



人工智能与深度学习

2-2 只有一个隐藏层的神经网络 II

(多个训练样本)

王中雷

经济学院、王亚南经济研究院、邹至庄经济研究院

厦门大学

(第一版)

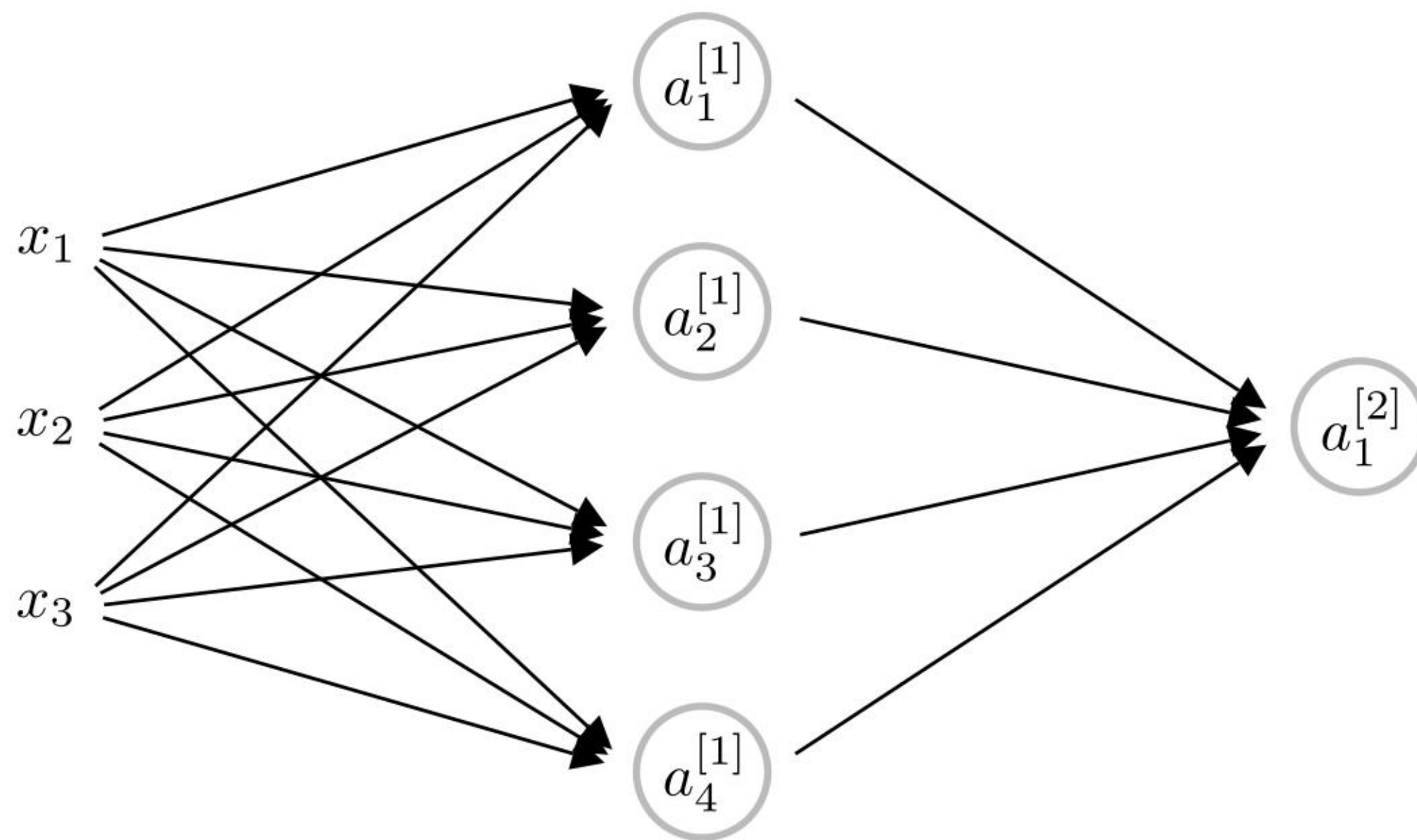
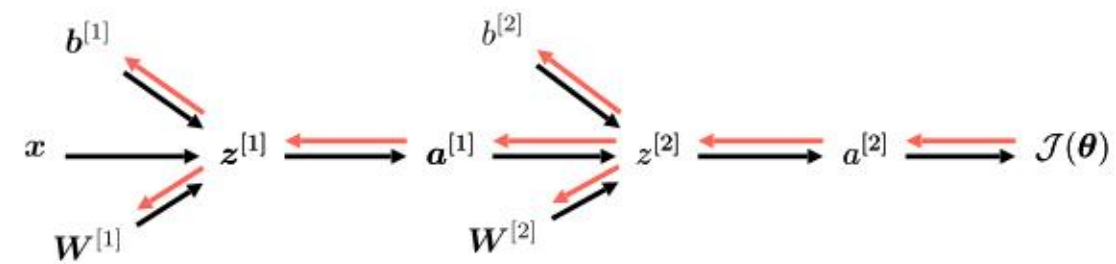
内容摘要

1. 前向传播

2. 后向传播

3. 批量梯度下降法

回顾



输入层 ($d^{[0]} = 3$)

第一隐藏层 ($d^{[1]} = 4$)

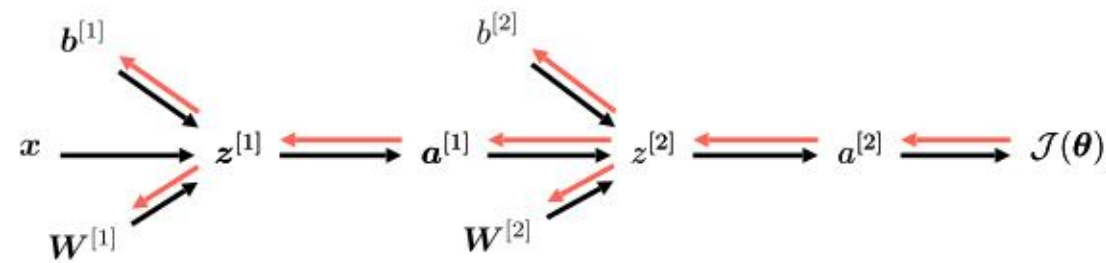
输出层 ($d^{[2]} = 1$)

