



AI and Deep Learning

2-2 Neural Network with One Hidden Layer II

(Multiple training examples)

Zhonglei Wang

WISE, SOE, CHOW

XMU

(Version v1)

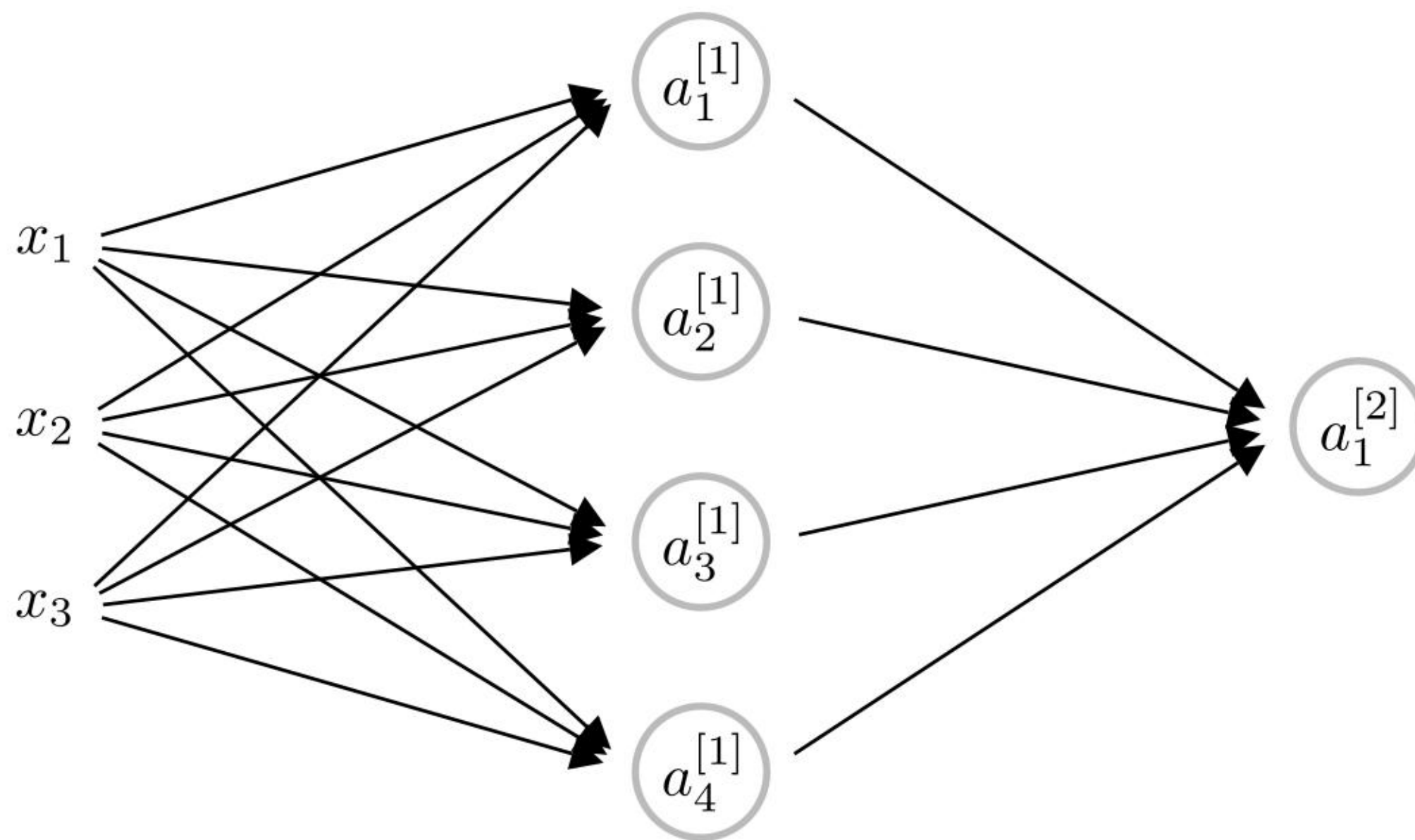
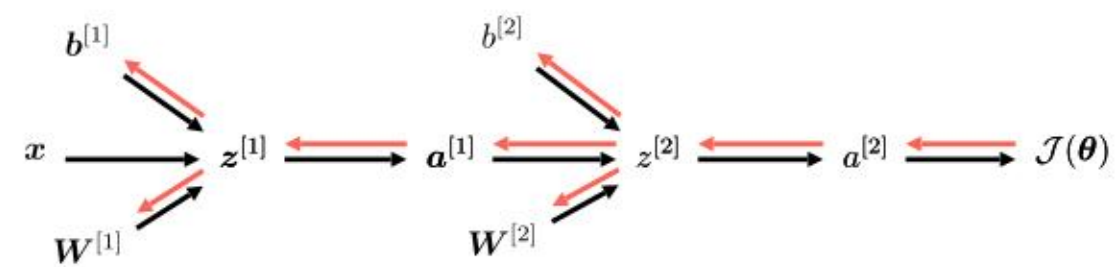
Contents

