



AI and Deep Learning

2-7 Softmax Regression

Zhonglei Wang

WISE, SOE, CHOW

XMU

(Version v1)

Contents

1. Model

2. Forward propagation

3. Backpropagation

4. Example

Introduction

1. Binary classification problems

- $y \in \{0, 1\}$
- Model is built for $P(y = 1 \mid \mathbf{x})$
- Loss function: Cross entropy

2. In practice, we may have $K(> 2)$ classes: $y \in \{0, 1, \dots, K - 1\}$

- Consider one-hot representation
- $\mathbf{y} = (0, \dots, 1, \dots, 0)^T \in \mathbb{R}^{K \times 1}$
- If $y = k$, only the $(k + 1)$ th element of $\mathbf{y}$ is 1

Model

1. For a feature $\mathbf{x}$,
 - We consider a neural network with K outputs
 - Each output is a score (conditional probability) for the corresponding class
2. That is, the only difference is the number of neurons in the last layer
 - For **binary** classification problems, we only have **one neuron** for the output layer
 - For **general** classification problems, we have $d^{[L]} = K$ **neurons** for the output layer

Softmax function

1. For a d -dimensional vector $\mathbf{z} = (z_1, \dots, z_d)^T \in \mathbb{R}^{d \times 1}$, the softmax function is

$$\text{Softmax}(\mathbf{z}) = \begin{pmatrix} \exp(z_1) / \sum_{i=1}^d \exp(z_i) \\ \exp(z_2) / \sum_{i=1}^d \exp(z_i) \\ \vdots \\ \exp(z_d) / \sum_{i=1}^d \exp(z_i) \end{pmatrix}$$

2. Remarks

- Amplify gaps among $\{z_1, \dots, z_d\}$ due to the exponential function
- Location invariant: $\text{Softmax}(\mathbf{z} + c) = \text{Softmax}(\mathbf{z})$ for all $c \in \mathbb{R}$
- Only consider relative order

Softmax function

1. Denote $\mathbf{a} = \text{Softmax}(\mathbf{z}) \in \mathbb{R}^{d \times 1}$

- $\mathbf{a} = (a_1, \dots, a_d)^T$

2. Then, we have (check by yourself)

$$\frac{\partial \mathbf{a}^T}{\partial \mathbf{z}} = \left(\frac{\partial a_m}{\partial z_k} \right) = (a_m \cdot (\delta_{mk} - a_k))$$

- $k = 1, \dots, d$: Row index
- $m = 1, \dots, d$: Column index
- δ_{mk} : Indicator function

