



AI and Deep Learning

1 Course Introduction

Zhonglei Wang

WISE, SOE, CHOW

XMU

(Version v1)

Contents

1. What we have learnt

2. Learning goals

3. Cross entropy

What we have learnt

1. Types of datasets:

- Training data of size $n : \{(\mathbf{x}_i, y_i) : i = 1, \dots, n\}$
 - ▷ $\mathbf{x}_i \in \mathbb{R}^{d \times 1}$: Feature for the i th example
 - ▷ $y_i \in \mathbb{R}$: Label for the i th example
- Test data

2. Regressions:

- Linear regression: $y_i = b + \mathbf{x}_i^T \cdot \mathbf{w} + \epsilon_i \quad (i = 1, \dots, n)$
- Logistic regression: $\text{logit}P(y_i = 1 \mid \mathbf{x}_i) = b + \mathbf{x}_i^T \cdot \mathbf{w} \quad (i = 1, \dots, n)$
 - ▷ $\text{logit}p = \log(p) - \log(1 - p), \quad \forall p \in (0, 1)$
- Multinomial logistic regression (softmax regression)

Notations

1. n : Number of examples in the training data
2. d : Dimension of features
3. $\mathbf{X} = (\mathbf{x}_1, \dots, \mathbf{x}_n)^T \in \mathbb{R}^{n \times d}$: Design matrix
4. $\mathbf{Y} = (y_1, \dots, y_n)^T \in \mathbb{R}^{n \times 1}$: Vector of labels

Linear regression -- Model setup

1. Training dataset

- Feature $\mathbf{x}_i \in \mathbb{R}^{d \times 1}$, label $y_i \in \mathbb{R}$ ($i = 1, \dots, n$)

2. Goal

- Learn the relationship between feature $\mathbf{x}$ and label y

3. True model (used to generate data)

- $y_i = b_0 + \mathbf{x}_i^T \cdot \mathbf{w}_0 + \epsilon_i$ ($i = 1, \dots, n$)
 - ▷ $\boldsymbol{\theta}_0 = (b_0, \mathbf{w}_0^T)^T$: True model parameters
 - ▷ ϵ_i : Noise satisfying $E(\epsilon_i \mid \mathbf{x}_i) = 0$, and $\text{var}(\epsilon_i \mid \mathbf{x}_i) = \sigma^2 > 0$

Linear regression -- Model setup

1. (Assumed) Proposed model (used to fit data)

- $E(y_i \mid \mathbf{x}_i) = b + \mathbf{x}_i^T \cdot \mathbf{w} \quad (i = 1, \dots, n)$
 - ▷ $\boldsymbol{\theta} = (b, \mathbf{w}^T)^T$: (Proposed) Model parameters
 - ▷ By fitting a model, we mean to estimate model parameters by training data
 - ▷ Have the same form with the true model, but we need to estimate parameters

Linear regression -- Parameter estimation

1. **Loss function:** For the i th example,

- l_2 -loss: $\mathcal{L}(y_i, \hat{y}_i) = (y_i - \hat{y}_i)^2 / 2$
- $\hat{y}_i = b + \mathbf{x}_i^T \cdot \mathbf{w}$
- $\hat{y}_i$ is a function of $\boldsymbol{\theta} = (b, \mathbf{w}^T)^T$
- $\mathcal{L}(y_i, \hat{y}_i)$ is also a function of model parameters $\boldsymbol{\theta}$
- However, in the above expressions, we omit its argument for simplicity
- The loss function is a convex function of the model parameter $\boldsymbol{\theta}$

Linear regression -- Parameter estimation

1. Cost function (Mean Squared Error, MSE)

$$\mathcal{J}(\boldsymbol{\theta}) = \frac{1}{n} \sum_{i=1}^n \mathcal{L}(y_i, \hat{y}_i) = \frac{1}{2n} \sum_{i=1}^n (y_i - \hat{y}_i)^2 = \frac{1}{2n} \sum_{i=1}^n (y_i - b - \mathbf{x}_i^T \cdot \mathbf{w})^2$$

2. Sometimes, we do not distinguish a loss function and a cost function

3. Solve

$$\frac{\mathcal{J}}{\partial b} = 0, \quad \frac{\mathcal{J}}{\partial \mathbf{w}} = 0$$

4. Pitfall: Low computational efficiency due to summation

5. **Question:** Why is that a disaster for Python to use for-loops for summation, especially when the sample size is extremely large?

Vectorization

1. Consider

$$\sum_{i=1}^m a_i b_i$$

- $a_i \in \mathbb{R}, \quad b_i \in \mathbb{R}$

2. For Python and other softwares, the computation efficiency of using for-loops for summation is quite low
3. However, matrix calculation is very efficient
4. Thus, we use vectorization for the summation

Vectorization

1. Vectorization:

$$\sum_{i=1}^m \mathbf{a}_i \mathbf{b}_i = \mathbf{a}^T \cdot \mathbf{b}$$

- $\mathbf{a} = (\mathbf{a}_1, \dots, \mathbf{a}_m)^T \in \mathbb{R}^{m \times 1}$
- $\mathbf{b} = (\mathbf{b}_1, \dots, \mathbf{b}_m)^T \in \mathbb{R}^{m \times 1}$

2. Avoid for-loops when linear algebra can be used

Vectorization

1. Example 1: Consider

$$\sum_{i=1}^m a_i$$

- $a_i \in \mathbb{R}$

2. Vectorization:

$$\sum_{i=1}^m a_i = \sum_{i=1}^m a_i \times \mathbf{1} = \mathbf{a}^T \cdot \mathbf{1} (= \mathbf{1}^T \cdot \mathbf{a})$$