1. Forward propagation

2. Backpropagation

3. (Batch) Gradient descent algorithm

Recall

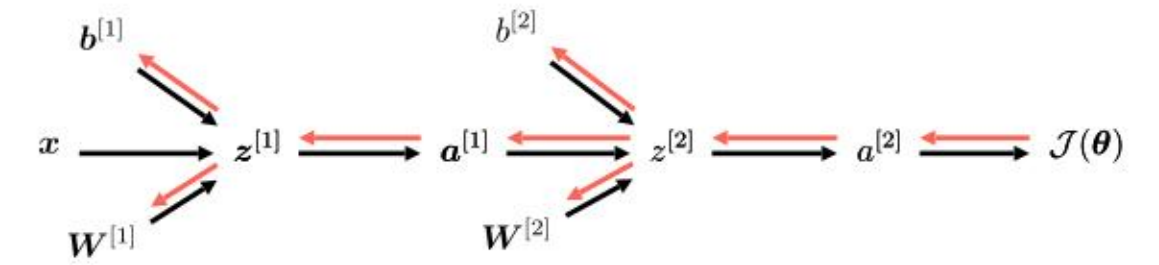


Input layer ($d^{[0]} = 3$)

1st hidden layer ($d^{[1]} = 4$)

Output layer ($d^{[2]} = 1$)

Recall

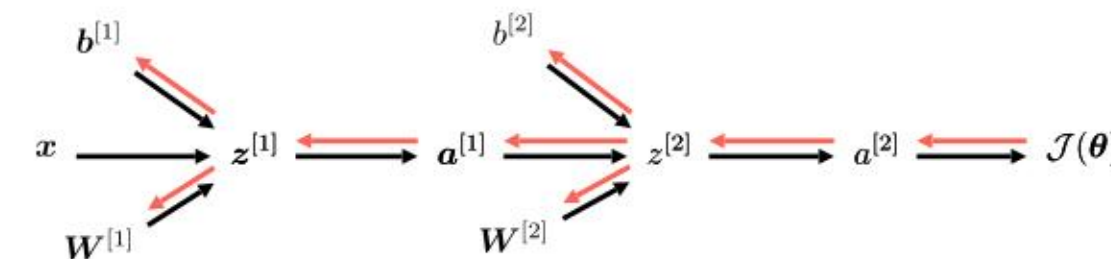


1. Recall that

- 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. We consider a binary classification problem with $y_i \in \{0, 1\}$

Recall



1. If we have only **ONE** training example (x, y) , forward propagation is

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

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

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

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

- **Broadcasting** is used to obtain $z^{[1]}$ and $a^{[1]}$
2. **Column vectors** are involved above
3. Cost function

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

Forward propagation

1. Suppose we have n training examples $\{(\mathbf{x}_i, y_i) : i = 1, \dots, n\}$

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

3. Forward propagation (using vectorization)

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

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

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

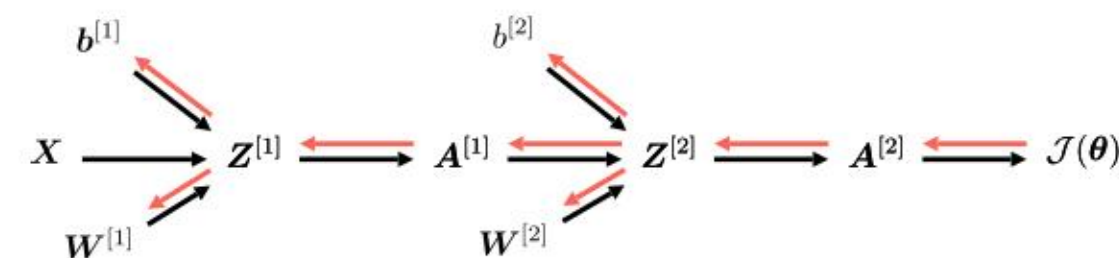
$$\mathbf{A}^{[2]} = \sigma(\mathbf{Z}^{[2]}) \in \mathbb{R}^{n \times d^{[2]}}$$

- **Broadcasting** in Python is implicitly used above

4. Actually, we stack vectors **row-wisely**

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

Forward propagation



1. Suppose we have n training examples $\{(\mathbf{x}_i, y_i) : i = 1, \dots, n\}$

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

3. Forward propagation (using vectorization)

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

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

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

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

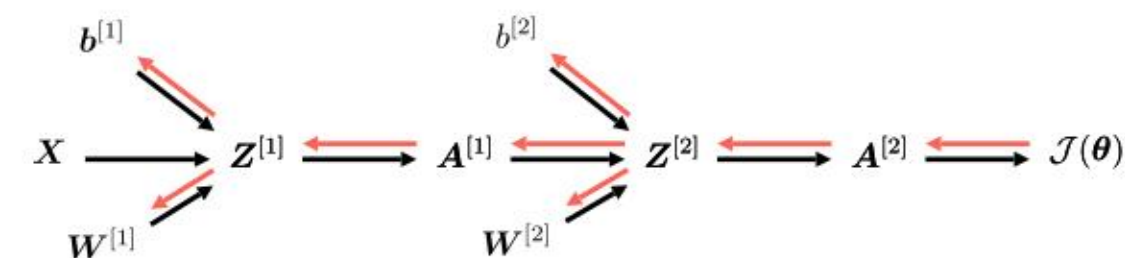
$$(\mathbf{Z}^{[1]} = (Z_{ij}^{[1]}); i = 1, \dots, n; j = 1, \dots, d^{[1]})$$

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

- **Broadcasting** in Python is implicitly used above

4. Actually, we stack vectors **row-wisely**

Forward propagation

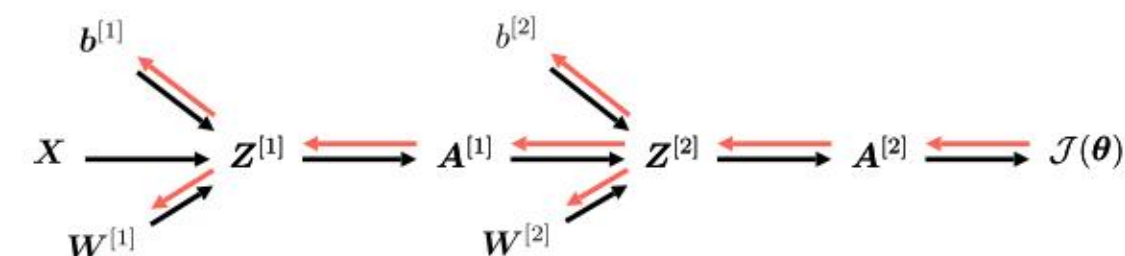


1. Broadcasting makes Python code simpler, faster, and more memory-efficient

- Code simplicity and readability
- Superior memory efficiency: Avoiding temporary copies
- Massive performance gains: Leveraging optimized C/Fortran code