$$\delta_{mk} = \begin{cases} 1, & m = k \\ 0, & m \neq k \end{cases}$$

Notations

1. Recall that

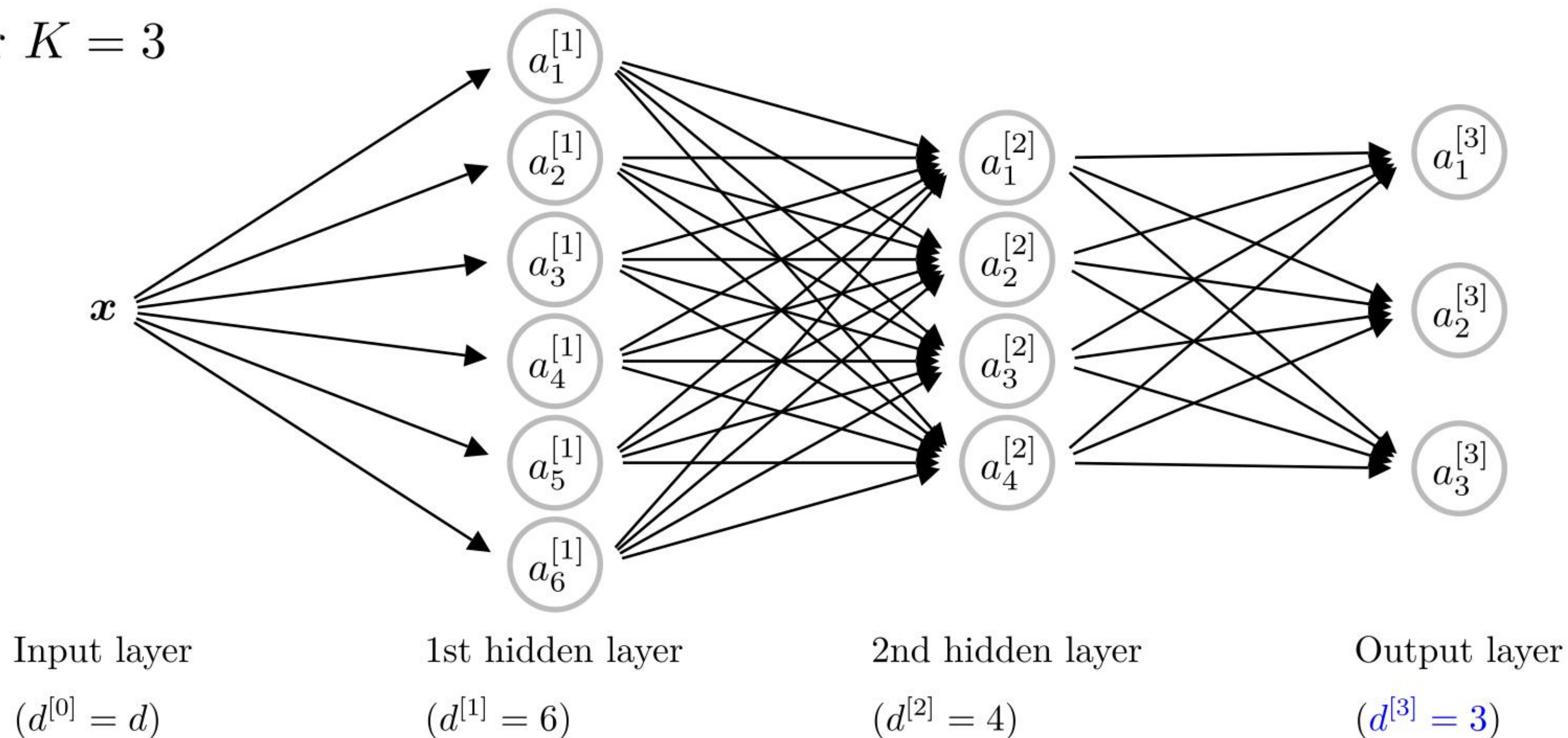
- L : Number of layers in the neural network
- $d^{[l]}$: Number of neurons in the l th layer ($l = 0, \dots, L$)
- $\mathbf{a}^{[l]} = (a_1^{[l]}, \dots, a_{d^{[l]}}^{[l]})^T \in \mathbb{R}^{d^{[l]} \times 1}$
- $\mathbf{W}^{[l]} = (\mathbf{w}_1^{[l]}, \dots, \mathbf{w}_{d^{[l]}}^{[l]})^T \in \mathbb{R}^{d^{[l]} \times d^{[l-1]}}$
- $\mathbf{b}^{[l]} = (b_1^{[l]}, \dots, b_{d^{[l]}}^{[l]})^T \in \mathbb{R}^{d^{[l]} \times 1}$

2. For **binary classification** problems, we have $d^{[L]} = 1$

3. For **general softmax regression** problems, we have $d^{[L]} = K$

Notations

1. For $K = 3$



Forward propagation

1. Let $\mathbf{A}^{[0]} = \mathbf{X} \in \mathbb{R}^{n \times d}$

2. For $l = 1, \dots, L - 1$,

$$\mathbf{Z}^{[l]} = (\mathbf{b}^{[l]})^T + \mathbf{A}^{[l-1]} \cdot (\mathbf{W}^{[l]})^T,$$
$$\mathbf{A}^{[l]} = \sigma^{[l]}(\mathbf{Z}^{[l]})$$

