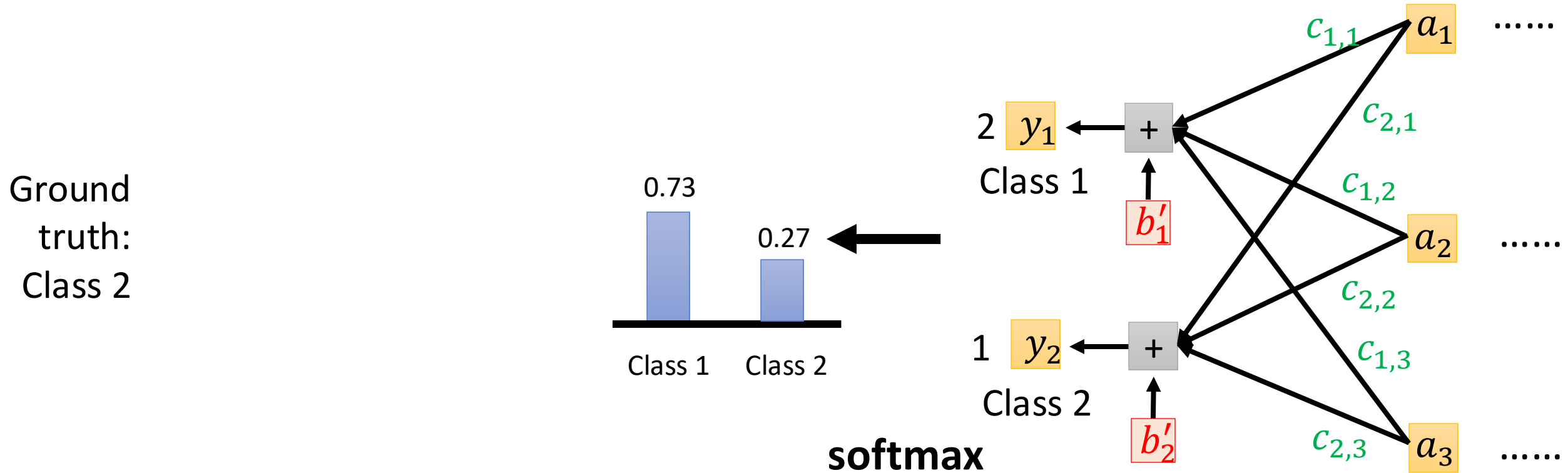


Loss = Evaluation Metrics?

- For classification, if we use accuracy as the loss function ...



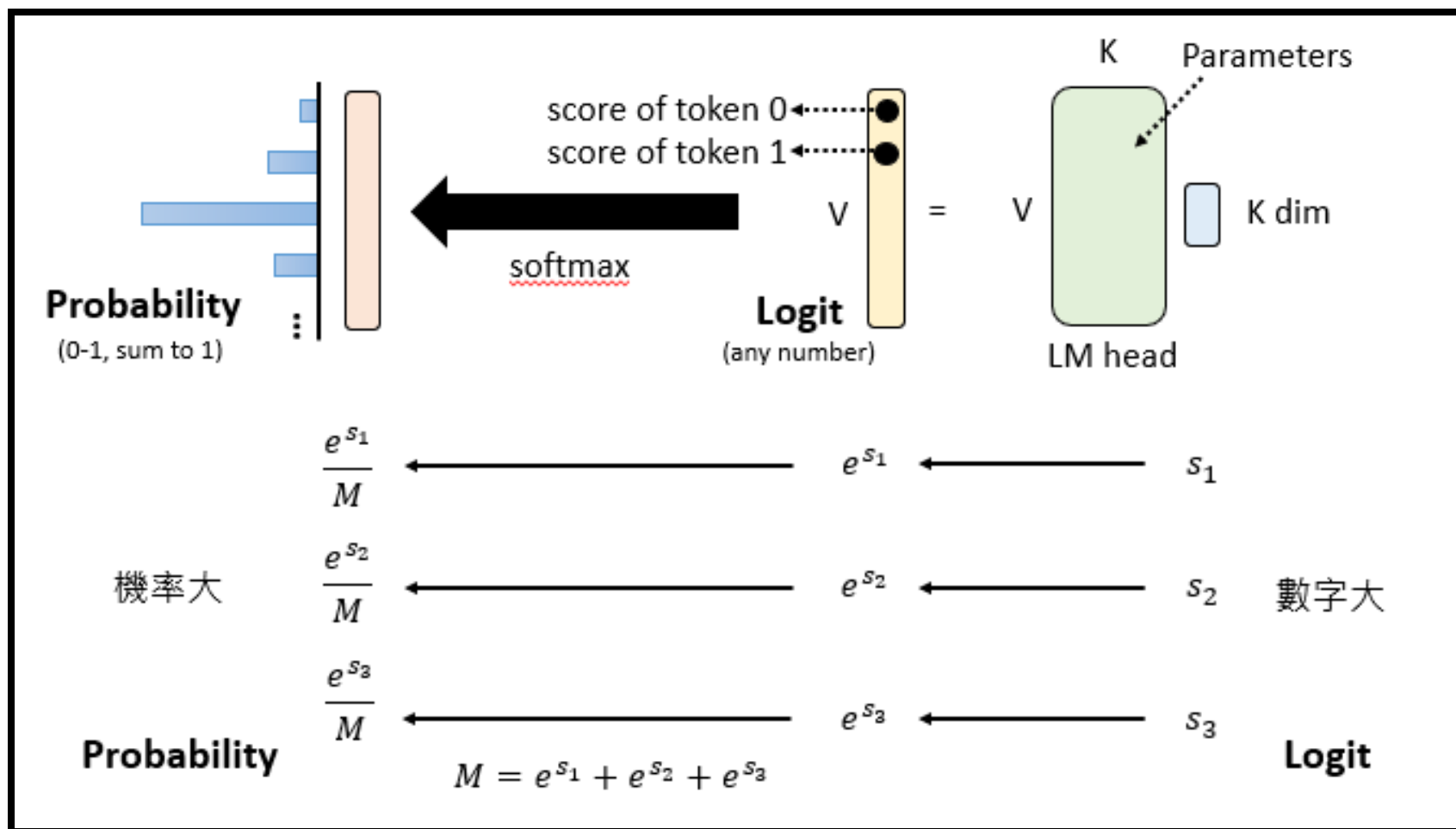
Cross-entropy as the loss in classification



什麼是 Softmax?