Forward propagation



1. For the previous example, we have $d^{[0]} = 3$, $d^{[1]} = 4$, $d^{[2]} = 1$

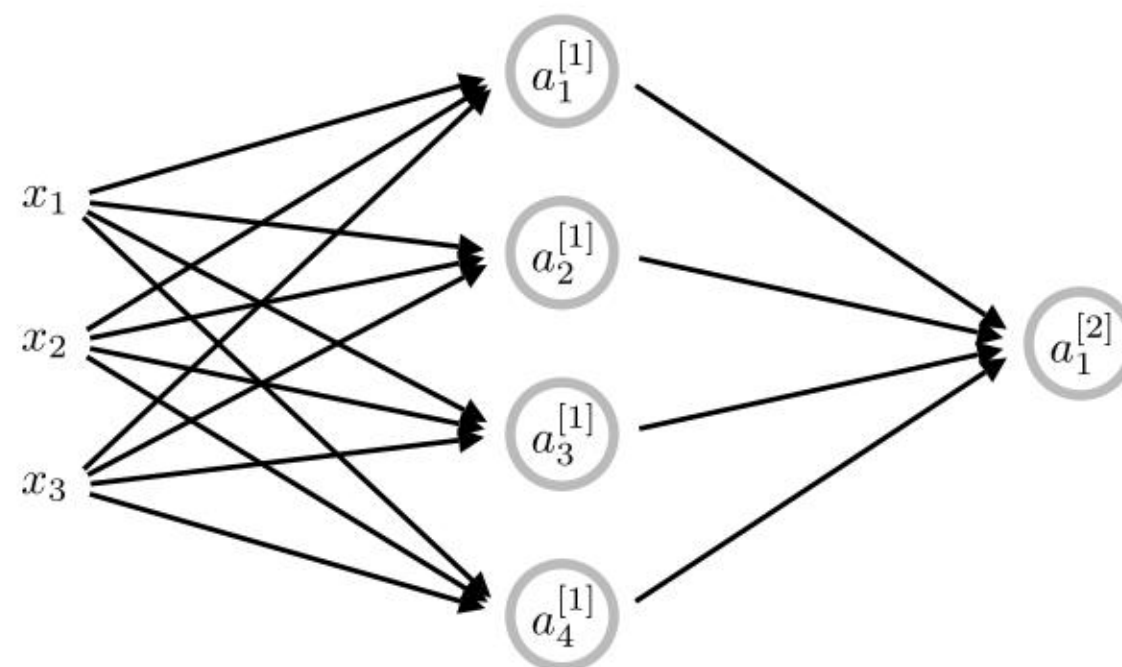
2. Thus,

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

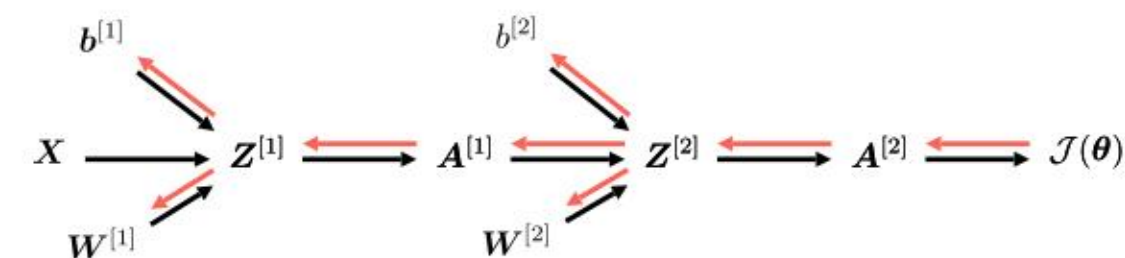
$$\mathbf{A}^{[1]} = \sigma(\mathbf{Z}^{[1]}) \in \mathbb{R}^{n \times 4},$$

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

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



Forward propagation



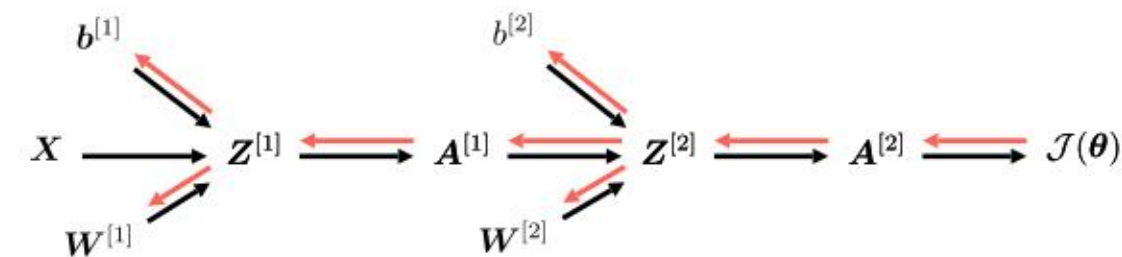
1. Cost function (cross entropy for binary classification problems)