回顾

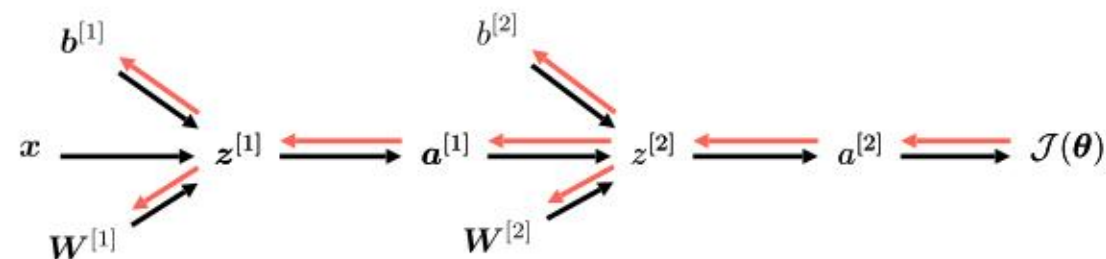
1. 回顾

- L : 神经网络模型的层数
- $d^{[l]}$: 第 l 层神经元的个数 ($l = 0, \dots, L$)
- $\mathbf{a}^{[l]} = (a_1^{[l]}, \dots, a_{d^{[l]}}^{[l]})^T \in \mathbb{R}^{d^{[l]} \times 1}$
- $\mathbf{W}^{[l]} = (\mathbf{w}_1^{[l]}, \dots, \mathbf{w}_{d^{[l]}}^{[l]})^T \in \mathbb{R}^{d^{[l]} \times d^{[l-1]}}$
- $\mathbf{b}^{[l]} = (b_1^{[l]}, \dots, b_{d^{[l]}}^{[l]})^T \in \mathbb{R}^{d^{[l]} \times 1}$

2. 我们考虑一个 0-1 二分类问题



回顾



1. 如果我们**只有一个**训练样本 (x, y) , 前向传播的计算如下

$$z^{[1]} = b^{[1]} + W^{[1]} \cdot x,$$

$$a^{[1]} = \sigma(z^{[1]}),$$

$$z^{[2]} = b^{[2]} + W^{[2]} \cdot a^{[1]},$$

$$a^{[2]} = \sigma(z^{[2]})$$

- Python 中的**广播机制**被用于计算 $z^{[1]}$ 和 $a^{[1]}$

2. 以上出现的是**列向量**

3. 代价函数为

$$\mathcal{J} = \mathcal{L} = - \{ y \log a^{[2]} + (1 - y) \log (1 - a^{[2]}) \}$$

前向传播

1. 假设我们有 n 个训练样本 $\{(\mathbf{x}_i, y_i) : i = 1, \dots, n\}$

2. 记 $\mathbf{X} = (\mathbf{x}_1, \dots, \mathbf{x}_n)^T \in \mathbb{R}^{n \times d}$

3. 基于向量化的前向传播，我们有

$$\mathbf{Z}^{[1]} = (\mathbf{b}^{[1]})^T + \mathbf{X} \cdot (\mathbf{W}^{[1]})^T \in \mathbb{R}^{n \times d^{[1]}},$$

$$\mathbf{A}^{[1]} = \sigma(\mathbf{Z}^{[1]}) \in \mathbb{R}^{n \times d^{[1]}},$$

$$\mathbf{Z}^{[2]} = \mathbf{b}^{[2]} + \mathbf{A}^{[1]} \cdot (\mathbf{W}^{[2]})^T \in \mathbb{R}^{n \times d^{[2]}},$$

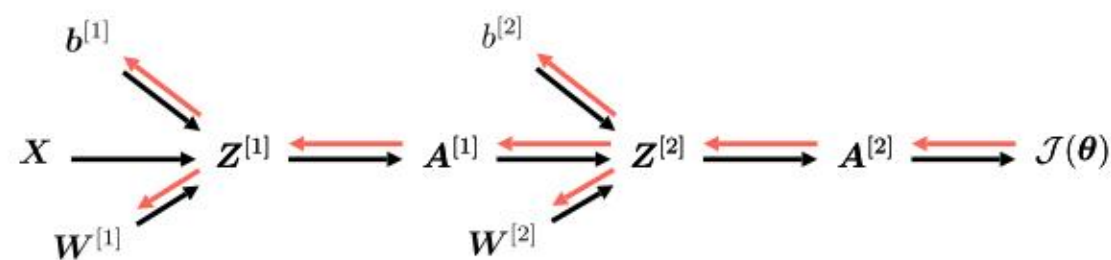
$$\mathbf{A}^{[2]} = \sigma(\mathbf{Z}^{[2]}) \in \mathbb{R}^{n \times d^{[2]}}$$

- Python 中的广播机制被用于以上计算

4. 实际上，我们将信息进行行拼接

$$\begin{aligned} \mathbf{Z}^{[1]} &= (\mathbf{b}^{[1]})^T + \begin{pmatrix} \mathbf{x}_1^T \\ \mathbf{x}_2^T \\ \vdots \\ \mathbf{x}_n^T \end{pmatrix} \cdot (\mathbf{W}^{[1]})^T \\ &= \begin{pmatrix} (\mathbf{b}^{[1]})^T + \mathbf{x}_1^T \cdot (\mathbf{W}^{[1]})^T \\ (\mathbf{b}^{[1]})^T + \mathbf{x}_2^T \cdot (\mathbf{W}^{[1]})^T \\ \vdots \\ (\mathbf{b}^{[1]})^T + \mathbf{x}_n^T \cdot (\mathbf{W}^{[1]})^T \end{pmatrix} \end{aligned}$$

前向传播



1. 假设我们有 n 个训练样本 $\{(\mathbf{x}_i, y_i) : i = 1, \dots, n\}$

2. 记 $\mathbf{X} = (\mathbf{x}_1, \dots, \mathbf{x}_n)^T \in \mathbb{R}^{n \times d}$

3. 基于向量化的前向传播，我们有

$$\mathbf{Z}^{[1]} = (\mathbf{b}^{[1]})^T + \mathbf{X} \cdot (\mathbf{W}^{[1]})^T \in \mathbb{R}^{n \times d^{[1]}},$$

$$\mathbf{A}^{[1]} = \sigma(\mathbf{Z}^{[1]}) \in \mathbb{R}^{n \times d^{[1]}},$$

$$\mathbf{Z}^{[2]} = \mathbf{b}^{[2]} + \mathbf{A}^{[1]} \cdot (\mathbf{W}^{[2]})^T \in \mathbb{R}^{n \times d^{[2]}},$$

$$\mathbf{A}^{[2]} = \sigma(\mathbf{Z}^{[2]}) \in \mathbb{R}^{n \times d^{[2]}}$$

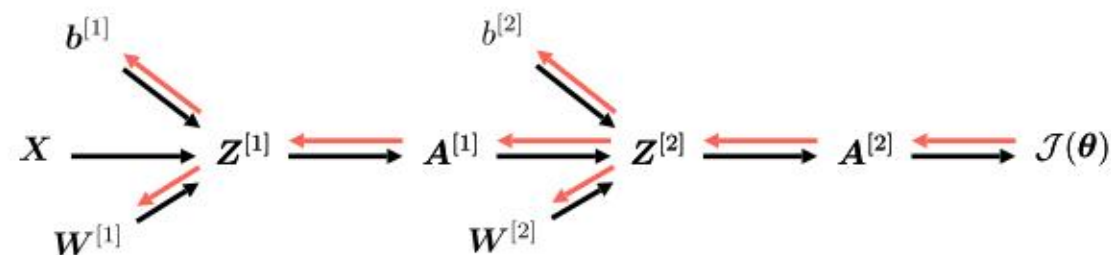
$$\mathbf{A}^{[1]} = \sigma(\mathbf{Z}^{[1]}) = (\sigma(Z_{ij}^{[1]}))$$