- $\sigma^{[l]}(z)$: Activation function for the l th layer
- Broadcasting is used for activation functions

Forward propagation

1. For the last layer,

$$\mathbf{W}^{[L]} \in \mathbb{R}^{K \times d_{L-1}}, \quad \mathbf{b} \in \mathbb{R}^{K \times 1}$$

2. Linear transformation is the same:

$$\mathbf{Z}^{[L]} = (\mathbf{b}^{[L]})^T + \mathbf{A}^{[L-1]} \cdot (\mathbf{W}^{[L]})^T$$

- $\mathbf{Z}^{[L]} \in \mathbb{R}^{n \times K}$

3. Softmax function is applied for activation:

$$\mathbf{A}^{[L]} = \text{Softmax}(\mathbf{Z}^{[L]})$$

- $\mathbf{A}^{[L]} \in \mathbb{R}^{n \times K}$

- Softmax function is applied to each row of $\mathbf{Z}^{[L]}$

- The sum of each row of $\mathbf{A}^{[L]}$ is 1

Forward propagation

1. The dimension of the $\mathbf{A}^{[L]}$, containing estimated probabilities in the last layer

- $\mathbf{A}^{[L]} \in \mathbb{R}^{n \times \mathbf{1}}$ for **binary classification** problems
- $\mathbf{A}^{[L]} \in \mathbb{R}^{n \times K}$ for **softmax regression** problems

2. The cost function for softmax regression is

$$\mathcal{J}(\boldsymbol{\theta}) = -n^{-1} \sum_{i=1}^n \sum_{k=1}^K y_{ik} \log a_{ik}^{[L]}$$

- $\boldsymbol{\theta}$: Model parameters
- $\mathbf{y}_i = (y_{i1}, \dots, y_{iK})^T$: One-hot representation for the i th example
- $a_{ik}^{[L]}$: Estimated probability for the $(k - 1)$ th class of the i th example

Backpropagation

1. We only need to focus on the backpropagation for the output layer
2. $\mathbf{dA}^{[L]}$ can be obtained from the cost function
3. The cost function for softmax regression is

$$\mathcal{J}(\boldsymbol{\theta}) = -n^{-1} \sum_{i=1}^n \sum_{k=1}^K y_{ik} \log a_{ik}^{[L]}$$

4. Thus, the (i, k) th component of $\mathbf{dA}^{[L]}$ is

$$\mathrm{d}a_{ik}^{[L]} = \frac{\partial \mathcal{J}}{\partial a_{ik}^{[L]}} = -n^{-1} \frac{y_{ik}}{a_{ik}^{[L]}}$$

5. $\mathbf{dA}^{[L]} = -n^{-1} \mathbf{Y} / \mathbf{A}^{[L]}$ (element-wise division by Python broadcasting)

Backpropagation

1. Next, we focus on obtaining $d\mathbf{Z}^{[L]}$
2. For $i = 1, \dots, n$ and $k = 1, \dots, K$, $z_{ik}^{[L]}$ is only used to get the i th row of $A^{[L]}$
3. For $h \neq i$, we have

$$\frac{\partial a_{hm}^{[L]}}{\partial z_{ik}^{[L]}} = 0 \quad (m = 1, \dots, K)$$

4. For $i = 1, \dots, n; k = 1, \dots, K$, we have

$$dz_{ik}^{[L]} = \frac{\partial \mathcal{J}}{\partial z_{ik}^{[L]}} = \sum_{h=1}^n \sum_{m=1}^K \frac{\partial \mathcal{J}}{\partial a_{hm}^{[L]}} \frac{\partial a_{hm}^{[L]}}{\partial z_{ik}^{[L]}} = \sum_{m=1}^K \frac{\partial \mathcal{J}}{\partial a_{im}^{[L]}} \frac{\partial a_{im}^{[L]}}{\partial z_{ik}^{[L]}}$$