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

- $a_i^{[2]}$: Estimation for $P(y_i | \mathbf{x}_i)$

2. The form of the cost function is determined by a specific problem

Backpropagation



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

2. Recall: For only one training example $(\mathbf{x}, y)$, we have

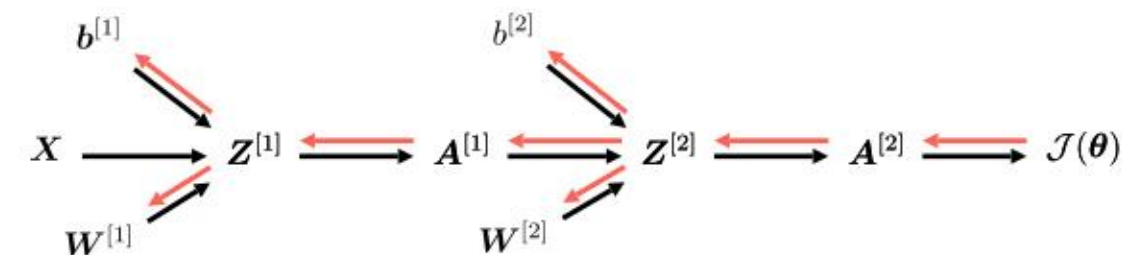
$$\begin{aligned} dz^{[2]} &= a^{[2]} - y, \quad db^{[2]} = dz^{[2]}, \quad d\mathbf{W}^{[2]} = dz^{[2]} \cdot (\mathbf{a}^{[1]})^T, \\ d\mathbf{z}^{[1]} &= dz^{[2]} \cdot \mathbf{D} \cdot (\mathbf{W}^{[2]})^T, \quad d\mathbf{b}^{[1]} = dz^{[1]}, \quad d\mathbf{W}^{[1]} = d\mathbf{z}^{[1]} \cdot \mathbf{x}^T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

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

3. Note that

- $dz^{[2]}$ can be obtained from $da^{[2]}$, determined by the cost function
- $d\mathbf{z}^{[1]}$ can be obtained from $d\mathbf{a}^{[1]}$, derived from $dz^{[2]}$ directly by the computation graph

Backpropagation



1. For general case, if we have n examples, the cost function is

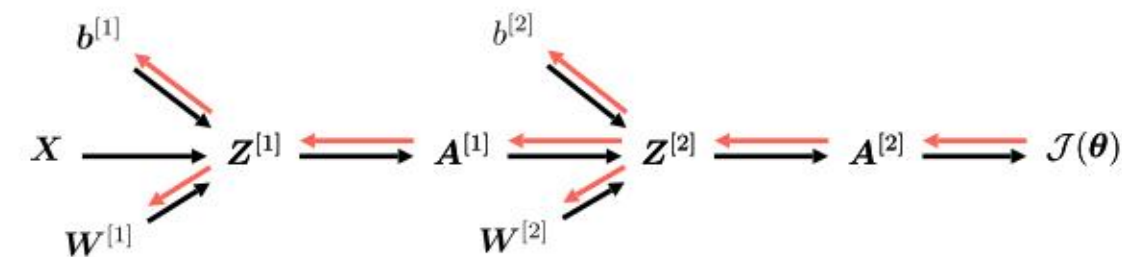
$$d\theta = \frac{1}{n} \sum_{i=1}^n \frac{\partial \mathcal{L}(y_i, a_i^{[2]})}{\partial \theta}$$

- $\mathcal{L}(y, a)$: Loss function based on true label y and its estimator a
- Recall that $\mathcal{L}(y, a) = -\{y \log a + (1 - y) \log(1 - a)\}$ for binary classification problems

2. Thus, to obtain $d\theta$, we only need to take average of derivatives for training examples

3. We use vectorization

Backpropagation

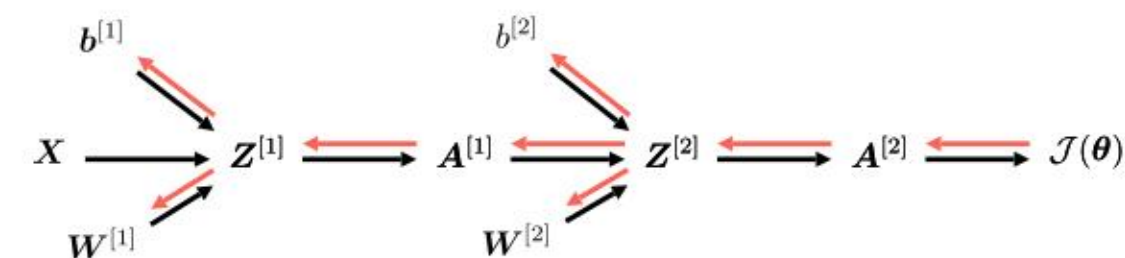


1. Then, we have

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

- $\sigma'(z)$: Derivative of $\sigma(z)$
 - ▷ $\sigma'(z) = \sigma(z)\{1 - \sigma(z)\}$ for sigmoid function $\sigma(z) = \{1 + \exp(-z)\}^{-1}$
- $\sigma'(\mathbf{Z}^{[1]}) = (\sigma'(z_{ij}^{[1]})) \in \mathbb{R}^{n \times d^{[1]}}$: Derivatives for each element in $\mathbf{Z}^{[1]}$
 - ▷ $\mathbf{Z}^{[1]} = (z_{ij}^{[1]}) \quad (i = 1, \dots, n; j = 1, \dots, d^{[1]})$
- $\circ$: Hadamard production (element-wise production)