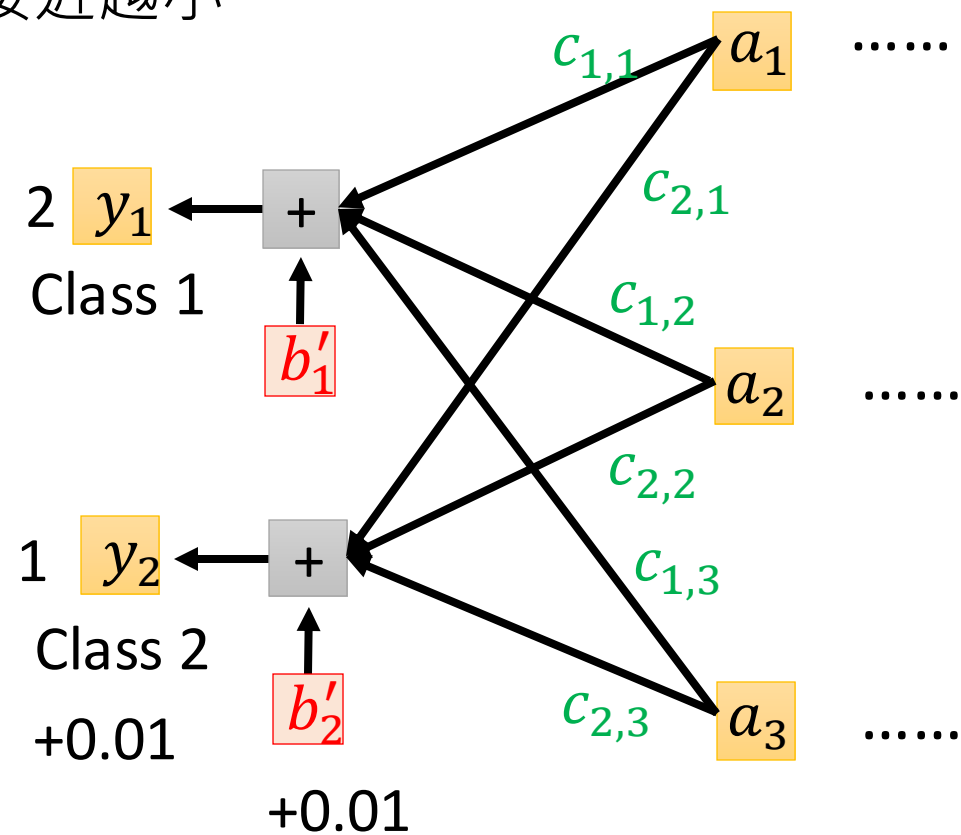
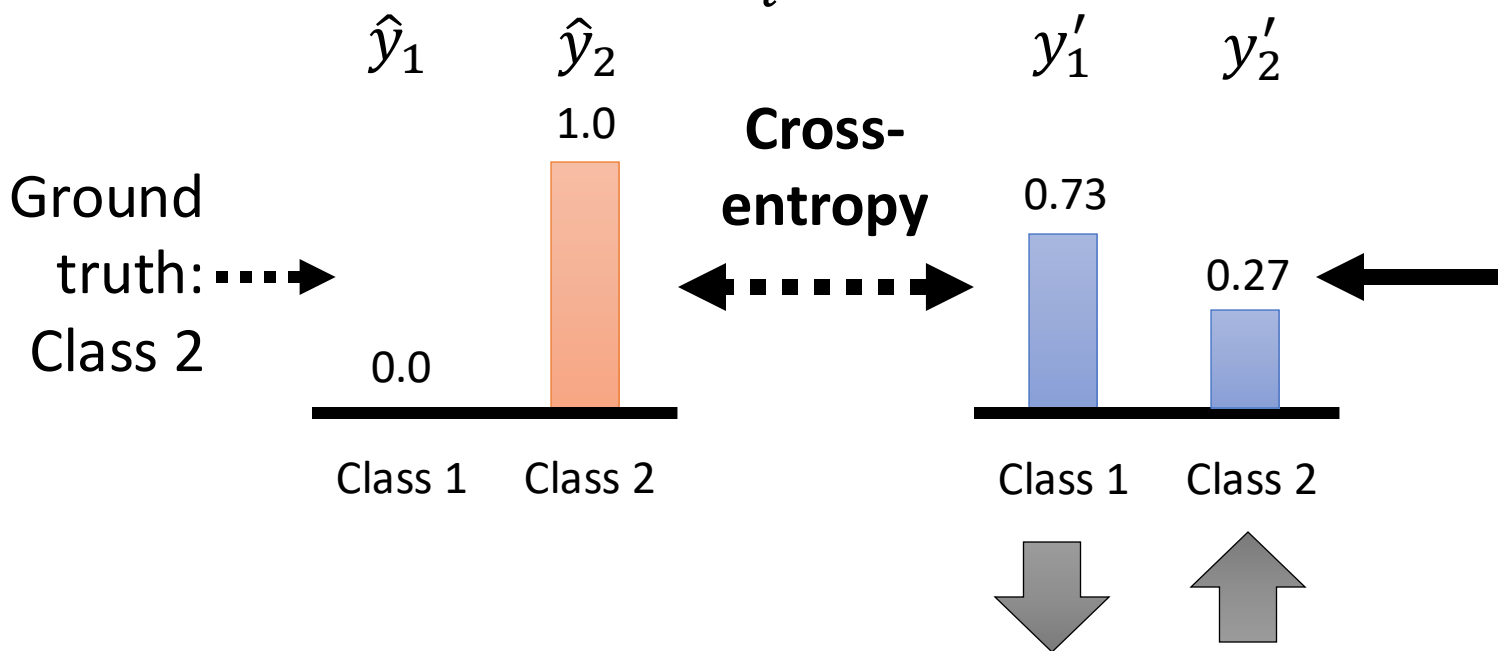
請見本課程第三講

<https://youtu.be/8iFvM7WUUs8?si=VdfmHEdUEm1TJOIR>



$$l = - \sum_i \hat{y}_i \log(y'_i)$$

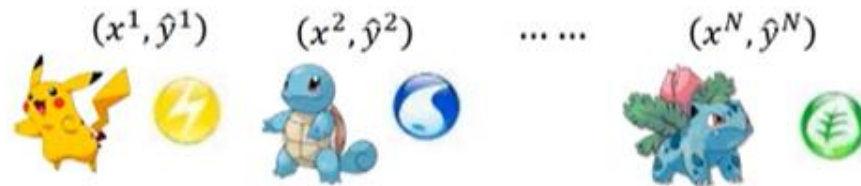
兩個分布越接近越小



Why Cross-entropy?