5. To obtain $d\mathbf{Z}^{[L]}$, it is enough to get its elements $\{dz_{ik}^{[L]} : i = 1, \dots, n; k = 1, \dots, K\}$

Backpropagation

1. We have obtained

$$\frac{\partial \mathcal{J}}{\partial a_{im}^{[L]}} = -n^{-1} \frac{y_{im}}{a_{im}^{[L]}}, \quad \frac{\partial a_{im}^{[L]}}{\partial z_{ik}^{[L]}} = a_{im}^{[L]} \cdot (\delta_{mk} - a_{ik}^{[L]})$$

2. Thus, we have

$$\begin{aligned} dz_{ik}^{[L]} &= \sum_{m=1}^K \frac{\partial \mathcal{J}}{\partial a_{im}^{[L]}} \frac{\partial a_{im}^{[L]}}{\partial z_{ik}^{[L]}} = -n^{-1} \sum_{m=1}^K \frac{y_{im}}{a_{im}^{[L]}} \cdot a_{im}^{[L]} \cdot (\delta_{mk} - a_{ik}^{[L]}) \\ &= -n^{-1} \sum_{m=1}^K y_{im} \cdot \delta_{mk} + n^{-1} a_{ik}^{[L]} \cdot \sum_{m=1}^K y_{im} = -n^{-1} (y_{ik} - a_{ik}^{[L]}) \end{aligned}$$

- The last equality holds since $\sum_{m=1}^K y_{im} = 1$

3. Thus, we have $d\mathbf{Z}^{[L]} = -n^{-1}(\mathbf{Y} - \mathbf{A}^{[L]})$ (error matrix from the output layer)

Backpropagation

1. To sum up, for the output layer, we have

$$d\mathbf{A}^{[L]} = -n^{-1}\mathbf{Y}/\mathbf{A}^{[L]} \quad (\text{element-wise division}),$$

$$d\mathbf{Z}^{[L]} = -n^{-1}(\mathbf{Y} - \mathbf{A}^{[L]}) \quad (\text{exactly the same as binary models}),$$

$$d\mathbf{W}^{[L]} = (d\mathbf{Z}^{[L]})^T \cdot \mathbf{A}^{[L-1]},$$

$$d\mathbf{b}^{[l]} = (d\mathbf{Z}^{[l]})^T \cdot \mathbf{1},$$

$$d\mathbf{A}^{[L-1]} = d\mathbf{Z}^{[L]} \cdot \mathbf{W}^{[L]}$$

- $\mathbf{1} \in \mathbb{R}^{n \times 1}$: Column vector of 1's

Backpropagation

1. The hidden layers for softmax regression is exactly the same as binary regression
2. We can directly use the backpropagation result for binary regression models
3. That is, by assuming $\mathbf{dA}^{[l]}$ is available for $l = L - 1, \dots, 1$, we have