- $\mathbf{a} = (a_1, \dots, a_m)^T \in \mathbb{R}^{m \times 1}$
- $\mathbf{1} = (1, \dots, 1)^T \in \mathbb{R}^{m \times 1}$: Vector of 1's with length m

Vectorization

1. Example 2: Consider

$$\sum_{i=1}^m \mathbf{a}_i \cdot \mathbf{b}_i$$

- $\mathbf{a}_i \in \mathbb{R}$: Scalar, which may be 1
- $\mathbf{b}_i \in \mathbb{R}^{k \times 1}$

2. Vectorization:

$$\sum_{i=1}^m \mathbf{a}_i \cdot \mathbf{b}_i = \mathbf{B} \cdot \mathbf{a}$$

- $\mathbf{a} = (\mathbf{a}_1, \dots, \mathbf{a}_m)^T \in \mathbb{R}^{m \times 1}$
- $\mathbf{B} = (\mathbf{b}_1, \dots, \mathbf{b}_m) \in \mathbb{R}^{k \times m}$

Vectorization

1. Example 3: Consider

$$\sum_{i=1}^m \mathbf{a}_i \cdot \mathbf{b}_i^T$$

- $\mathbf{a}_i \in \mathbb{R}$: Scalar, which may be 1
- $\mathbf{b}_i \in \mathbb{R}^{k \times 1}$

2. Vectorization:

$$\sum_{i=1}^m \mathbf{a}_i \cdot \mathbf{b}_i^T = \mathbf{a}^T \cdot \mathbf{B}$$

- $\mathbf{a} = (\mathbf{a}_1, \dots, \mathbf{a}_m)^T \in \mathbb{R}^{m \times 1}$
- $\mathbf{B} = (\mathbf{b}_1, \dots, \mathbf{b}_m)^T \in \mathbb{R}^{m \times k}$

Vectorization

1. Example 4: Consider

$$\sum_{i=1}^m \mathbf{a}_i \cdot \mathbf{b}_i^T$$

- $\mathbf{a}_i \in \mathbb{R}^{k \times 1}$
- $\mathbf{b}_i \in \mathbb{R}^{l \times 1}$

2. Vectorization:

$$\sum_{i=1}^m \mathbf{a}_i \cdot \mathbf{b}_i^T = \mathbf{A} \cdot \mathbf{B}$$

- $\mathbf{A} = (\mathbf{a}_1, \dots, \mathbf{a}_m) \in \mathbb{R}^{k \times m}$
- $\mathbf{B} = (\mathbf{b}_1, \dots, \mathbf{b}_m)^T \in \mathbb{R}^{m \times l}$

Vectorization

1. Denote $\hat{\mathbf{Y}} = b \cdot \mathbf{1}_n + \mathbf{X} \cdot \mathbf{w} \in \mathbb{R}^{n \times 1}$

- $\mathbf{1}_n = (1, \dots, 1)^T \in \mathbb{R}^{n \times 1}$: Vector of 1's with length n

2. Cost function

$$\begin{aligned}\mathcal{J}(\boldsymbol{\theta}) &= \frac{1}{2n} \sum_{i=1}^n (y_i - \hat{y}_i)^2 = \frac{1}{2n} (\mathbf{Y} - \hat{\mathbf{Y}})^T \cdot (\mathbf{Y} - \hat{\mathbf{Y}}) \\ &= \frac{1}{2n} (\mathbf{Y} - b \cdot \mathbf{1}_n - \mathbf{X} \cdot \mathbf{w})^T \cdot (\mathbf{Y} - b \cdot \mathbf{1}_n - \mathbf{X} \cdot \mathbf{w}) \\ &= \frac{1}{2n} (\mathbf{Y} - \tilde{\mathbf{X}} \cdot \boldsymbol{\theta})^T \cdot (\mathbf{Y} - \tilde{\mathbf{X}} \cdot \boldsymbol{\theta})\end{aligned}$$

- $\tilde{\mathbf{X}} = (\tilde{\mathbf{x}}_1, \dots, \tilde{\mathbf{x}}_n)^T$ (integrate b and $\mathbf{w}$ together)
- $\tilde{\mathbf{x}}_i = (1, \mathbf{x}_i^T)^T$