Detailed derivation



1. We first show the derivation of the derivatives associated with the output layer

2. Recall that the cost function is

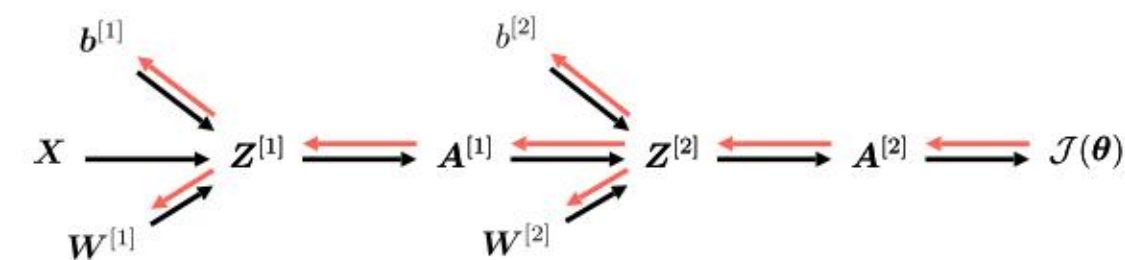
$$\mathcal{J} = -\frac{1}{n} \sum_{i=1}^n \left\{ y_i \log a_i^{[2]} + (1 - y_i) \log (1 - a_i^{[2]}) \right\}$$

3. To obtain $d\mathbf{A}^{[2]} = (\partial \mathcal{J} / \partial a_1^{[2]}, \dots, \partial \mathcal{J} / \partial a_n^{[2]})^T \in \mathbb{R}^{n \times 1}$, we treat $\mathcal{J}$ as a function of $\mathbf{A}^{[2]}$

- $\mathbf{A}^{[2]} = (a_1^{[2]}, \dots, a_n^{[2]})^T \in \mathbb{R}^{n \times 1}$
- $a_i^{[2]}$ is the estimated probability for $P(y_i = 1 \mid \mathbf{x}_i)$

4. We can easily show that $\partial \mathcal{J} / \partial a_i^{[2]} = n^{-1} (a_i^{[2]} - y) / \{a_i^{[2]} (1 - a_i^{[2]})\}$

Detailed derivation



1. Thus, we have

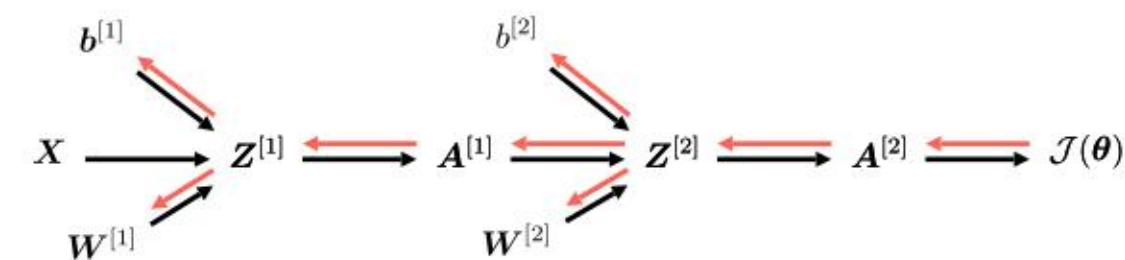
$$\begin{aligned} d\mathbf{A}^{[2]} &= n^{-1}((a_1^{[2]} - y_1)/\{a_1^{[2]}(1 - a_1^{[2]})\}, \dots, (a_n^{[2]} - y_n)/\{a_n^{[2]}(1 - a_n^{[2]})\})^T \\ &= n^{-1} \mathbf{D}^{-1} \cdot (\mathbf{A}^{[2]} - \mathbf{Y}) \end{aligned}$$

- $\mathbf{D} \in \mathbb{R}^{n \times n}$ is a diagonal matrix, whose (i, i) th entry is $a_i^{[2]}(1 - a_i^{[2]})$

2. To obtain $d\mathbf{Z}^{[2]} = (\partial \mathcal{J} / \partial z_1^{[2]}, \dots, \partial \mathcal{J} / \partial z_n^{[2]})^T \in \mathbb{R}^{n \times 1}$, we use the chain rule and treat $\mathcal{J}$ as a function of $\mathbf{A}^{[2]}$, and $\mathbf{A}^{[2]}$ as a function of $\mathbf{Z}^{[2]}$

- $\mathbf{Z}^{[2]} = (z_1^{[2]}, \dots, z_n^{[2]})^T \in \mathbb{R}^{n \times 1}$
- $z_i^{[2]}$ is the result after linear transformation in the second layer for the i th training example

Detailed derivation



1. Since both $\mathbf{A}^{[2]}$ and $\mathbf{Z}^{[2]}$ are column vectors, we have

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

2. Notice that

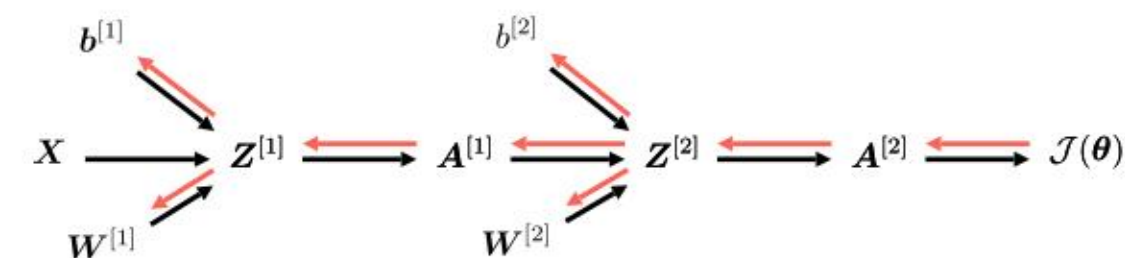
$$\frac{\partial (\mathbf{A}^{[2]})^T}{\partial \mathbf{Z}^{[2]}} = \mathbf{D}$$