$$\mathbf{dZ}^{[l]} = \mathbf{dA}^{[l]} \circ \sigma^{[l]'}(\mathbf{Z}^{[l]}),$$

$$\mathbf{dW}^{[l]} = (\mathbf{dZ}^{[l]})^T \cdot \mathbf{A}^{[l-1]},$$

$$\mathbf{db}^{[l]} = (\mathbf{dZ}^{[l]})^T \cdot \mathbf{1},$$

$$\mathbf{dA}^{[l-1]} = \mathbf{dZ}^{[l]} \cdot \mathbf{W}^{[l]}$$

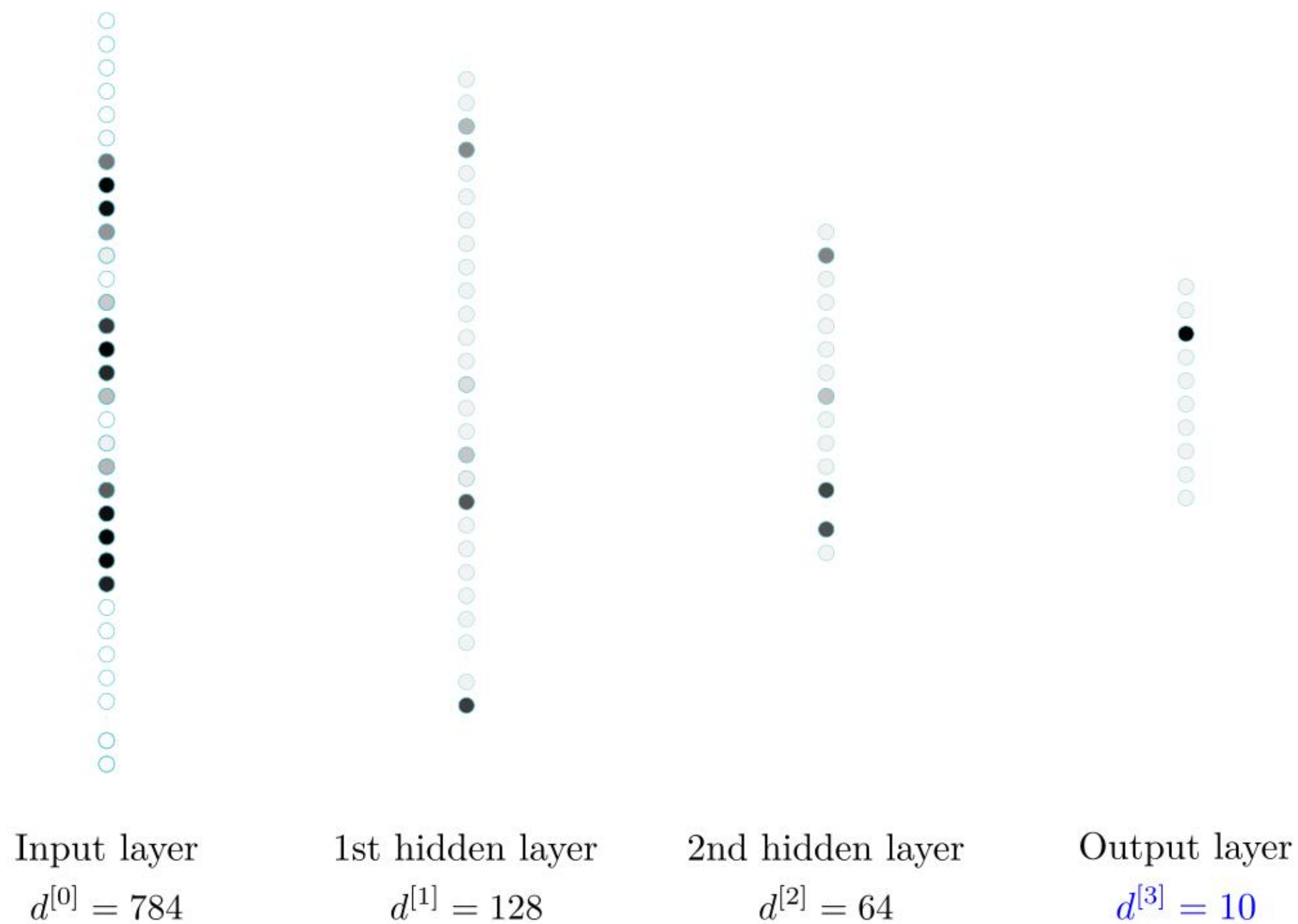
- $\mathbf{1} \in \mathbb{R}^{n \times 1}$: Column vector of 1's

Example

Test image



Model (FNN with 2 hidden layers)



Summary

1. Four aspects for softmax regression

- Model structure: K outputs, estimating a conditional probability for each class
- Forward propagation: Hidden layers are the same, but apply softmax to the output layer
- Cost function: Cross entropy
- Backpropagation: Hidden layers are the same, and scaled estimated error for the output layer, $d\mathbf{Z}^{[L]} = n^{-1}(\mathbf{A}^{[L]} - \mathbf{Y})$