Vectorization

1. Consider

$$\frac{\partial \mathcal{J}}{\partial \boldsymbol{\theta}} = 0$$

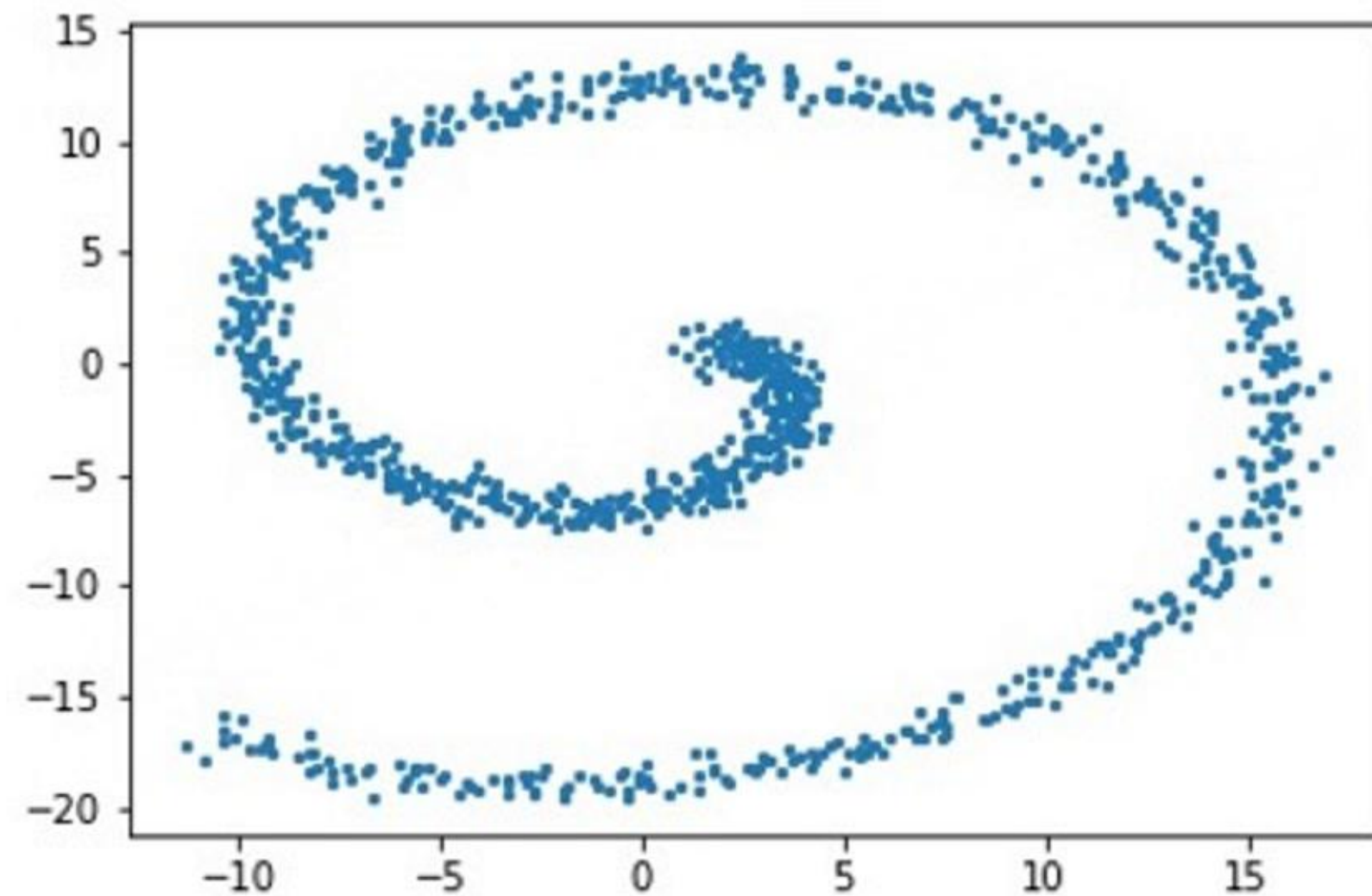
2. Obtain the **normal equation**

$$\tilde{\mathbf{X}}^T \cdot \tilde{\mathbf{X}} \cdot \boldsymbol{\theta} = \tilde{\mathbf{X}}^T \cdot \mathbf{Y}$$

3. Solution (**closed-form or analytical solution**)

$$\hat{\boldsymbol{\theta}} = (\tilde{\mathbf{X}}^T \cdot \tilde{\mathbf{X}})^{-1} \tilde{\mathbf{X}}^T \cdot \mathbf{Y}$$

Example



Logistic regression -- Model setup

1. Training data

- Feature $\mathbf{x}_i \in \mathbb{R}^{d \times 1}$, labels $y_i \in \{0, 1\}$ ($i = 1, \dots, n$)

2. Goal

- Learn the relationship between feature $\mathbf{x}$ and label y

3. True model (used to generate data)

$$E(y_i \mid \mathbf{x}_i) = P(y_i = 1 \mid \mathbf{x}_i) = \sigma(\mathbf{b}_0 + \mathbf{x}_i^T \cdot \mathbf{w}_0) \quad (i = 1, \dots, n)$$

- $\sigma(z) = \{1 + \exp(-z)\}^{-1}$ (why do we need this transformation?)

Logistic regression -- Model setup

1. (Assumed) Proposed model (used to fit data)

- $P(y_i = 1 \mid \mathbf{x}_i) = \sigma(\mathbf{b} + \mathbf{x}_i^T \cdot \mathbf{w}) \quad (i = 1, \dots, n)$
 - ▷ $\boldsymbol{\theta} = (\mathbf{b}, \mathbf{w}^T)^T$: (Proposed) Model parameters
 - ▷ By fitting a model, we mean to estimate parameters of the proposed model
 - ▷ Have the same form with the true model, but we need to estimate parameters

2. **Question:** What if we use a linear regression model to analyze binary classification problems?

Logistic regression -- Parameter estimation

1. **Loss function:** For the i th example,

- Cross entropy (minus log-likelihood)

$$\mathcal{L}(y_i, a_i) = -\{y_i \log a_i + (1 - y_i) \log(1 - a_i)\}$$

▷ $a_i = \sigma(z_i)$ (estimated probability)

▷ $z_i = b + \mathbf{x}_i^T \cdot \mathbf{w}$

- a_i is a function of $\boldsymbol{\theta} = (b, \mathbf{w}^T)^T$
- Cross entropy is also a function of $\boldsymbol{\theta}$
- However, in the above expressions, we omit its argument for simplicity
- The above cross entropy is a convex function of the model parameter $\boldsymbol{\theta}$

Logistic regression -- Parameter estimation

1. Cost function

$$\begin{aligned}\mathcal{J}(\boldsymbol{\theta}) &= n^{-1} \sum_{i=1}^n \mathcal{L}\{y_i, a_i\} \\ &= -n^{-1} \sum_{i=1}^n \{y_i \log a_i + (1 - y_i) \log\{1 - a_i\}\}\end{aligned}$$