How to do Classification

- Training data for Classification



Classification as Regression?

Binary classification as example

Training: Class 1 means the target is 1; Class 2 means the target is -1

Testing: closer to 1 → class 1; closer to -1 → class 2

Created with EverCam
<http://www.camdemy.com>

ML Lecture 4: Classification

<https://youtu.be/fZAZUYEeIMg>

Logistic Regression

Step 1: $f_{w,b}(x) = \sigma \left(\sum_i w_i x_i + b \right)$
Output: between 0 and 1

Linear Regression

$f_{w,b}(x) = \sum_i w_i x_i + b$
Output: any value

Step 2: Training data: $(x^n, \hat{y}^n)$
 $\hat{y}^n$: 1 for class 1, 0 for class 2
 $L(f) = \sum_n C(f(x^n), \hat{y}^n)$

Training data: $(x^n, \hat{y}^n)$
 $\hat{y}^n$: a real number
 $L(f) = \frac{1}{2} \sum_n (f(x^n) - \hat{y}^n)^2$

Cross entropy:

$$C(f(x^n), \hat{y}^n) = -[\hat{y}^n \ln f(x^n) + (1 - \hat{y}^n) \ln(1 - f(x^n))]$$

Question: Why don't we simply use square error as linear regression?

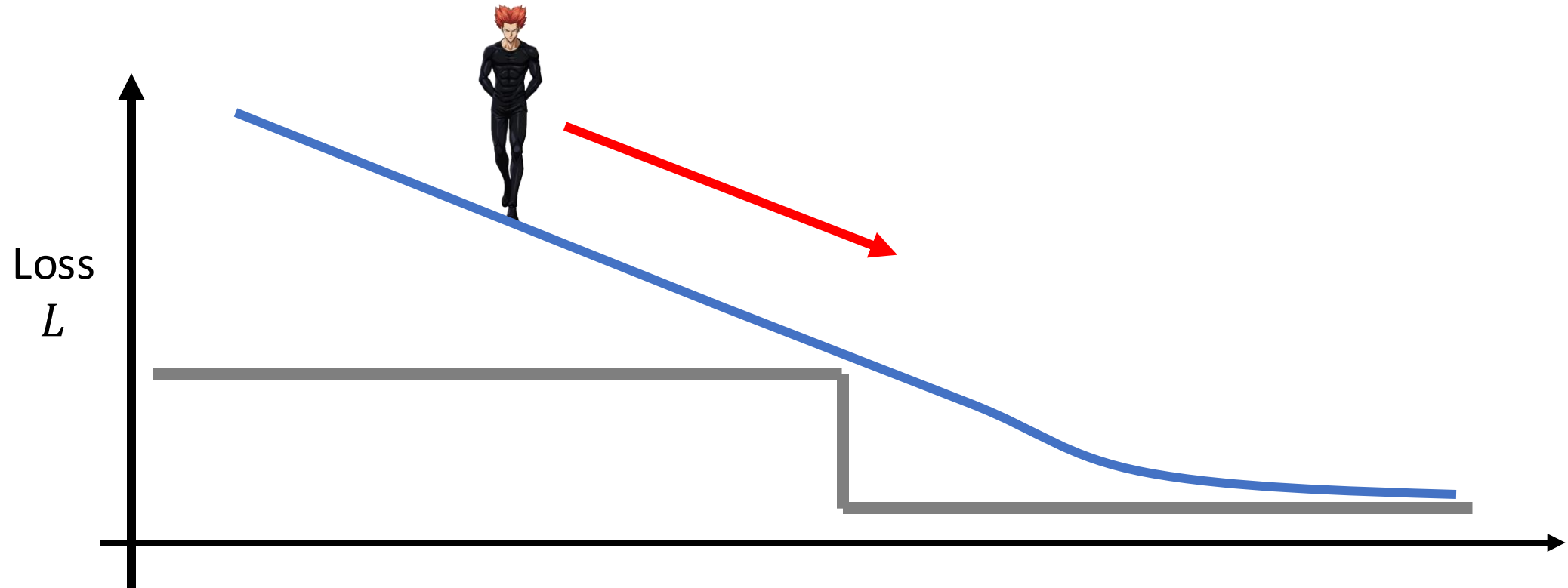
Created with EverCam
<http://www.camdemy.com>

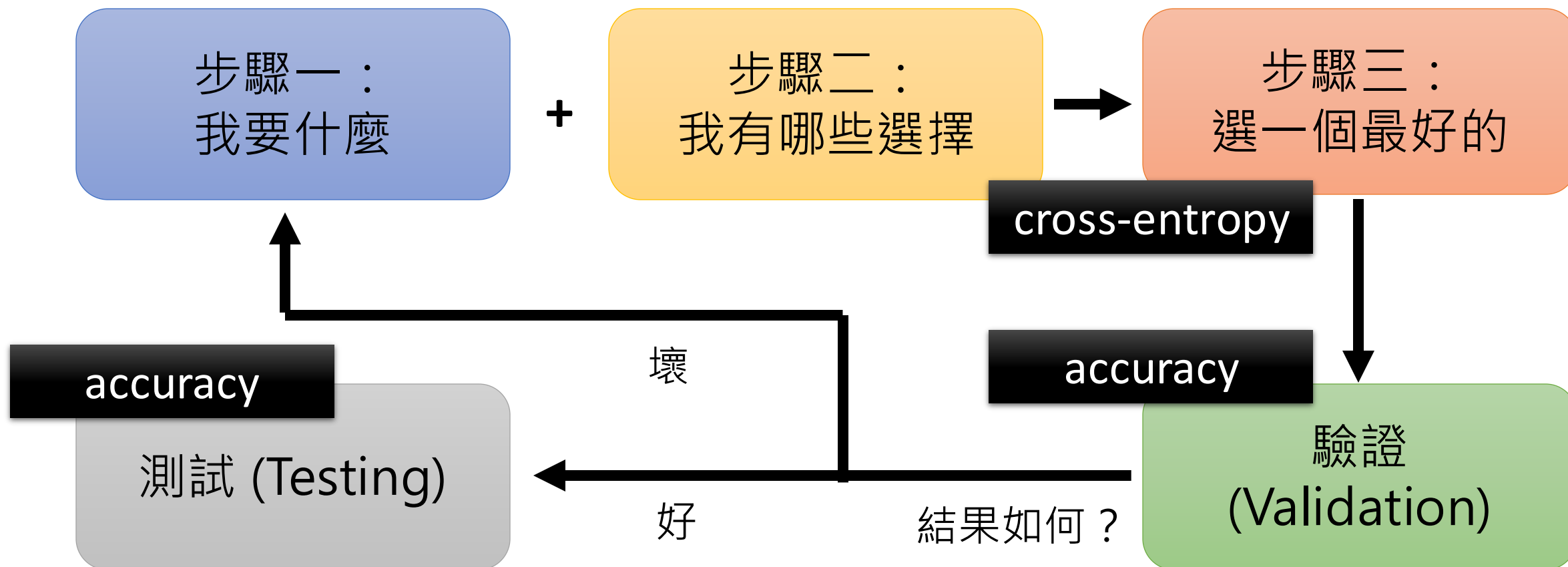
ML Lecture 5: Logistic Regression

<https://youtu.be/hSXFuypLukA>

Loss = Evaluation Metrics?