- Since $a_i^{[2]}$ is only a function of $z_i^{[2]}$, $\partial (\mathbf{A}^{[2]})^T / \partial \mathbf{Z}^{[2]}$ is diagonal
- Since $a_i^{[2]} = \sigma(z_i^{[2]})$, $\partial a_i^{[2]} / \partial z_i^{[2]} = a_i^{[2]}(1 - a_i^{[2]})$, exactly the (i, i) th entry of $\mathbf{D}$

3. Thus, we have

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

Detailed derivation

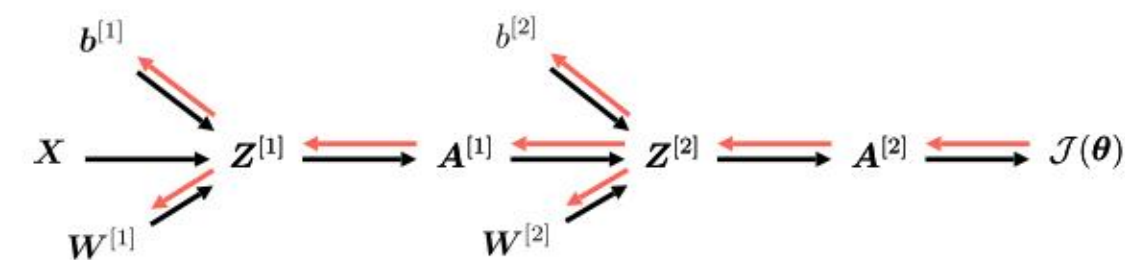


1. To obtain $db^{[2]} = \partial \mathcal{J} / \partial b^{[2]}$, we use the chain rule and treat $\mathcal{J}$ as a function of $\mathbf{Z}^{[2]}$,
and $\mathbf{Z}^{[2]}$ as a function of $b^{[2]}$
2. Thus, we have

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

- $\mathbf{1} = (1, \dots, 1)^T$ is a column vector of length n with all elements being 1
- Each element in $\mathbf{Z}^{[2]}$ is a function of $b^{[2]}$: $z_i^{[2]} = b^{[2]} + \mathbf{W}^{[2]} \cdot \mathbf{a}_i^{[1]}$
- Thus, $\partial z_i^{[2]} / \partial b^{[2]} = 1$ and $\partial \mathbf{Z}^{[2]} / \partial b^{[2]} = (\partial z_1^{[2]} / \partial b^{[2]}, \dots, \partial z_n^{[2]} / \partial b^{[2]})^T = \mathbf{1}$

Detailed derivation

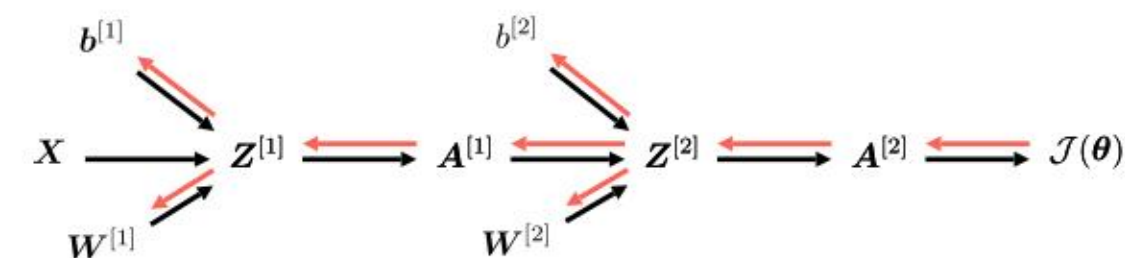


1. To obtain $d\mathbf{W}^{[2]} = \partial \mathcal{J} / \partial \mathbf{W}^{[2]}$, we use the chain rule and treat $\mathcal{J}$ as a function of $\mathbf{Z}^{[2]}$, and $\mathbf{Z}^{[2]}$ as a function of $\mathbf{W}^{[2]}$
2. Since $\mathbf{W}^{[2]}$ is a row vector, not a column one, we have

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

- Each element in $\mathbf{Z}^{[2]}$ is a function of $\mathbf{W}^{[2]}$: $z_i^{[2]} = b^{[2]} + \mathbf{W}^{[2]} \cdot \mathbf{a}_i^{[1]}$
- Thus, $\partial z_i^{[2]} / \partial \mathbf{W}^{[2]} = (\mathbf{a}_i^{[1]})^T$
- The last equation holds by vectorization