- $\boldsymbol{\theta} = (b, \boldsymbol{w}^T)^T$
- $a_i = \sigma(b + \boldsymbol{x}_i^T \cdot \boldsymbol{w})$

Logistic regression -- Parameter estimation

1. Solve

$$\frac{\mathcal{J}}{\partial \boldsymbol{\theta}} = 0$$

2. No analytical solution

Newton-Raphson algorithm

Step 1. Randomly initialize $\boldsymbol{\theta}^{(0)}$

Step 2. Based on $\boldsymbol{\theta}^{(t)}$ obtain

$$\nabla \mathcal{J}(\boldsymbol{\theta}^{(t)}) = \frac{\partial \mathcal{J}}{\partial \boldsymbol{\theta}}(\boldsymbol{\theta}^{(t)}),$$
$$H(\mathcal{J})(\boldsymbol{\theta}^{(t)}) = \frac{\partial^2 \mathcal{J}}{\partial \boldsymbol{\theta} \partial \boldsymbol{\theta}^T}(\boldsymbol{\theta}^{(t)})$$

Step 3. Update parameter

$$\boldsymbol{\theta}^{(t+1)} = \boldsymbol{\theta}^{(t)} - \{H(\mathcal{J})(\boldsymbol{\theta}^{(t)})\}^{-1} \cdot \nabla \mathcal{J}(\boldsymbol{\theta}^{(t)})$$

Step 4. Go back to Step 2 until convergence, or reaching maximum number of iterations

Remarks

1. Cost function $\mathcal{J}$ is univariate
2. In this course, we implicitly assume that the dimension of $\partial\mathcal{J}/\partial\boldsymbol{\theta}$ is the same as $\boldsymbol{\theta}$
3. That is, if $\boldsymbol{\theta}$ is a column vector, so is $\partial\mathcal{J}/\partial\boldsymbol{\theta}$
4. Then, the third step of Newton-Raphson algorithm is mathematically rigorous
5. In other textbooks, $\partial\mathcal{J}/\partial\boldsymbol{\theta}$ is defined as a row vector, even if $\boldsymbol{\theta}$ is a column vector
6. We don't adopt that convention in this course

Discussion

1. Fast convergence by incorporating a Hessian matrix
2. Computation efficiency is sacrificed at the same time
3. Feasible when the parameter dimension is low
4. Not applicable for deep learning models
5. What algorithms should we use for deep learning models?

(Batch) Gradient descent algorithm

Step 1. Randomly initialize $\boldsymbol{\theta}^{(0)}$

Step 2. Based on $\boldsymbol{\theta}^{(t)}$ obtain

$$\nabla \mathcal{J}(\boldsymbol{\theta}^{(t)}) = \frac{\partial \mathcal{J}}{\partial \boldsymbol{\theta}}(\boldsymbol{\theta}^{(t)})$$

Step 3. Update parameter

$$\boldsymbol{\theta}^{(t+1)} = \boldsymbol{\theta}^{(t)} - \alpha \cdot \nabla \mathcal{J}(\boldsymbol{\theta}^{(t)})$$

Step 4. Go back to Step 2 until convergence, or reaching maximum number of iterations

(Batch) Gradient descent algorithm

1. α : Learning rate

- Controls the speed of convergence
- Affects the performance of the model
- To be discussed

Comparison of algorithms

1. Newton-Raphson algorithm

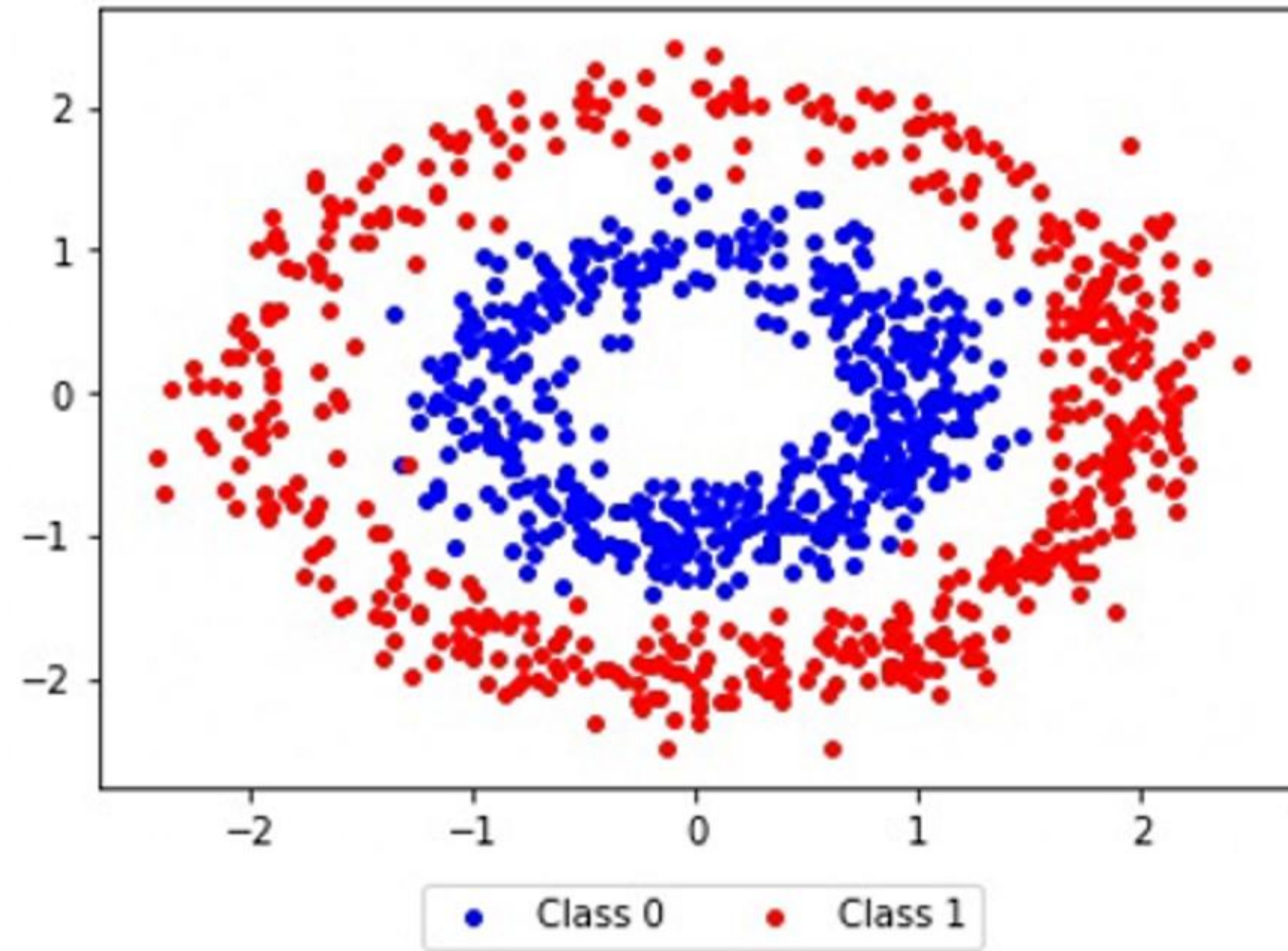
- Advantages:
 - ▷ Fast convergence rate
 - ▷ High accuracy
- Disadvantages:
 - ▷ Low computation efficiency due to the Hessian matrix
 - ▷ Low stability if the Hessian matrix is ill-conditioned