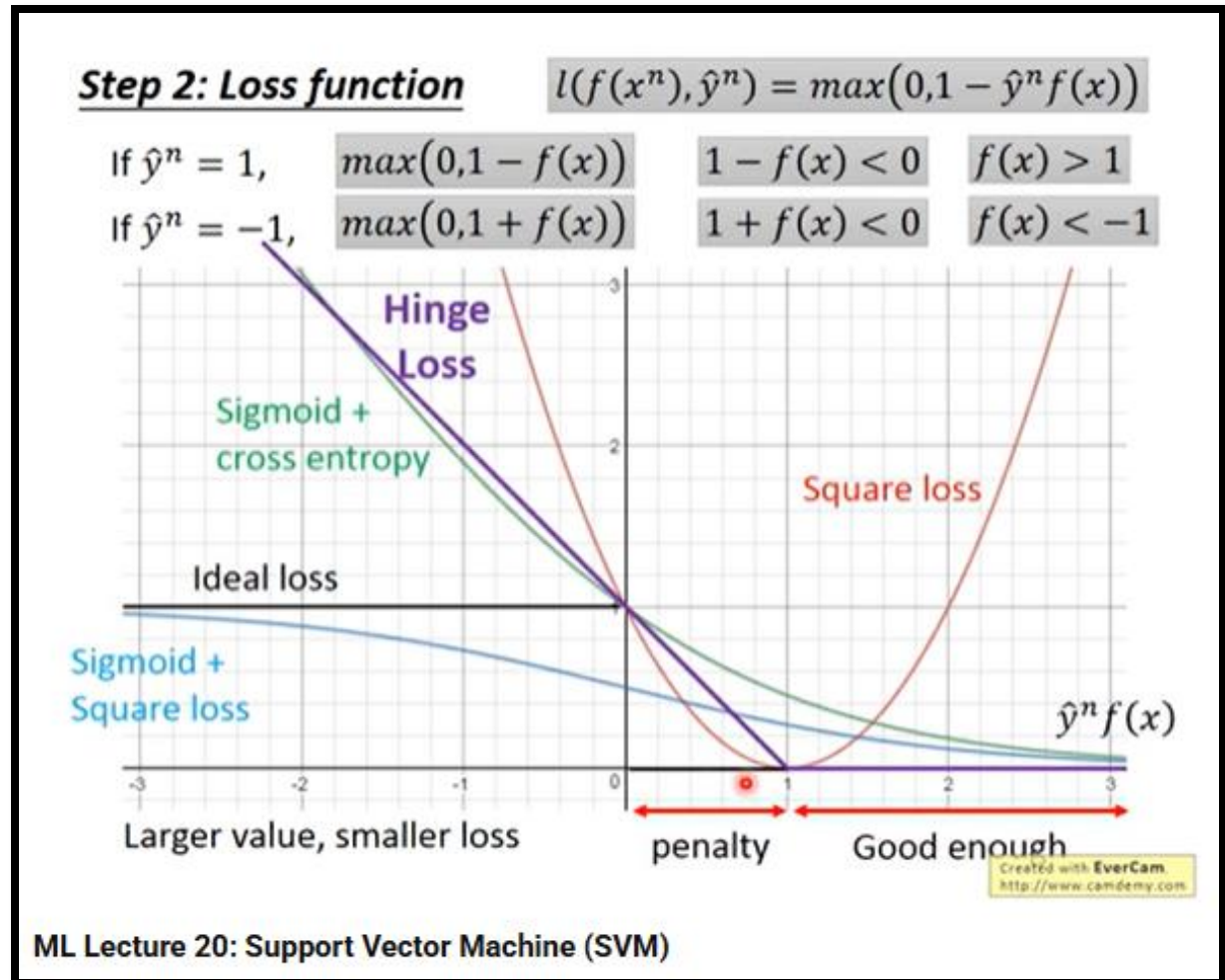
- For classification, if we use cross-entropy as the loss function ...





方法名	改了那一個步驟	帶來什麼好處
Do not use accuracy as loss	我要找甚麼	Better Optimization (Not for Generalization)

有沒有其他可能性？



<https://youtu.be/QSEPStBgwRQ?si=5oBXckoLuqbaroKY>

Need More Training Data ...

Overfitting

$$L = \sum_{x \sim \text{Train}} l(f(x), \hat{y}) \quad \cdots \cdots \cdots \rightarrow \quad L = \sum_{x \sim \text{Test}} l(f(x), \hat{y})$$



Collecting more training data

(be careful about the representative)