$$(\mathbf{Z}^{[1]} = (Z_{ij}^{[1]}); i = 1, \dots, n; j = 1, \dots, d^{[1]})$$

- Python 中的广播机制被用于以上计算

4. 实际上，我们将信息进行行拼接

前向传播

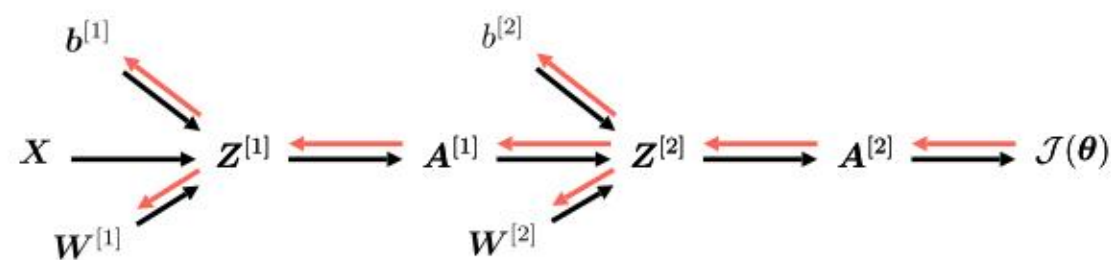


1. 广播机制能够使 Python 代码更加简洁、高效，并能够节省内存

- 代码简洁并易于阅读
- 避免了临时变量的存储，内存优势明显
- 通过 C/Fortran 语言的优化，大规模提升计算性能



前向传播



1. 对于右下角的例子，我们有 $d^{[0]} = 3, d^{[1]} = 4, d^{[2]} = 1$

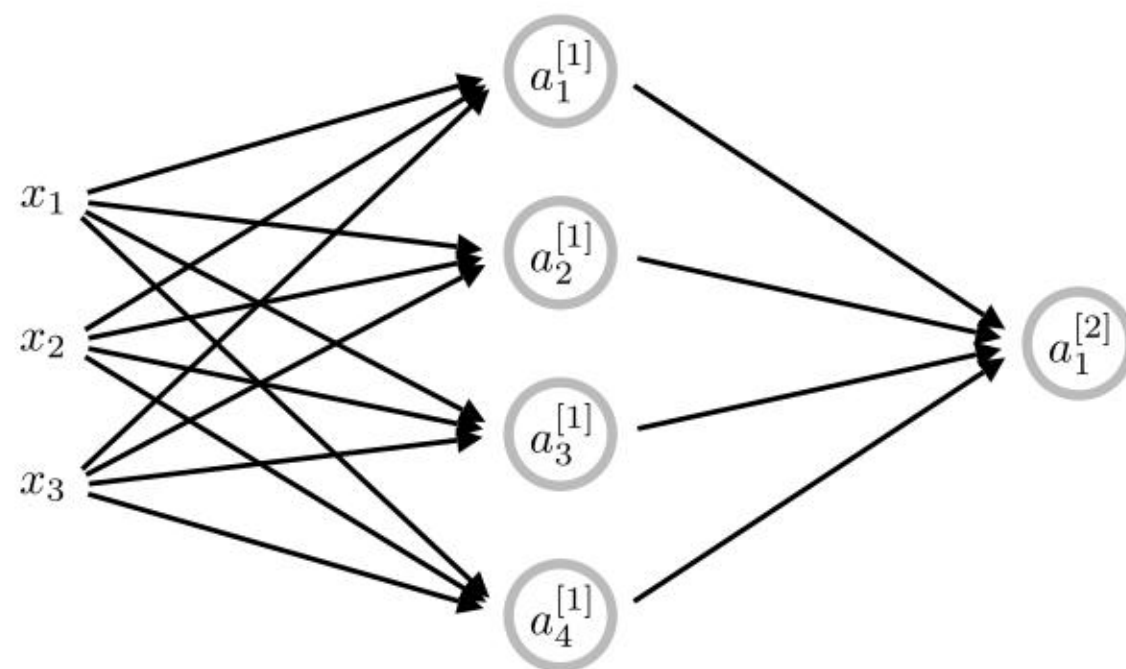
2. 因此，我们有

$$\mathbf{Z}^{[1]} = (\mathbf{b}^{[1]})^T + \mathbf{X} \cdot (\mathbf{W}^{[1]})^T \in \mathbb{R}^{n \times 4},$$

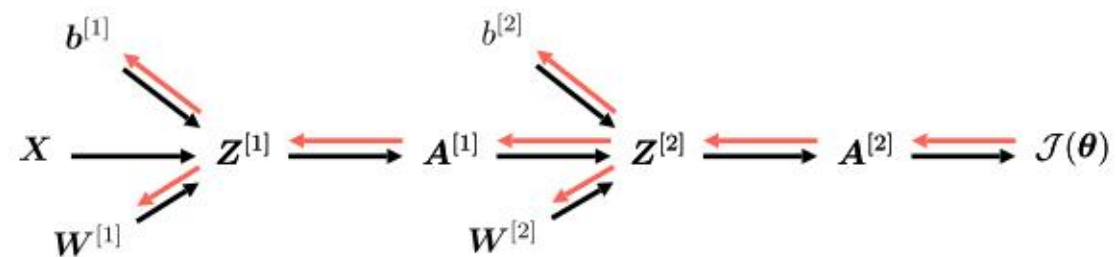
$$\mathbf{A}^{[1]} = \sigma(\mathbf{Z}^{[1]}) \in \mathbb{R}^{n \times 4},$$

$$\mathbf{Z}^{[2]} = \mathbf{b}^{[2]} + \mathbf{A}^{[1]} \cdot (\mathbf{W}^{[2]})^T \in \mathbb{R}^{n \times 1},$$

$$\mathbf{A}^{[2]} = \sigma(\mathbf{Z}^{[2]}) \in \mathbb{R}^{n \times 1}$$



前向传播



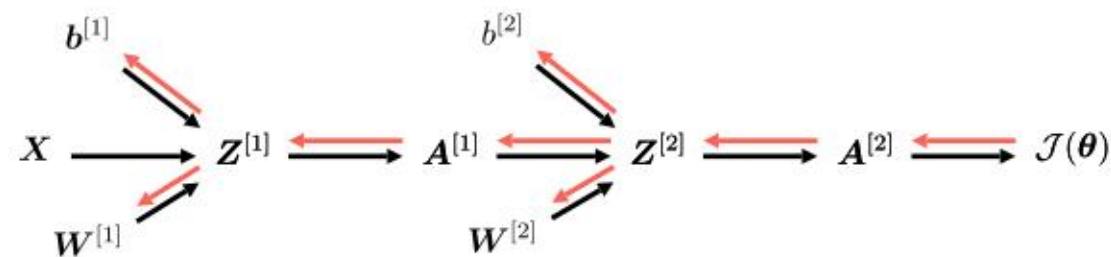
1. 二分类问题的交叉熵代价函数为

$$\mathcal{J} = -\frac{1}{n} \sum_{i=1}^n \left\{ y_i \log a_i^{[2]} + (1 - y_i) \log (1 - a_i^{[2]}) \right\}$$