Comparison of algorithms (Cont'd)

1. (Batch) Gradient descent algorithm

- Advantages:
 - ▷ High computation efficiency
 - ▷ Easy to implement
- Disadvantages:
 - ▷ Low convergence rate
 - ▷ Sensitive to learning rate

Example



Summary

1. The following shows typical steps for data modeling

Problem	Proposed model	Cost function	Algorithm
$y \in \mathbb{R}$	$E(y \mathbf{x}) = b + \mathbf{x}^T \cdot \mathbf{w}$	$n^{-1} \sum_{i=1}^n (y_i - b - \mathbf{x}_i^T \cdot \mathbf{w})^2$	Normal equation
$y \in \{0, 1\}$	$E(y \mathbf{x}) = \sigma(b + \mathbf{x}^T \cdot \mathbf{w})$	$-n^{-1} \sum_{i=1}^n \{y_i \log a_i + (1 - y_i) \log(1 - a_i)\}$	Gradient descent

2. Gradient descent algorithms are commonly used for AI models

3. The essential step for gradient descent algorithms is to obtain

$$\frac{\partial \mathcal{J}}{\partial \boldsymbol{\theta}}$$

Problems

1. Up to now, we have assumed the same form for the true model and the proposed model
2. However, it is commonly not the case in practice
3. Model **misspecification**
 - The proposed (statistical) models are **too simple**
 - True model, however, is **COMPLEX**

Learning goals

1. Fully connected neural network (**multiple layer perceptron**)
2. Convolutional Neural Network (**CNN**, ...)
3. Sequential modeling (**RNN**, **LSTM**, **transformers**, ...)
4. (If we have time,) More topics (**GNN**, **GCN**, ...)