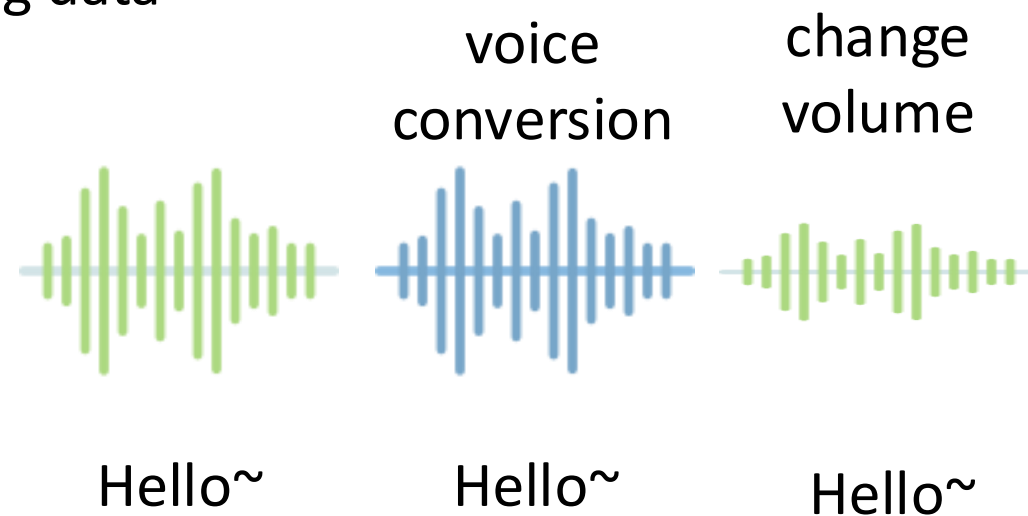
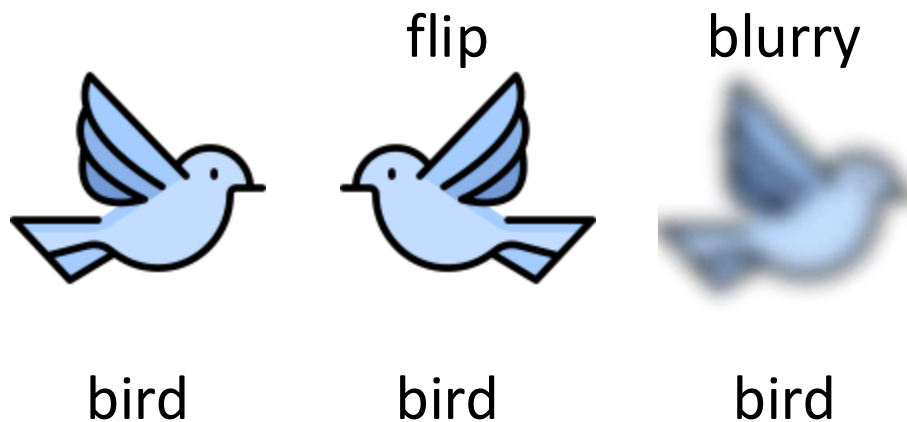
方法名	改了那一個步驟	帶來什麼好處
More training data	我要找甚麼	Better Generalization Not for Optimization

Data Augmentation

Overfitting

$$L = \sum_{x \sim \text{Train}} l(f(x), \hat{y}) \quad \cdots \cdots \cdots \rightarrow \quad L = \sum_{x \sim \text{Test}} l(f(x), \hat{y})$$

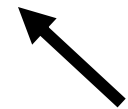
Collecting more training data



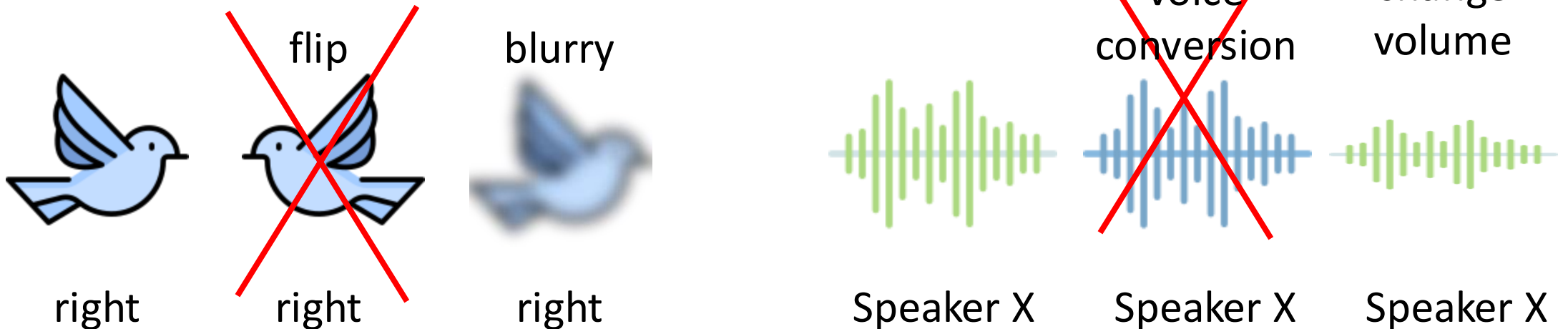
Be careful about what you augment

Overfitting

$$L = \sum_{x \sim \text{Train}} l(f(x), \hat{y}) \quad \cdots \cdots \cdots \rightarrow \quad L = \sum_{x \sim \text{Test}} l(f(x), \hat{y})$$