- $a_i^{[2]}$: 对 $P(y_i = 1 \mid \mathbf{x}_i)$ 的估计值

2. 损失函数的形式由具体问题决定

后向传播



1. 记 $d\cdot = \frac{\partial \mathcal{J}}{\partial \cdot}$

2. 回顾: 对于一个样本点 $(\mathbf{x}, y)$, 我们有

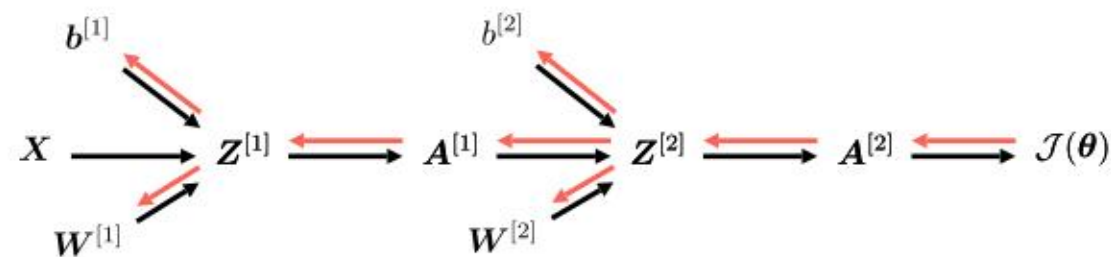
$$\begin{aligned} dz^{[2]} &= a^{[2]} - y, \quad db^{[2]} = dz^{[2]}, \quad d\mathbf{W}^{[2]} = dz^{[2]} \cdot (\mathbf{a}^{[1]})^T, \\ d\mathbf{z}^{[1]} &= dz^{[2]} \cdot \mathbf{D} \cdot (\mathbf{W}^{[2]})^T, \quad d\mathbf{b}^{[1]} = d\mathbf{z}^{[1]}, \quad d\mathbf{W}^{[1]} = d\mathbf{z}^{[1]} \cdot \mathbf{x}^T \end{aligned}$$

- $\mathbf{D} = \text{diag}(\{a_j^{[1]}(1 - a_j^{[1]}) : j = 1, \dots, d^{[1]}\})$

3. 我们有

- $dz^{[2]}$ 可通过 $da^{[2]}$ 得到, 而后者通过代价函数得到
- $d\mathbf{z}^{[1]}$ 可通过 $d\mathbf{a}^{[1]}$ 得到, 而后者可利用计算图并基于 $dz^{[2]}$ 直接得到

后向传播



1. 一般地，如果有 n 个训练样本，则我们有

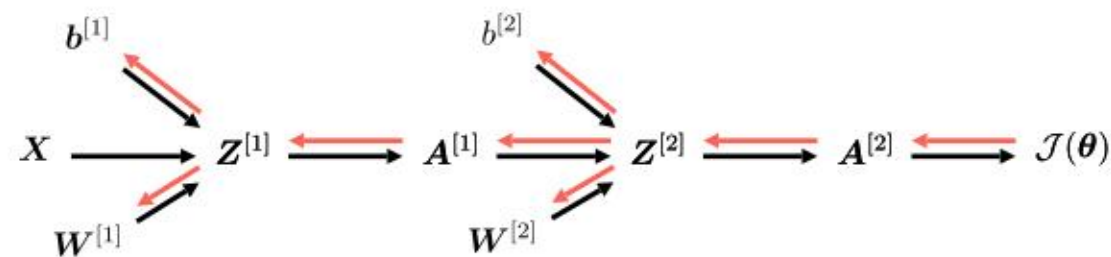
$$d\theta = \frac{1}{n} \sum_{i=1}^n \frac{\partial \mathcal{L}(y_i, a_i^{[2]})}{\partial \theta}$$

- $\mathcal{L}(y, a)$: 基于观测标签 y 及其估计值 a 的损失函数
- 针对于二分类问题，我们有 $\mathcal{L}(y, a) = -\{y \log a + (1 - y) \log(1 - a)\}$

2. 因此，为了得到 $d\theta$ ，我们只需要对训练样本对应的偏导数取平均即可

3. 利用向量化

后向传播

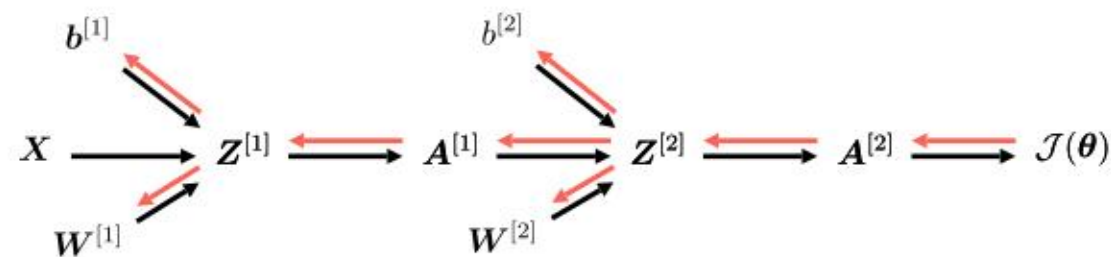


1. 基于向量化，我们有

$$\begin{aligned} d\mathbf{Z}^{[2]} &= n^{-1}(\mathbf{A}^{[2]} - \mathbf{Y}), & db^{[2]} &= (d\mathbf{Z}^{[2]})^T \cdot \mathbf{1}, & d\mathbf{W}^{[2]} &= (d\mathbf{Z}^{[2]})^T \cdot \mathbf{A}^{[1]}, \\ d\mathbf{Z}^{[1]} &= (d\mathbf{Z}^{[2]} \cdot \mathbf{W}^{[2]}) \circ \sigma'(\mathbf{Z}^{[1]}), & db^{[1]} &= (d\mathbf{Z}^{[1]})^T \cdot \mathbf{1}, & d\mathbf{W}^{[1]} &= (d\mathbf{Z}^{[1]})^T \cdot \mathbf{X} \end{aligned}$$

- $\sigma'(z)$: 对应于 $\sigma(z)$ 的导数
 - ▷ 对于 sigmoid 函数 $\sigma(z) = \{1 + \exp(-z)\}^{-1}$, 我们有 $\sigma'(z) = \sigma(z)\{1 - \sigma(z)\}$
- $\sigma'(\mathbf{Z}^{[1]}) = (\sigma'(z_{ij}^{[1]})) \in \mathbb{R}^{n \times d^{[1]}}$: 利用广播机制, 计算 $\mathbf{Z}^{[1]}$ 每个元素对应的导数值
 - ▷ $\mathbf{Z}^{[1]} = (z_{ij}^{[1]}) \quad (i = 1, \dots, n; j = 1, \dots, d^{[1]})$
- $\circ$: Hadamard 乘积 (矩阵对应位置元素相乘)