Detailed derivation

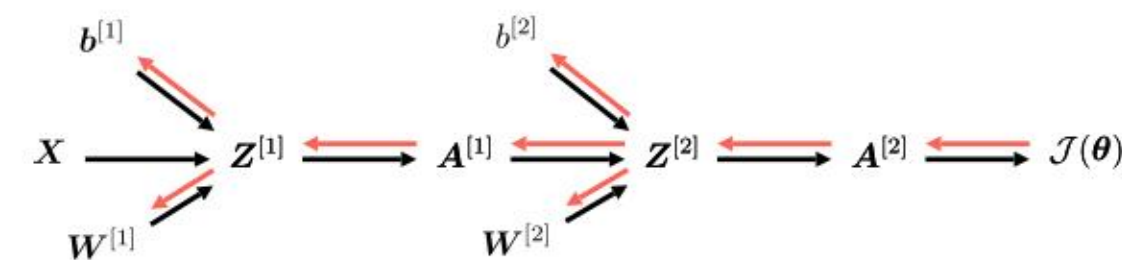


1. Since $(\mathbf{W}^{[2]})^T$ is a column vector, another way to obtain $d\mathbf{W}^{[2]}$ is

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

- Each element in $\mathbf{Z}^{[2]}$ is a function of $\mathbf{W}^{[2]}$: $z_i^{[2]} = b^{[2]} + \mathbf{W}^{[2]} \cdot \mathbf{a}_i^{[1]}$
- Thus, $\partial z_i^{[2]} / \partial (\mathbf{W}^{[2]})^T = \mathbf{a}_i^{[1]}$ and $\partial (\mathbf{Z}^{[2]})^T / \partial (\mathbf{W}^{[2]})^T = (\mathbf{a}_1^{[1]}, \dots, \mathbf{a}_n^{[1]}) = (\mathbf{A}^{[1]})^T$

Detailed derivation



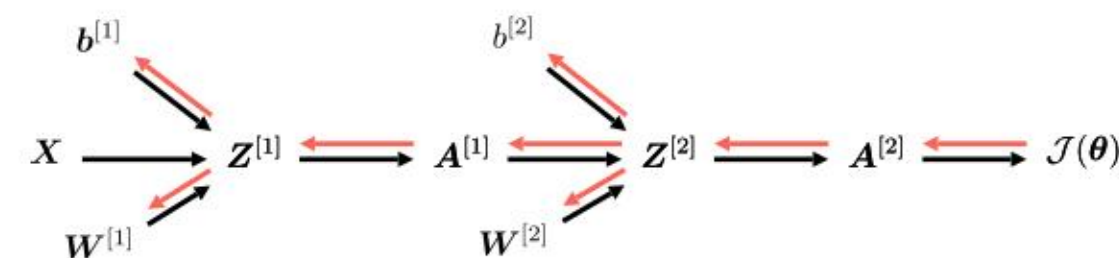
1. To obtain $d\mathbf{A}^{[1]} = \partial\mathcal{J}/\partial\mathbf{A}^{[1]}$, we use the chain rule and treat $\mathcal{J}$ as a function of $\mathbf{Z}^{[2]}$,
and $\mathbf{Z}^{[2]}$ as a function of $\mathbf{A}^{[1]}$

2. Since $\mathbf{A}^{[1]} \in \mathbb{R}^{n \times d^{[1]}}$ is a matrix, not a column one, we have

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

- $\partial z_i^{[2]}/\partial \mathbf{A}^{[1]} \in \mathbb{R}^{n \times d^{[1]}}$ has the same dimension with $\mathbf{A}^{[1]}$
- $z_i^{[2]}$ is a only function of the i th row of $\mathbf{A}^{[1]}$: $z_i^{[2]} = b^{[2]} + \mathbf{W}^{[2]} \cdot \mathbf{a}_i^{[1]}$
- Thus, $\partial z_i^{[2]}/\partial (\mathbf{a}_i^{[1]})^T = \mathbf{W}^{[2]}$ and $\partial z_i^{[2]}/\partial \mathbf{A}^{[1]} = \mathbf{e}_i \cdot \mathbf{W}^{[2]}$
- $\mathbf{e}_i = (0, \dots, 1, \dots, 0)^T$ is a column vector of length n ; all elements are zero but the i th being 1

Detailed derivation

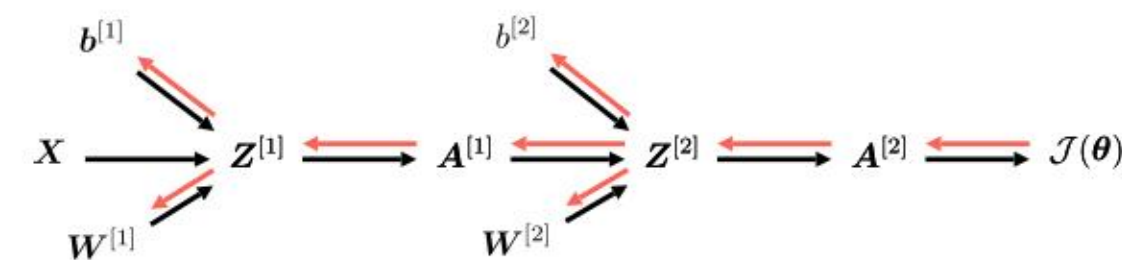


1. To get $d\mathbf{Z}^{[1]} = \partial\mathcal{J}/\partial\mathbf{Z}^{[1]}$, we treat $\mathcal{J}$ as a function of $\mathbf{A}^{[1]}$, and $\mathbf{A}^{[1]}$ as a function of $\mathbf{Z}^{[1]}$
2. Since $\mathbf{A}^{[1]} \in \mathbb{R}^{n \times d^{[1]}}$, we have

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

- $\partial a_{ij}^{[1]}/\partial \mathbf{Z}^{[1]} \in \mathbb{R}^{n \times d^{[1]}}$ has the same dimension with $\mathbf{Z}^{[1]}$
- $a_{ij}^{[1]}$ is a only function of the (i, j) th element of $\mathbf{Z}^{[1]}$: $a_{ij}^{[1]} = \sigma(z_{ij}^{[1]})$
- Thus, $\partial a_{ij}^{[1]}/\partial z_{ij}^{[1]} = \sigma'(z_{ij}^{[1]})$ and $\partial a_{ij}^{[1]}/\partial \mathbf{Z}^{[1]} = \sigma'(z_{ij}^{[1]}) \cdot \mathbf{e}_i \cdot \mathbf{h}_j^T$
- $\mathbf{e}_i = (0, \dots, 1, \dots, 0)^T$ is a column vector of length n ; all elements are zero but the i th being 1
- $\mathbf{h}_j = (0, \dots, 1, \dots, 0)^T$ is a column vector of length $d^{[1]}$; all are zero but the j th being 1