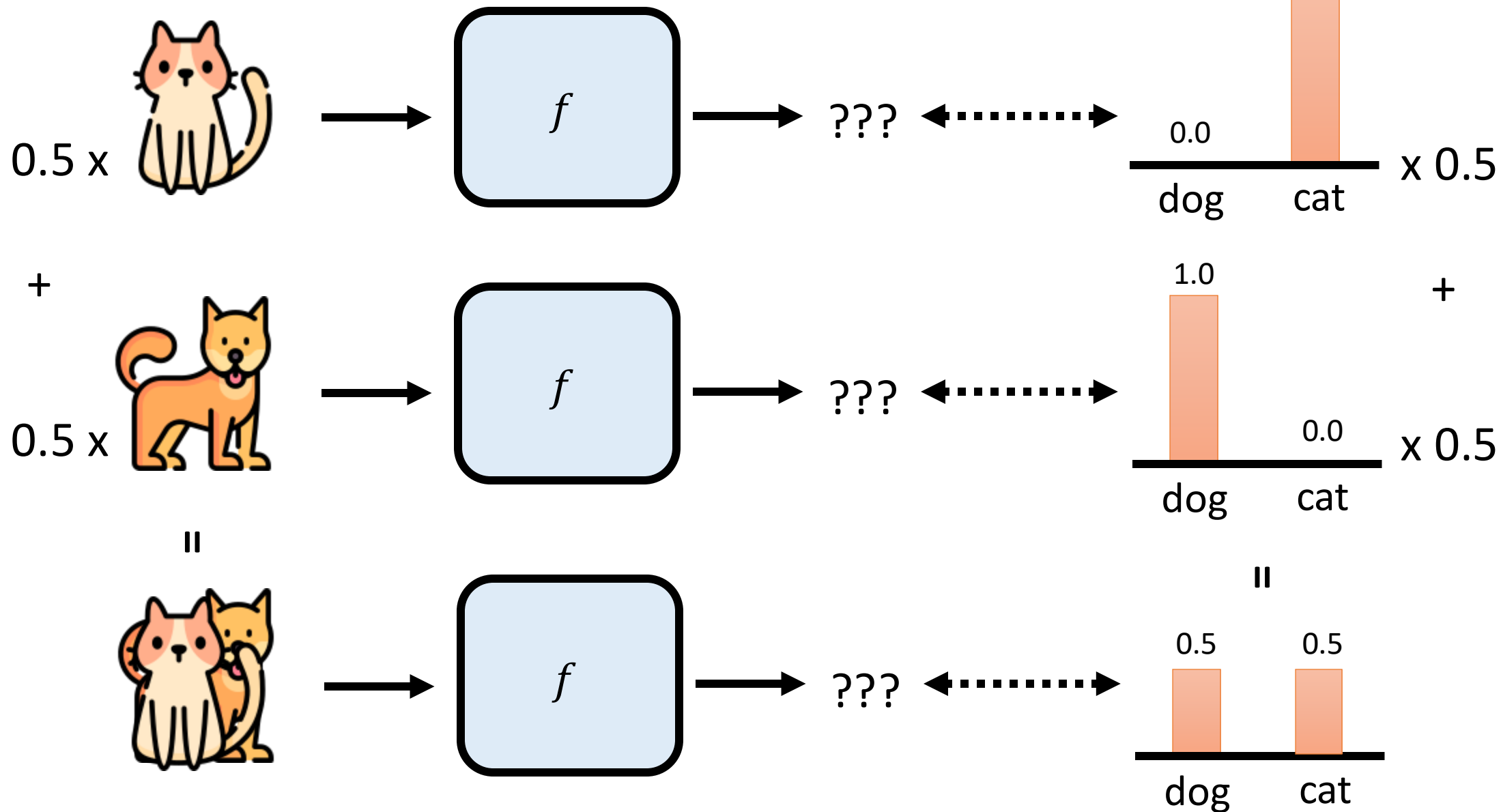


Collecting more training data



Mixup

<https://arxiv.org/abs/1710.09412>



Data Augmentation

Overfitting

$$L = \sum_{x \sim \text{Train}} l(f(x), \hat{y}) \quad \cdots \cdots \cdots \rightarrow \quad L = \sum_{x \sim \text{Test}} l(f(x), \hat{y})$$

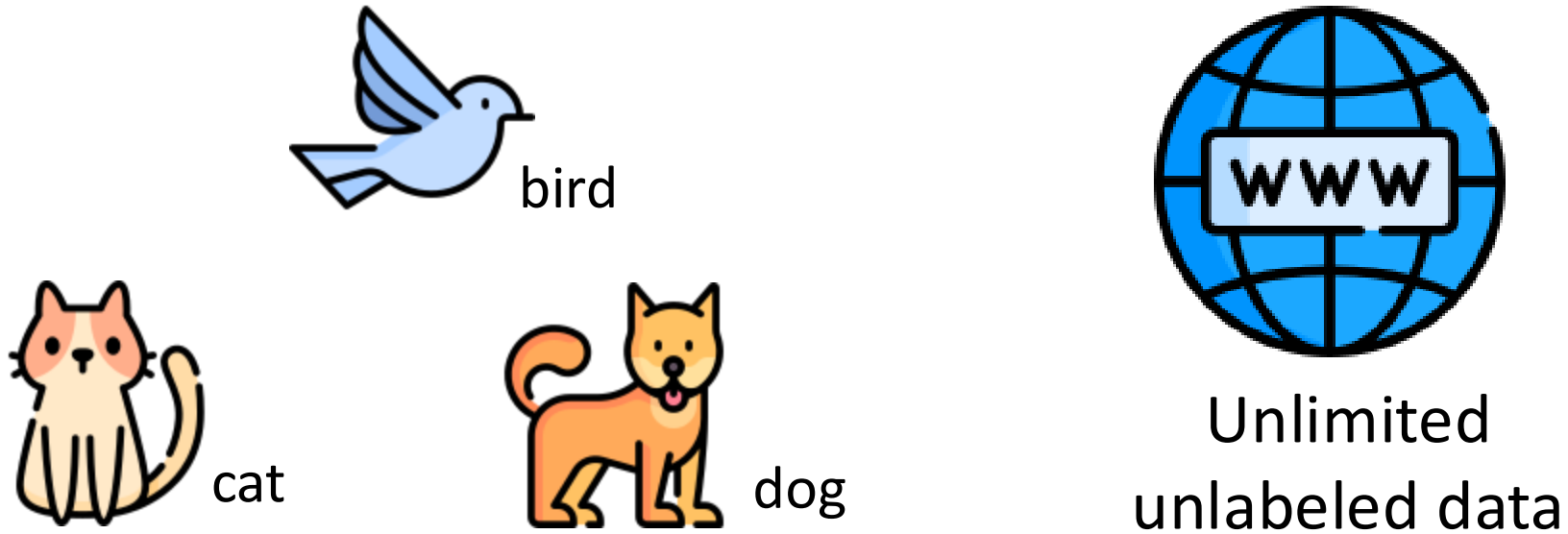


Collecting more training data

(be careful about the representative)

方法名	改了那一個步驟	帶來什麼好處
Data Augmentation (e.g. Mixup)	我要找甚麼	Better Generalization Not for Optimization

Semi-supervised Learning

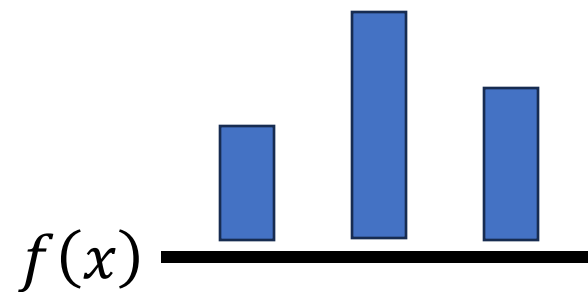


$$L = \sum_{x \sim \text{Train}} l(f(x), \hat{y}) + \lambda \sum_{x \sim \text{Unlabel}} \text{?????}$$

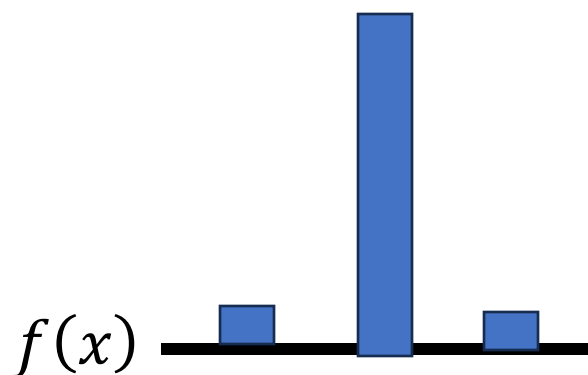
$$L = \sum_{x \sim \text{Train}} l(f(x), \hat{y}) + \lambda \sum_{x \sim \text{Unlabel}} l'(f(x))$$