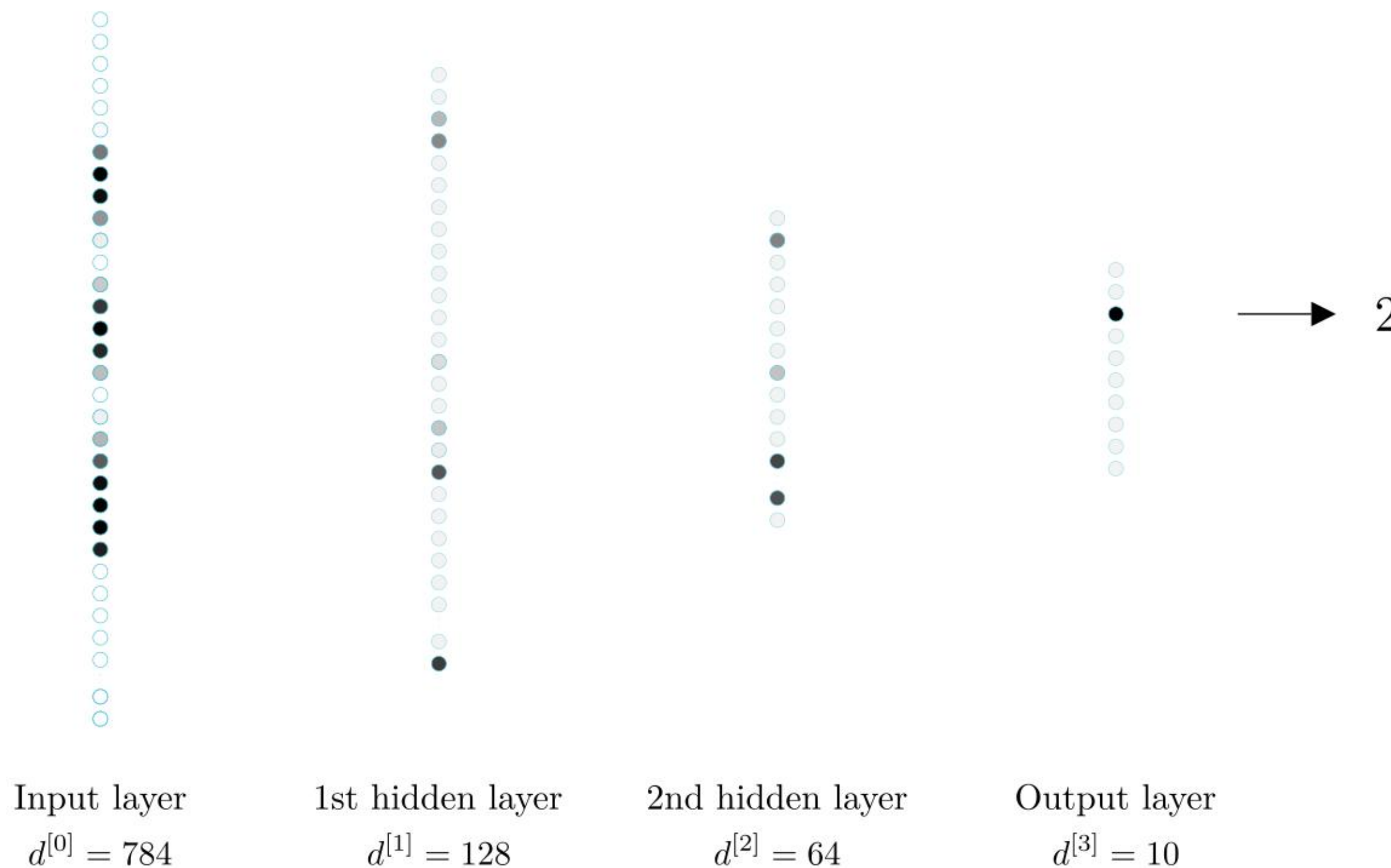
FNN Example

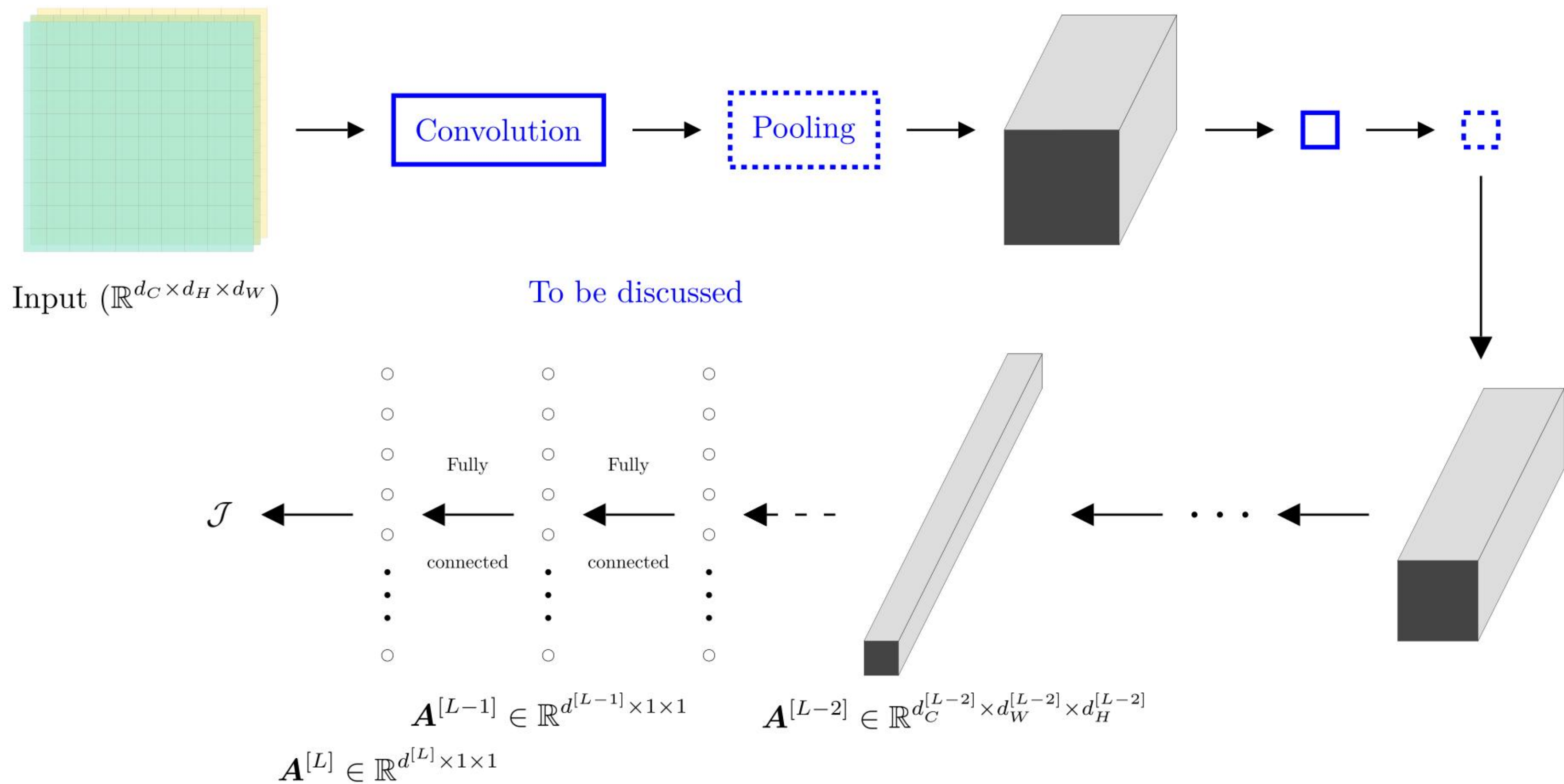
Test image



Model (FNN with 2 hidden layers)