Detailed derivation

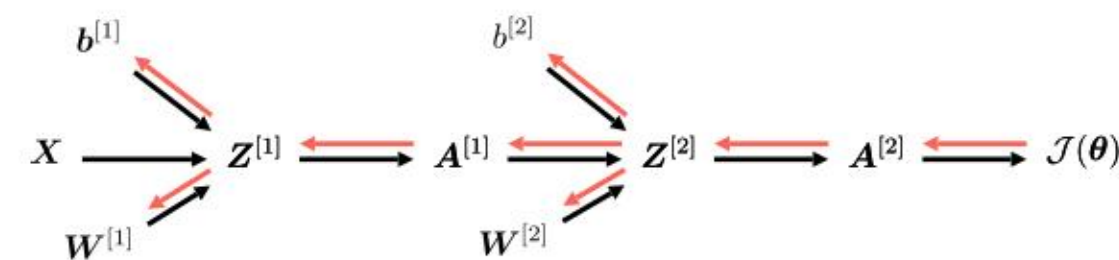


1. To obtain $d\mathbf{b}^{[1]} = \partial\mathcal{J}/\partial\mathbf{b}^{[1]}$, we treat $\mathcal{J}$ as a function of $\mathbf{Z}^{[1]} = (z_1^{[1]}, \dots, z_n^{[1]})^T$,
and $\mathbf{Z}^{[1]}$ as a function of $\mathbf{b}^{[1]}$
2. Since $\mathbf{b}^{[1]} \in \mathbb{R}^{d^{[1]} \times 1}$ is a column vector, we have

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

- Each row of $\mathbf{Z}^{[1]}$ is a function of $\mathbf{b}^{[1]}$: $(z_i^{[1]})^T = (\mathbf{b}^{[1]})^T + \mathbf{x}_i^T \cdot (\mathbf{W}^{[1]})^T$
- Thus, $\partial(z_i^{[1]})^T / \partial \mathbf{b}^{[1]} = \mathbf{I}$
- $\mathbf{I}$ is an identity matrix of dimension $d^{[1]} \times d^{[1]}$

Detailed derivation

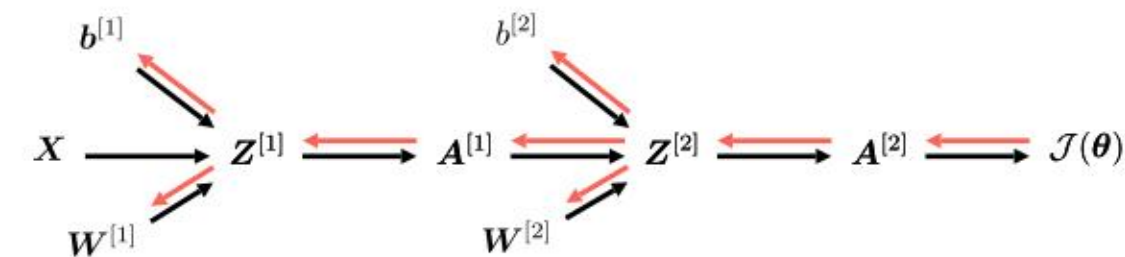


1. To obtain $d\mathbf{W}^{[1]} = \partial\mathcal{J}/\partial\mathbf{W}^{[1]}$, we treat $\mathcal{J}$ as a function of $\mathbf{Z}^{[1]} = (z_{ij}^{[1]})$, and $\mathbf{Z}^{[1]}$ as a function of $\mathbf{W}^{[1]}$
2. Since $\mathbf{W}^{[1]} \in \mathbb{R}^{d^{[1]} \times d^{[0]}}$ is a matrix, we have

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

- $z_{ij}^{[1]}$ is only a function of the j th row of $\mathbf{W}^{[1]} = (\mathbf{w}_1^{[1]}, \dots, \mathbf{w}_{d^{[1]}}^{[1]})^T$: $z_{ij}^{[1]} = b_j^{[1]} + (\mathbf{w}_j^{[1]})^T \cdot \mathbf{x}_i$
- Thus, $\partial z_{ij}^{[1]} / \partial \mathbf{W}^{[1]} = \mathbf{h}_j \cdot \mathbf{x}_i^T$
- $\mathbf{h}_j = (0, \dots, 1, \dots, 0)^T$ is a column vector of length $d^{[1]}$; all are zero but the j th being 1

Remarks



1. For convenience, let $\mathbf{A}^{[0]} = \mathbf{X}$
2. Forward propagation and backpropagation can be coded up using for-loops
 - **Forward propagation:** Based on the current model parameters, $\mathbf{A}^{[l-1]}$ is the “input” for $\mathbf{Z}^{[l]}$ and $\mathbf{A}^{[l]}$
 - **Backpropagation:** $d\mathbf{A}^{[l]}$ is the “input” to obtain $dbZ^{[l]}$, $db^{[l]}$, $d\mathbf{W}^{[l]}$, $d\mathbf{Z}^{[l-1]}$ and $d\mathbf{A}^{[l-1]}$

(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

Summary

1. Extend the single-hidden-layer neural network from single-sample computation to vectorized training with multiple training samples.
 - Data: Stack features row-wisely to obtain a design matrix $\mathbf{X}$
 - Forward propagation: Use current parameters to obtain cost, and broadcasting is applied to improve computation efficiency
 - Backpropagation: Obtain average gradients by a chain rule and vectorization
 - Gradient descent algorithm