$\downarrow$ $\searrow$
 entropy distribution

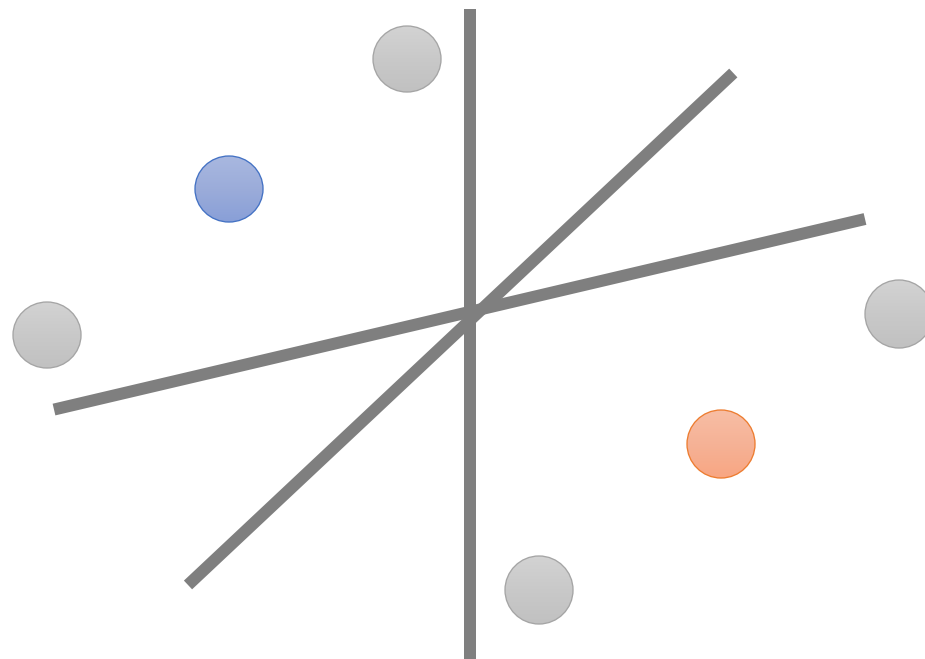
larger entropy



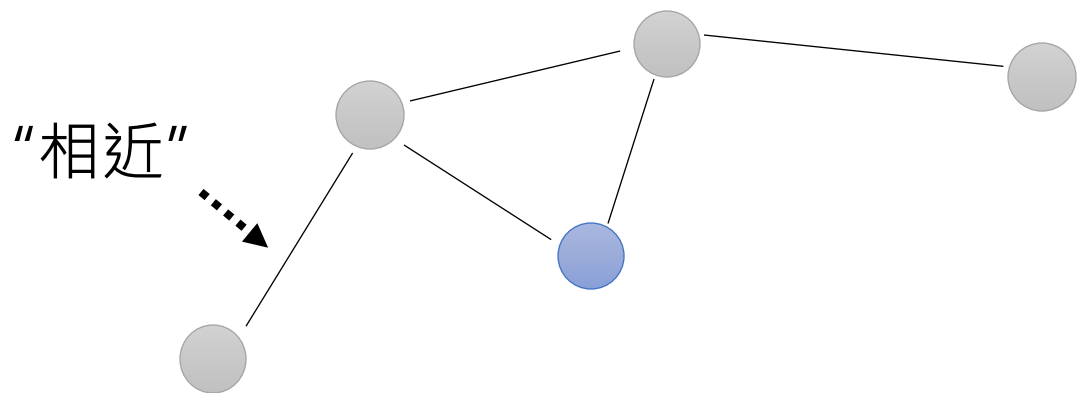
smaller entropy



非黑即白
的世界

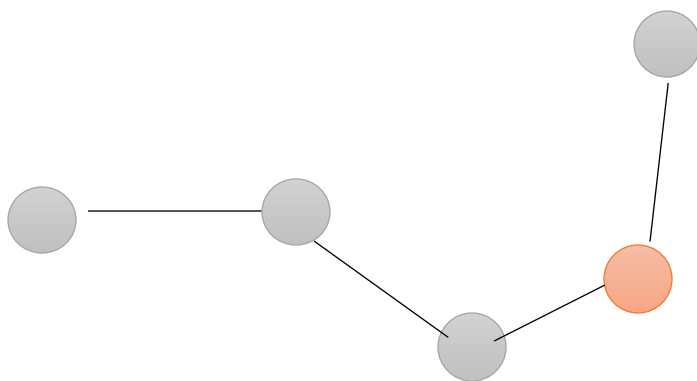


$$L = \sum_{x \sim \text{Train}} l(f(x), \hat{y}) + \lambda \sum_{x, x' \sim \text{Unlabel}} g(x, x') l'(f(x), f(x'))$$



$$g(x, x') = \begin{cases} 1 & \text{connect} \\ 0 & \text{otherwise} \end{cases}$$

物以類聚
的世界



$$l'(f(x), f(x'))$$

$f(x), f(x')$ 的 "距離"

Semi-supervised Learning

$$L = \sum_{x \sim \text{Train}} l(f(x), \hat{y}) + \lambda \sum_{x \sim \text{Unlabel}} l'(f(x))$$

$$L = \sum_{x \sim \text{Train}} l(f(x), \hat{y}) + \lambda \sum_{x, x' \sim \text{Unlabel}} g(x, x') l'(f(x), f(x'))$$

方法名	改了那一個步驟	帶來什麼好處
Semi-supervised (e.g., Entropy, Graph)	我要找甚麼	Better Generalization Not for Optimization