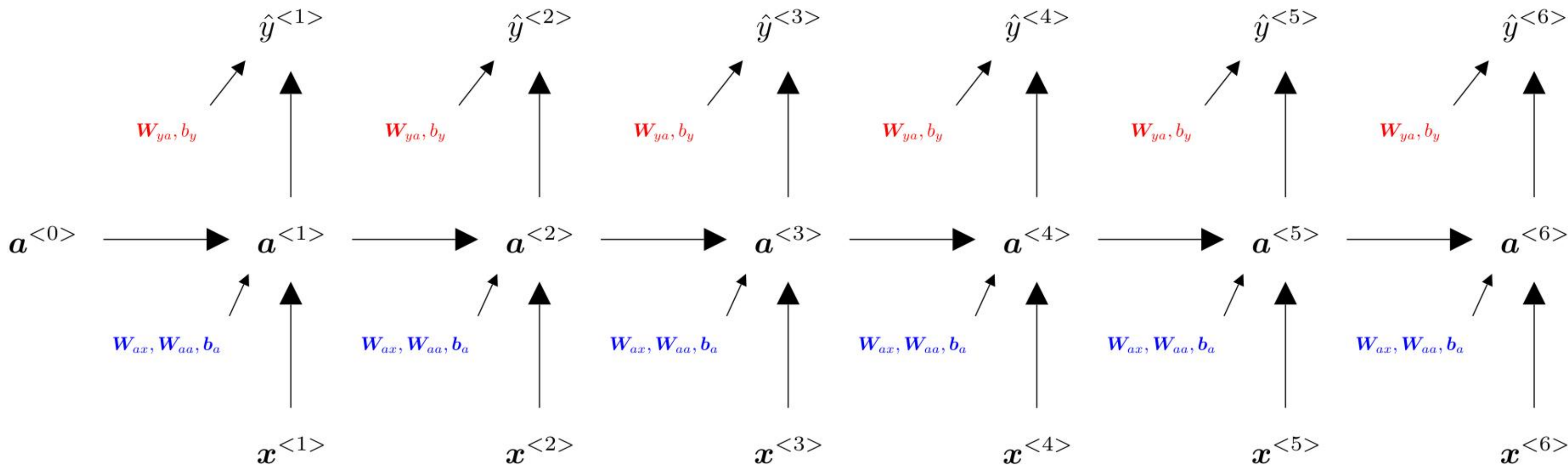
Estimated result





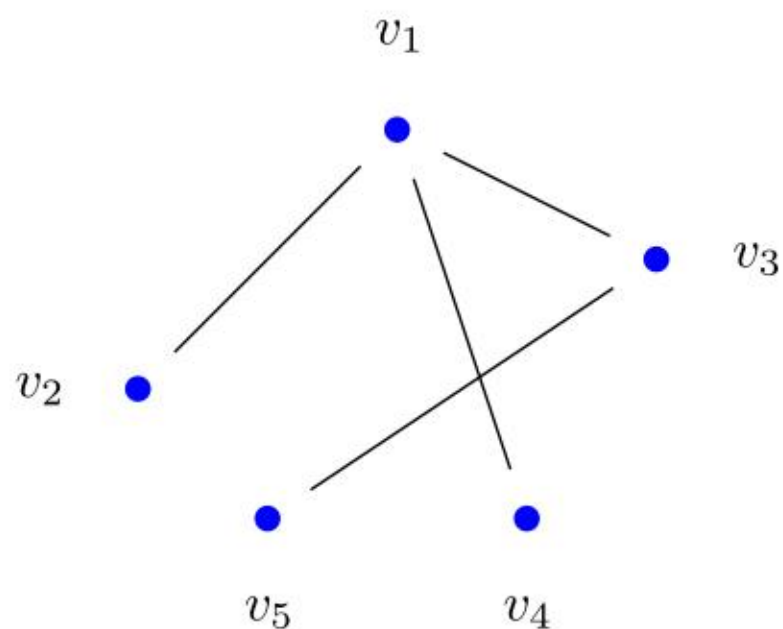
Building block for RNN

$$\hat{y}^{<i>} = \sigma_{ya}(\mathbf{W}_{ya} \cdot \mathbf{a}^{<i>} + b_y) \quad (i = 1, \dots, 6)$$



$$\mathbf{a}^{<i>} = \sigma_{ax}(\mathbf{W}_{ax} \cdot \mathbf{x}^{<i>} + \mathbf{W}_{aa} \cdot \mathbf{a}^{<i-1>} + b_a) \quad (i = 1, \dots, 6)$$

Undirected graph



Vertex set: $\mathcal{V} = \{v_i : i = 1, \dots, n\}$

Edge set: $\mathcal{E}$

Undirected graph: $\mathcal{G} = \mathcal{V} \cup \mathcal{E}$

Adjacent matrix (binary or weighted)

$$A = \begin{pmatrix} 0 & 1 & 1 & 1 & 0 \\ 1 & 0 & 0 & 0 & 0 \\ 1 & 0 & 0 & 0 & 1 \\ 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \end{pmatrix}$$

Degree matrix (diagonal, summation of each row in A)

$$D = \begin{pmatrix} 3 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 2 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \end{pmatrix}$$