推导细节



1. 我们首先展示，输出层相关导数的计算细节

2. 回顾：代价函数为

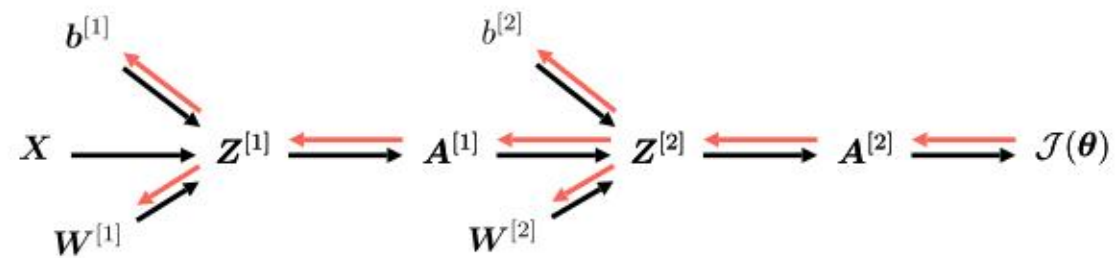
$$\mathcal{J} = -\frac{1}{n} \sum_{i=1}^n \left\{ y_i \log a_i^{[2]} + (1 - y_i) \log (1 - a_i^{[2]}) \right\}$$

3. 为了得到 $d\mathbf{A}^{[2]} = (\partial \mathcal{J} / \partial a_1^{[2]}, \dots, \partial \mathcal{J} / \partial a_n^{[2]})^T \in \mathbb{R}^{n \times 1}$ ，我们将 $\mathcal{J}$ 看成 $\mathbf{A}^{[2]}$ 的函数

- $\mathbf{A}^{[2]} = (a_1^{[2]}, \dots, a_n^{[2]})^T \in \mathbb{R}^{n \times 1}$
- $a_i^{[2]}$ 为对 $P(y_i = 1 \mid \mathbf{x}_i)$ 的估计量

4. 我们可得到 $\partial \mathcal{J} / \partial a_i^{[2]} = n^{-1} (a_i^{[2]} - y_i) / \{a_i^{[2]} (1 - a_i^{[2]})\}$

推导细节



1. 因此，我们有

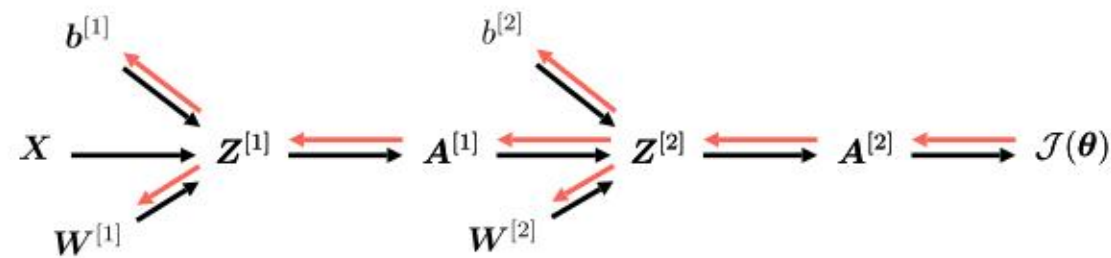
$$\begin{aligned} d\mathbf{A}^{[2]} &= n^{-1}((a_1^{[2]} - y_1)/\{a_1^{[2]}(1 - a_1^{[2]})\}, \dots, (a_n^{[2]} - y_n)/\{a_n^{[2]}(1 - a_n^{[2]})\})^T \\ &= n^{-1} \mathbf{D}^{-1} \cdot (\mathbf{A}^{[2]} - \mathbf{Y}) \end{aligned}$$

- $\mathbf{D} \in \mathbb{R}^{n \times n}$ 为对角阵，其第 i 个对角元素为 $a_i^{[2]}(1 - a_i^{[2]})$

2. 为了得到 $d\mathbf{Z}^{[2]} = (\partial \mathcal{J} / \partial z_1^{[2]}, \dots, \partial \mathcal{J} / \partial z_n^{[2]})^T \in \mathbb{R}^{n \times 1}$ ，我们利用链式法则，并将 $\mathcal{J}$ 看成 $\mathbf{A}^{[2]}$ 的函数，将 $\mathbf{A}^{[2]}$ 看成 $\mathbf{Z}^{[2]}$ 的函数

- $\mathbf{Z}^{[2]} = (z_1^{[2]}, \dots, z_n^{[2]})^T \in \mathbb{R}^{n \times 1}$
- $z_i^{[2]}$ 是第 i 个训练样本，在模型第二层线性变换的结果

推导细节



1. 由于 $\mathbf{A}^{[2]}$ 和 $\mathbf{Z}^{[2]}$ 都是列向量，我们有

$$d\mathbf{Z}^{[2]} = \frac{\partial \mathcal{J}}{\partial \mathbf{Z}^{[2]}} = \frac{\partial (\mathbf{A}^{[2]})^T}{\partial \mathbf{Z}^{[2]}} \cdot \frac{\partial \mathcal{J}}{\partial \mathbf{A}^{[2]}} = \frac{\partial (\mathbf{A}^{[2]})^T}{\partial \mathbf{Z}^{[2]}} \cdot n^{-1} \cdot \mathbf{D}^{-1} \cdot (\mathbf{A}^{[2]} - \mathbf{Y})$$

2. 注意到

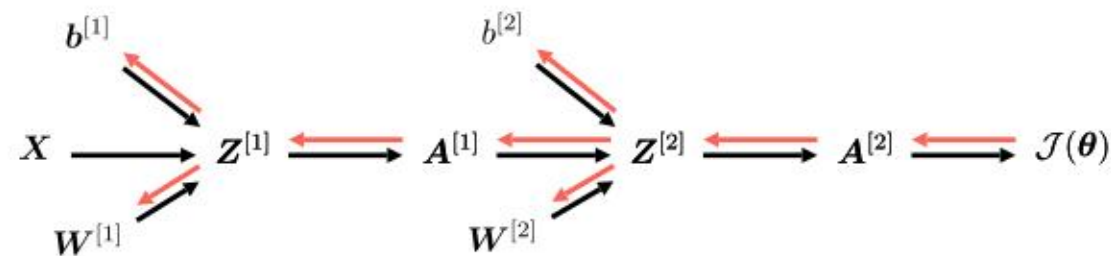
$$\frac{\partial (\mathbf{A}^{[2]})^T}{\partial \mathbf{Z}^{[2]}} = \mathbf{D}$$

- 由于 $a_i^{[2]}$ 只是关于 $z_i^{[2]}$ 的函数, $\partial (\mathbf{A}^{[2]})^T / \partial \mathbf{Z}^{[2]}$ 的结果为对角阵
- 由于 $a_i^{[2]} = \sigma(z_i^{[2]})$, 故 $\partial a_i^{[2]} / \partial z_i^{[2]} = a_i^{[2]}(1 - a_i^{[2]})$, 其正是 $\mathbf{D}$ 的第 i 个对角元素