Parameter Regularization

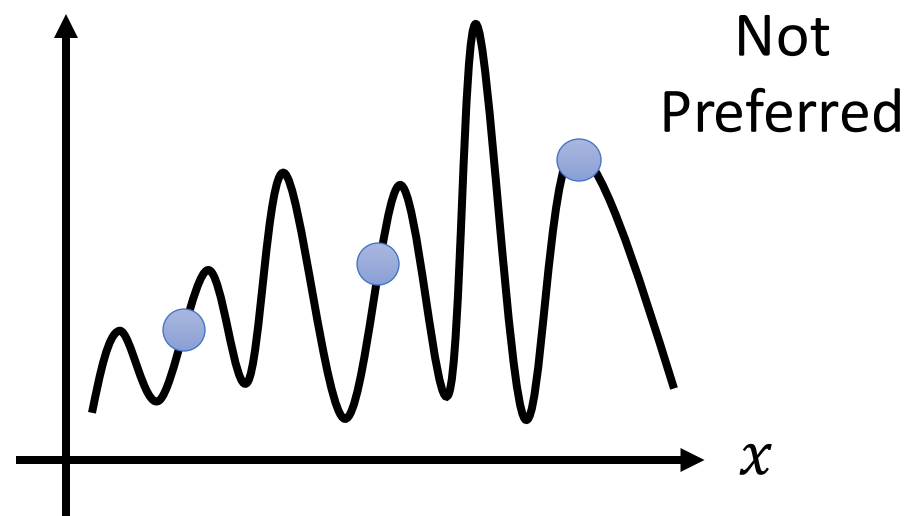
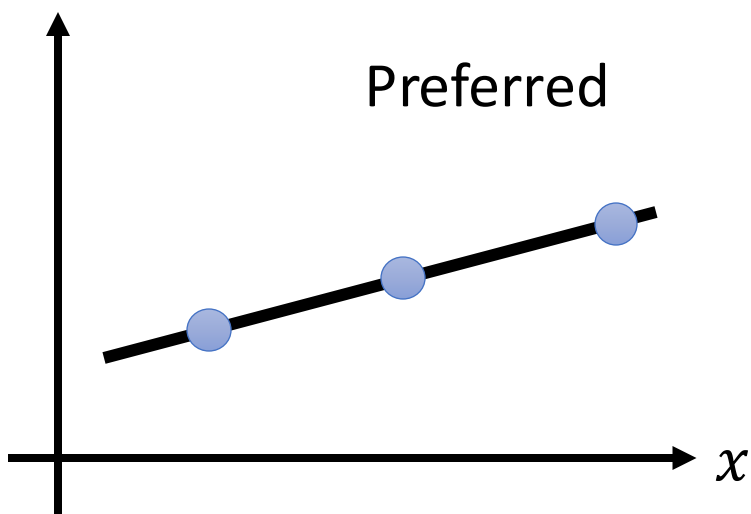
希望 Parameter 越接近 0 越好

$$L'(\boldsymbol{\theta}) = \boxed{L(\boldsymbol{\theta})} + \lambda \boxed{R(\boldsymbol{\theta})}$$

跟資料有關

對於 Parameter 的偏好 (與資料無關)

$$R(\boldsymbol{\theta}) = \sum_i (\theta_i)^2$$



Parameter Regularization

希望 Parameter 越接近 0 越好

$$L'(\boldsymbol{\theta}) = \boxed{L(\boldsymbol{\theta})} + \lambda \boxed{R(\boldsymbol{\theta})}$$

跟資料有關

對於 Parameter 的偏好 (與資料無關)

$$R(\boldsymbol{\theta}) = \sum_i (\theta_i)^2$$

$$\mathbf{g} = \nabla L(\boldsymbol{\theta}) \quad \boldsymbol{\theta} \leftarrow \boldsymbol{\theta} - \eta \mathbf{g}$$

$$\begin{aligned} \mathbf{g}' &= \nabla L'(\boldsymbol{\theta}) \quad \boldsymbol{\theta} \leftarrow \boldsymbol{\theta} - \eta \mathbf{g}' = \boldsymbol{\theta} - \eta (\mathbf{g} + 2\lambda \boldsymbol{\theta}) = \boldsymbol{\theta} - \eta \mathbf{g} - 2\eta \lambda \boldsymbol{\theta} \\ &= \nabla L(\boldsymbol{\theta}) + \lambda \nabla R(\boldsymbol{\theta}) = \nabla L(\boldsymbol{\theta}) + 2\lambda \boldsymbol{\theta} \quad \mathbf{g} \\ &= \underline{(1 - 2\eta \lambda)} \boldsymbol{\theta} - \eta \mathbf{g} \\ &\quad \text{weight decay} \end{aligned}$$

Parameter Regularization

希望 Parameter 越接近 0 越好

$$L'(\boldsymbol{\theta}) = \boxed{L(\boldsymbol{\theta})} + \lambda \boxed{R(\boldsymbol{\theta})}$$

跟資料有關

對於 Parameter 的偏好 (與資料無關)

$$R(\boldsymbol{\theta}) = \sum_i (\theta_i)^2$$

方法名	改了那一個步驟	帶來什麼好處
Parameter Regularization	我要找甚麼	Better Generalization Not for Optimization
Weight Decay	找最好的函式	Better Generalization Not for Optimization

AdamW

<https://arxiv.org/abs/1711.05101>

Momentum

For each dimension i : $m_i^t = \beta m_i^{t-1} + (1 - \beta) g_i^t$

RMSProp

For each dimension i : $\sigma_i^t = \sqrt{\alpha (\sigma_i^{t-1})^2 + (1 - \alpha) (g_i^t)^2}$

Adam

$$\theta_i^{t+1} \leftarrow \theta_i^t - \frac{\eta^t}{\sigma_i^t} m_i^t$$

AdamW

$$\theta_i^{t+1} \leftarrow \underline{(1 - 2\eta^t \lambda)} \theta_i^t - \frac{\eta^t}{\sigma_i^t} m_i^t$$

方法名	改了那一個步驟	帶來什麼好處
AdamW	找最好的函式	Better Generalization (cf. Adam) Better Optimization (cf. Vanilla)

方法名	改了那一個步驟	帶來什麼好處
Adagrad, RMSProp, Momentum, Adam, etc.	找最好的函式	Better Optimization
AdamW	找最好的函式	Better Generalization (cf. Adam) Better Optimization (cf. Vanilla)
Dropout	找最好的函式	Better Generalization
Weight Decay	找最好的函式	Better Generalization
Initialization (e.g., pre-train)	找最好的函式	Better Optimization, Better Generalization
CNN (e.g., for image)	改變函式搜尋範圍	Better Generalization
Skip Connection	改變函式搜尋範圍	Better Optimization
Normalization	改變函式搜尋範圍	Better Optimization (Sometimes Better Generalization)
Do not use accuracy as loss	我要找什麼	Better Optimization
More training data	我要找什麼	Better Generalization
Data Augmentation (e.g. Mixup)	我要找什麼	Better Generalization
Semi-supervised (e.g., Entropy, Graph)	我要找什麼	Better Generalization
Parameter Regularization	我要找甚麼	Better Generalization