Review

1. $X \sim P$: Random variable X is generated from a distribution P
2. $E_{X \sim P}\{f(X)\}$: Expectation of $f(X)$
 - We use “ $X \sim P$ ” to highlight the true distribution of X
 - $f(x)$: Measurable function
3. In practice, we generally cannot observe the distribution P
4. Instead, we may only observe a random sample $\{x_i : i = 1, \dots, n\}$ generated from P

Review

1. Based on $\{x_i : i = 1, \dots, n\}$, its empirical distribution P_n is

Sample	x_1	x_2	x_3	$\dots$	x_n
Probability	$\frac{1}{n}$	$\frac{1}{n}$	$\frac{1}{n}$	$\dots$	$\frac{1}{n}$

2. Remarks

- Since $\{x_i : i = 1, \dots, n\}$ is random, **the empirical distribution is also random**
- The density of the empirical distribution (with respect to a counting measure) is

$$p_n(x) = n^{-1} \sum_{i=1}^n \delta(x = x_i)$$

▷ $\delta(x = y)$: indicator function, and it is 1 if $x = y$, and it is 0 otherwise

- It can be easily shown that the sample mean is $E_{X \sim P_n}(X) = n^{-1} \sum_{i=1}^n x_i$

Cross entropy

1. A fact: under regularity conditions,

$$E_{X \sim P} \{\log p(X)\} \geq E_{X \sim P} \{\log q(X)\}$$

- $p(x)$: Density function of the distribution P
- $q(x)$: Another density function of a certain distribution

2. (Information) Entropy comes from the information theory

- “Information content” of an event: Negative logarithm of its probability of occurrence
- “(Information) Entropy”: Average of information content

3. Cross entropy for two density functions p and q

$$H(p, q) = -E_{X \sim P} \{\log q(X)\} = -E \{\log q(X)\}$$

- Measures how well a density q approximates the distribution of X .

Cross entropy

1. Only observe a random sample $\{x_1, \dots, x_n\}$
2. Under generality conditions, cross entropy $H(p, q)$ can be approximated by (law of large numbers)

$$\hat{H}(p, q) = -\frac{1}{n} \sum_{i=1}^n \log q(x_i)$$

- Equivalently, it is the cross entropy between the empirical density

$$P_n(x) = n^{-1} \sum_{i=1}^n \delta(x = x_i) \text{ and } q(x)$$

3. Goal: Find a density function $q(x)$ to minimize $\hat{H}(p, q)$
 - That is, to approximate the empirical distribution of the random sample well

Revisit logistic regression

1. Assume $\{(\mathbf{x}_i, y_i) : i = 1, \dots, n\}$ is i.i.d. with a density function

$$p(\mathbf{x}, y) = p(\mathbf{x})p(y \mid \mathbf{x})$$

- $p(\mathbf{x})$: Unspecified marginal density for $\mathbf{x}$
- $p(y \mid \mathbf{x}) = \sigma(\mathbf{b}_0 + \mathbf{x}^\top \mathbf{w}_0)$: Parametric conditional density for y

2. Goal: Find density $q(\mathbf{x}, y)$ to minimize

$$\hat{H}(p, q) = -\frac{1}{n} \sum_{i=1}^n \log q(\mathbf{x}_i, y_i)$$

Revisit logistic regression

1. A little more math

$$\begin{aligned} -\frac{1}{n} \sum_{i=1}^n \log q(\mathbf{x}_i, y_i) &= -\frac{1}{n} \sum_{i=1}^n \{\log q(\mathbf{x}_i) + \log q(y_i \mid \mathbf{x}_i)\} \\ &= -\frac{1}{n} \sum_{i=1}^n \log q(\mathbf{x}_i) - \frac{1}{n} \sum_{i=1}^n [y_i \log a_i(\boldsymbol{\theta}) + (1 - y_i) \log \{1 - a_i(\boldsymbol{\theta})\}] \end{aligned}$$

- $q(\mathbf{x}, y) = q(\mathbf{x})q(y \mid \mathbf{x})$
- $q(\mathbf{x})$: Unspecified
- $q(y \mid \mathbf{x}) = a(\boldsymbol{\theta})^y \{1 - a(\boldsymbol{\theta})\}^{1-y}$
- $a(\boldsymbol{\theta}) = \sigma(b + \mathbf{x}^T \cdot \mathbf{w})$ with parameter $\boldsymbol{\theta} = (b, \mathbf{w}^T)^T$

Revisit logistic regression

1. Thus,

$$-\frac{1}{n} \sum_{i=1}^n \log q(\mathbf{x}_i, y_i) = -\frac{1}{n} \sum_{i=1}^n \log q(\mathbf{x}_i) - \frac{1}{n} \sum_{i=1}^n [y_i \log a_i(\boldsymbol{\theta}) + (1 - y_i) \log \{1 - a_i(\boldsymbol{\theta})\}]$$

- The **first** part is of no interest
- We focus on estimating parameters in the **second** part

2. Thus, it is equivalent to minimizing

$$-\frac{1}{n} \sum_{i=1}^n [y_i \log a_i(\boldsymbol{\theta}) + (1 - y_i) \log \{1 - a_i(\boldsymbol{\theta})\}]$$

- The cost function for logistic regression

Summary

1. Reviewed linear and logistic regressions
 - Goal: A unified procedure to model data
 - ▷ Data $\rightarrow$ (Proposed) Model $\rightarrow$ Cost (Loss) function $\rightarrow$ Optimization
2. Vectorization to improve computation efficiency
3. Two optimization algorithms: Newton-Raphson vs gradient descent
4. Drawbacks of simple models: Model misspecification
5. Cross entropy for classification problems