3. 因此，我们有

$$d\mathbf{Z}^{[2]} = n^{-1}(\mathbf{A}^{[2]} - \mathbf{Y})$$

推导细节

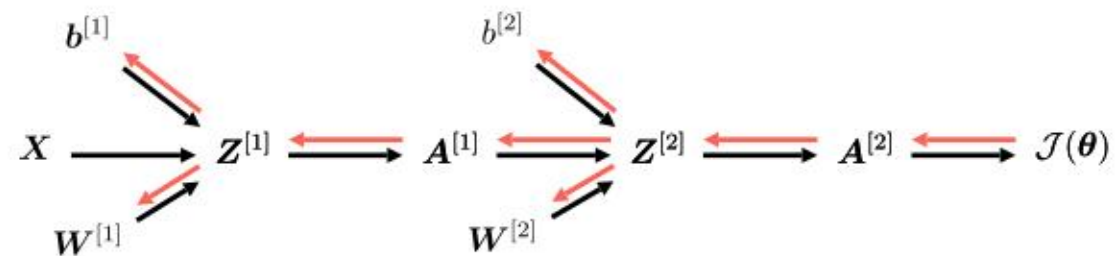


1. 为了得到 $db^{[2]} = \partial \mathcal{J} / \partial b^{[2]}$ ，我们利用链式法则，并将 $\mathcal{J}$ 看成 $\mathbf{Z}^{[2]}$ 的函数，将 $\mathbf{Z}^{[2]}$ 看成 $b^{[2]}$ 的函数
2. 因此，我们有

$$db^{[2]} = \frac{\partial \mathcal{J}}{\partial b^{[2]}} = \frac{\partial (\mathbf{Z}^{[2]})^T}{\partial b^{[2]}} \cdot \frac{\partial \mathcal{J}}{\partial \mathbf{Z}^{[2]}} = \mathbf{1}^T \cdot d\mathbf{Z}^{[2]} = (d\mathbf{Z}^{[2]})^T \cdot \mathbf{1}$$

- $\mathbf{1} = (1, \dots, 1)^T$ 是一个元素均为 1 的长度为 n 的列向量
- $\mathbf{Z}^{[2]}$ 每个元素都是 $b^{[2]}$ 的函数: $z_i^{[2]} = b^{[2]} + \mathbf{W}^{[2]} \cdot \mathbf{a}_i^{[1]}$
- 因此, $\partial z_i^{[2]} / \partial b^{[2]} = 1$ 以及 $\partial \mathbf{Z}^{[2]} / \partial b^{[2]} = (\partial z_1^{[2]} / \partial b^{[2]}, \dots, \partial z_n^{[2]} / \partial b^{[2]})^T = \mathbf{1}$

推导细节

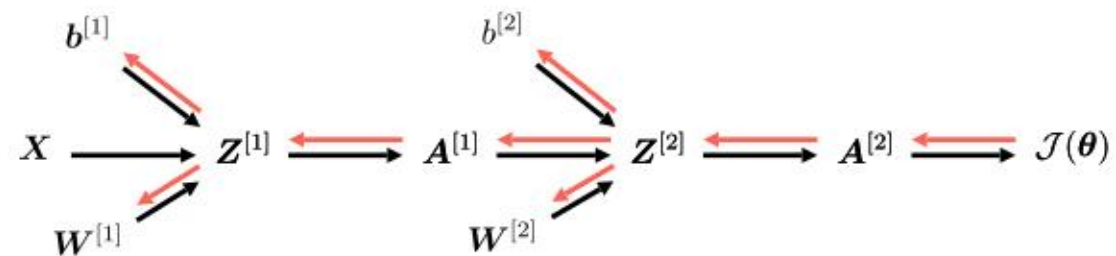


1. 为了得到 $d\mathbf{W}^{[2]} = \partial \mathcal{J} / \partial \mathbf{W}^{[2]}$ ，我们利用链式法则，并将 $\mathcal{J}$ 看成 $\mathbf{Z}^{[2]}$ 的函数，将 $\mathbf{Z}^{[2]}$ 看成 $\mathbf{W}^{[2]}$ 的函数
2. 由于 $\mathbf{W}^{[2]}$ 是一个行向量，而不是列向量，我们有

$$d\mathbf{W}^{[2]} = \sum_{i=1}^n \frac{\partial z_i^{[2]}}{\partial \mathbf{W}^{[2]}} \cdot \frac{\partial \mathcal{J}}{\partial z_i^{[2]}} = \left(\frac{\partial \mathcal{J}}{\partial z_1^{[2]}}, \dots, \frac{\partial \mathcal{J}}{\partial z_n^{[2]}} \right) \cdot \begin{pmatrix} (\mathbf{a}_1^{[1]})^T \\ \vdots \\ (\mathbf{a}_n^{[1]})^T \end{pmatrix} = (d\mathbf{Z}^{[2]})^T \cdot \mathbf{A}^{[1]}$$

- $\mathbf{Z}^{[2]}$ 每个元素都是 $\mathbf{W}^{[2]}$ 的函数: $z_i^{[2]} = b^{[2]} + \mathbf{W}^{[2]} \cdot \mathbf{a}_i^{[1]}$
- 因此, $\partial z_i^{[2]} / \partial \mathbf{W}^{[2]} = (\mathbf{a}_i^{[1]})^T$
- 利用向量化, 我们得到最后一个等式

推导细节

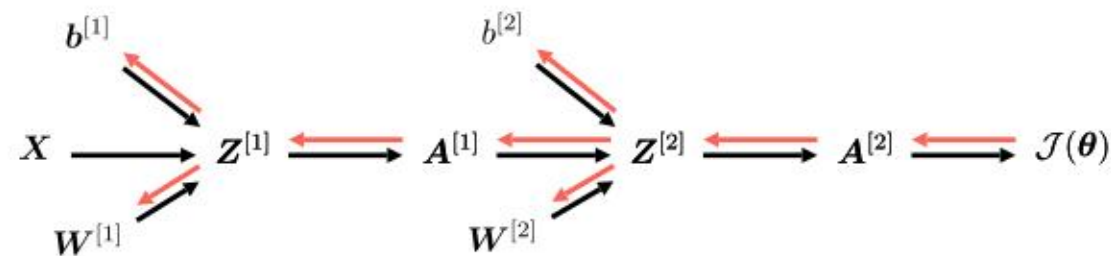


1. 由于 $(\mathbf{W}^{[2]})^T$ 是列向量，得到 $d\mathbf{W}^{[2]}$ 的第二种方法为

$$d\mathbf{W}^{[2]} = \left(\frac{\partial \mathcal{J}}{\partial (\mathbf{W}^{[2]})^T} \right)^T = \left(\frac{\partial (\mathbf{Z}^{[2]})^T}{\partial (\mathbf{W}^{[2]})^T} \cdot \frac{\partial \mathcal{J}}{\partial \mathbf{Z}^{[2]}} \right)^T = (d\mathbf{Z}^{[2]})^T \cdot \mathbf{A}^{[1]}$$

- $\mathbf{Z}^{[2]}$ 每个元素都是 $\mathbf{W}^{[2]}$ 的函数: $z_i^{[2]} = b^{[2]} + \mathbf{W}^{[2]} \cdot \mathbf{a}_i^{[1]}$
- 因此, $\partial z_i^{[2]} / \partial (\mathbf{W}^{[2]})^T = \mathbf{a}_i^{[1]}$ 以及 $\partial (\mathbf{Z}^{[2]})^T / \partial (\mathbf{W}^{[2]})^T = (\mathbf{a}_1^{[1]}, \dots, \mathbf{a}_n^{[1]}) = (\mathbf{A}^{[1]})^T$

推导细节



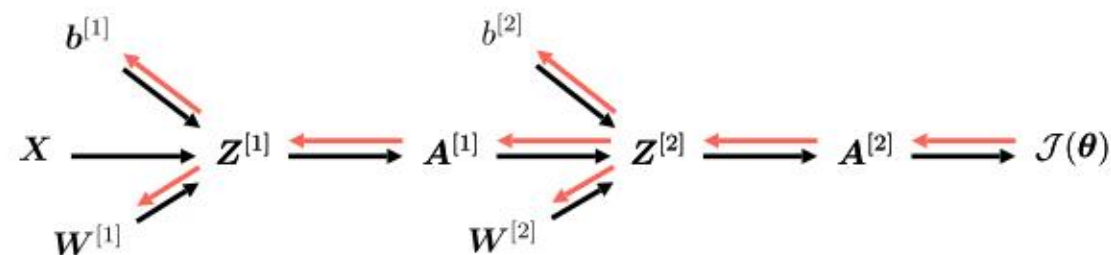
1. 为了得到 $d\mathbf{A}^{[1]} = \partial\mathcal{J}/\partial\mathbf{A}^{[1]}$ ，我们利用链式法则，并将 $\mathcal{J}$ 看成 $\mathbf{Z}^{[2]}$ 的函数，将 $\mathbf{Z}^{[2]}$ 看成 $\mathbf{A}^{[1]}$ 的函数

2. 由于 $\mathbf{A}^{[1]} \in \mathbb{R}^{n \times d^{[1]}}$ 是一个矩阵，而不是列向量，我们有

$$d\mathbf{A}^{[1]} = \sum_{i=1}^n \frac{\partial z_i^{[2]}}{\partial \mathbf{A}^{[1]}} \cdot \frac{\partial \mathcal{J}}{\partial z_i^{[2]}} = \sum_{i=1}^n \frac{\partial \mathcal{J}}{\partial z_i^{[2]}} \cdot \mathbf{e}_i \cdot \mathbf{W}^{[2]} = \frac{\partial \mathcal{J}}{\partial \mathbf{Z}^{[2]}} \cdot \mathbf{W}^{[2]} = d\mathbf{Z}^{[2]} \cdot \mathbf{W}^{[2]}$$

- $\partial z_i^{[2]}/\partial \mathbf{A}^{[1]} \in \mathbb{R}^{n \times d^{[1]}}$ 与 $\mathbf{A}^{[1]}$ 具有相同维度
- $z_i^{[2]}$ 只是 $\mathbf{A}^{[1]}$ 第 i 行的函数: $z_i^{[2]} = b^{[2]} + \mathbf{W}^{[2]} \cdot \mathbf{a}_i^{[1]}$
- 因此, $\partial z_i^{[2]}/\partial (\mathbf{a}_i^{[1]})^T = \mathbf{W}^{[2]}$ 以及 $\partial z_i^{[2]}/\partial \mathbf{A}^{[1]} = \mathbf{e}_i \cdot \mathbf{W}^{[2]}$
- $\mathbf{e}_i = (0, \dots, 1, \dots, 0)^T$ 是一个长度为 n 的列向量; 除了第 i 个元素为 1 外, 其他元素均为 0

推导细节



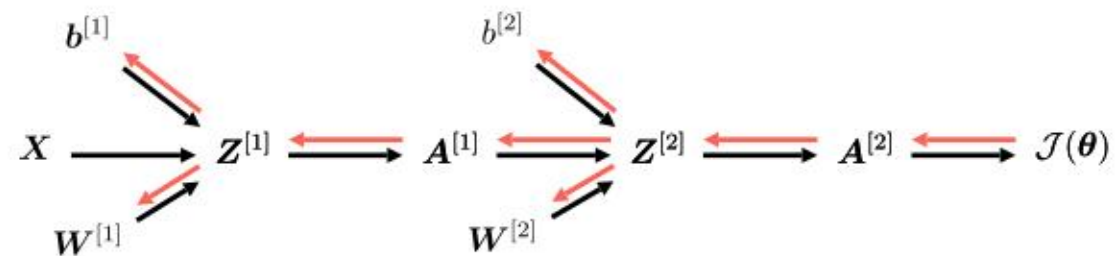
1. 为了得到 $d\mathbf{Z}^{[1]} = \partial \mathcal{J} / \partial \mathbf{Z}^{[1]}$, 我们将 $\mathcal{J}$ 看成 $\mathbf{A}^{[1]}$ 的函数, 将 $\mathbf{A}^{[1]}$ 看成 $\mathbf{Z}^{[1]}$ 的函数

2. 由于 $\mathbf{A}^{[1]} \in \mathbb{R}^{n \times d^{[1]}}$, 我们有

$$d\mathbf{Z}^{[1]} = \sum_{i=1}^n \sum_{j=1}^{d^{[1]}} \frac{\partial a_{ij}^{[1]}}{\partial \mathbf{Z}^{[1]}} \cdot \frac{\partial \mathcal{J}}{\partial a_{ij}^{[1]}} = \sum_{i=1}^n \sum_{j=1}^{d^{[1]}} \frac{\partial \mathcal{J}}{\partial a_{ij}^{[1]}} \cdot \sigma'(z_{ij}^{[1]}) \cdot \mathbf{e}_i \cdot \mathbf{h}_j^T = d\mathbf{A}^{[1]} \circ \sigma'(\mathbf{Z}^{[1]})$$

- $\partial a_{ij}^{[1]} / \partial \mathbf{Z}^{[1]} \in \mathbb{R}^{n \times d^{[1]}}$ 与 $\mathbf{Z}^{[1]}$ 具有相同维度
- $a_{ij}^{[1]}$ 只是 $\mathbf{Z}^{[1]}$ 的第 (i, j) 个元素的函数: $a_{ij}^{[1]} = \sigma(z_{ij}^{[1]})$
- 因此, $\partial a_{ij}^{[1]} / \partial z_{ij}^{[1]} = \sigma'(z_{ij}^{[1]})$ 以及 $\partial a_{ij}^{[1]} / \partial \mathbf{Z}^{[1]} = \sigma'(z_{ij}^{[1]}) \cdot \mathbf{e}_i \cdot \mathbf{h}_j^T$
- $\mathbf{e}_i = (0, \dots, 1, \dots, 0)^T$ 是一个长度为 n 的列向量; 除了第 i 个元素为 1 外, 其他元素均为 0
- $\mathbf{h}_j = (0, \dots, 1, \dots, 0)^T$ 是一个长度为 $d^{[1]}$ 的列向量; 除了第 j 个元素为 1 外, 其他元素均为 0

推导细节



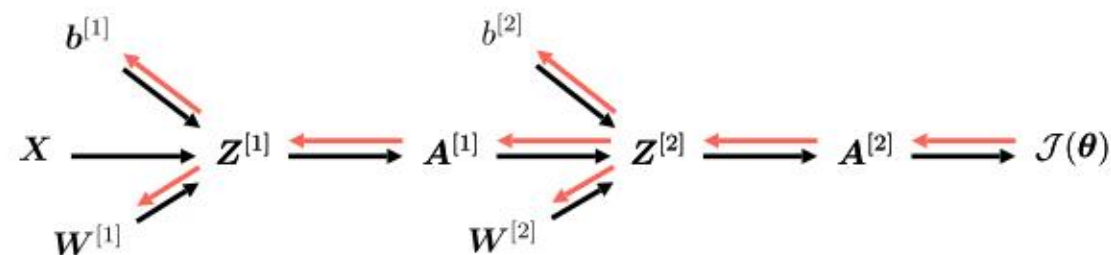
1. 为了得到 $d\mathbf{b}^{[1]} = \partial \mathcal{J} / \partial \mathbf{b}^{[1]}$, 我们将 $\mathcal{J}$ 看成 $\mathbf{Z}^{[1]} = (z_1^{[1]}, \dots, z_n^{[1]})^T$ 的函数, 将 $\mathbf{Z}^{[1]}$ 看成 $\mathbf{b}^{[1]}$ 的函数

2. 由于 $\mathbf{b}^{[1]} \in \mathbb{R}^{d^{[1]} \times 1}$ 是列向量, 我们有

$$d\mathbf{b}^{[1]} = \sum_{i=1}^n \frac{\partial (z_i^{[1]})^T}{\partial \mathbf{b}^{[1]}} \cdot \frac{\partial \mathcal{J}}{\partial z_i^{[1]}} = \sum_{i=1}^n \mathbf{I} \cdot \frac{\partial \mathcal{J}}{\partial z_i^{[1]}} = \sum_{i=1}^n \frac{\partial \mathcal{J}}{\partial z_i^{[1]}} = (d\mathbf{Z}^{[1]})^T \cdot \mathbf{1}$$

- $\mathbf{Z}^{[1]}$ 每一行都是 $\mathbf{b}^{[1]}$ 的函数: $(z_i^{[1]})^T = (\mathbf{b}^{[1]})^T + \mathbf{x}_i^T \cdot (\mathbf{W}^{[1]})^T$
- 因此, $\partial (z_i^{[1]})^T / \partial \mathbf{b}^{[1]} = \mathbf{I}$
- $\mathbf{I}$: 维度为 $d^{[1]} \times d^{[1]}$ 的单位阵

推导细节



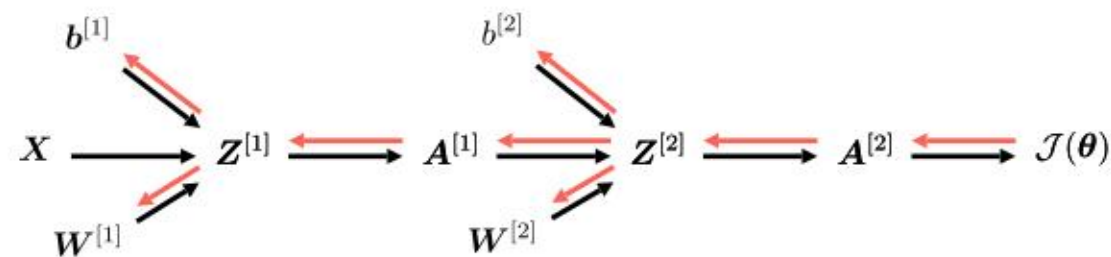
1. 为了得到 $d\mathbf{W}^{[1]} = \partial \mathcal{J} / \partial \mathbf{W}^{[1]}$, 我们将 $\mathcal{J}$ 看成 $\mathbf{Z}^{[1]} = (z_{ij}^{[1]})$ 的函数, 将 $\mathbf{Z}^{[1]}$ 看成 $\mathbf{W}^{[1]}$ 的函数

2. 由于 $\mathbf{W}^{[1]} \in \mathbb{R}^{d^{[1]} \times d^{[0]}}$ 是一个矩阵, 我们有

$$d\mathbf{W}^{[1]} = \sum_{i=1}^n \sum_{j=1}^{d^{[1]}} \frac{\partial z_{ij}^{[1]}}{\partial \mathbf{W}^{[1]}} \cdot \frac{\partial \mathcal{J}}{\partial z_{ij}^{[1]}} = \sum_{i=1}^n \sum_{j=1}^{d^{[1]}} \frac{\partial \mathcal{J}}{\partial z_{ij}^{[1]}} \cdot \mathbf{h}_j \cdot \mathbf{x}_i^T = \sum_{i=1}^n \frac{\partial \mathcal{J}}{\partial \mathbf{z}_i^{[1]}} \cdot \mathbf{x}_i^T = (d\mathbf{Z}^{[1]})^T \cdot \mathbf{X}$$

- $z_{ij}^{[1]}$ 只是 $\mathbf{W}^{[1]} = (\mathbf{w}_1^{[1]}, \dots, \mathbf{w}_{d^{[1]}}^{[1]})^T$ 第 j 行的函数: $z_{ij}^{[1]} = b_j^{[1]} + (\mathbf{w}_j^{[1]})^T \cdot \mathbf{x}_i$
- 因此, $\partial z_{ij}^{[1]} / \partial \mathbf{W}^{[1]} = \mathbf{h}_j \cdot \mathbf{x}_i^T$
- $\mathbf{h}_j = (0, \dots, 1, \dots, 0)^T$ 是一个长度为 $d^{[1]}$ 的列向量; 除了第 j 个元素为 1 外, 其他元素均为 0

备注



1. 简单起见，记 $\mathbf{A}^{[0]} = \mathbf{X}$

2. 前向传播和后向传播均可使用 for 循环实现

- **前向传播**：基于当前模型参数， $\mathbf{A}^{[l-1]}$ 可被视为计算 $\mathbf{Z}^{[l]}$ 和 $\mathbf{A}^{[l]}$ 的“输入”
- **后向传播**： $d\mathbf{A}^{[l]}$ 是用来计算 $d\mathbf{Z}^{[l]}$ 、 $d\mathbf{b}^{[l]}$ 、 $d\mathbf{W}^{[l]}$ 以及 $d\mathbf{A}^{[l-1]}$ 的“输入”

批量梯度下降法

步骤1. 随机初始化 $\theta^{(0)}$

步骤2. 基于 $\theta^{(t)}$ 计算

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

步骤3. 更新参数

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

步骤4. 回到步骤 2 直至收敛，或达到最大迭代次数

课程总结

1. 将单隐层神经网络由单样本计算推广到多样本的向量化训练

- 数据：将各训练样本的特征向量按行堆叠，构成设计矩阵 $\mathbf{X}$
- 前向传播：利用当前模型参数，计算代价函数值，借助广播机制提高计算效率
- 后向传播：基于链式法则和向量化，得到代价函数对参数的平均梯度值
- 梯度下降法