



AI and Deep Learning

2-1 Neural Network with One Hidden Layer I

(one training example)

Zhonglei Wang

WISE, SOE, CHOW

XMU

(Version v1)

Contents

1. Revisit logistic regression

2. Forward propagation

3. Backpropagation

4. (Batch) Gradient descent algorithm

Background

1. Illustrate basic concepts using **only ONE** training example $(\mathbf{x}, y)$ ($y \in \{0, 1\}$)
2. A neural network can be viewed as a function of features $\mathbf{x}$ with parameter $\boldsymbol{\theta}$
3. To obtain an estimator of the model parameter $\boldsymbol{\theta}$, we use a (batch) gradient descent algorithm
4. An essential step is

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

- $\boldsymbol{\theta}^{(t)}$: Current model parameter
- **What are the cost function and its partial derivatives?**

Dimension of partial derivatives

1. Consider $\mathbf{f}(\mathbf{x})$,

- $\mathbf{f}(\mathbf{x}) = (f_1(\mathbf{x}), \dots, f_k(\mathbf{x}))^T$: Column vector of length k
- $\mathbf{x} = (x_1, \dots, x_m)^T$: Column vector of length m

2. In this class, we define

$$\frac{\partial \mathbf{f}^T}{\partial \mathbf{x}} = \left(\frac{\partial f_1}{\partial \mathbf{x}}, \dots, \frac{\partial f_k}{\partial \mathbf{x}} \right) = \begin{pmatrix} \partial f_1(\mathbf{x}) / \partial x_1 & \dots & \partial f_k(\mathbf{x}) / \partial x_1 \\ \partial f_1(\mathbf{x}) / \partial x_2 & \dots & \partial f_k(\mathbf{x}) / \partial x_2 \\ \vdots & \ddots & \vdots \\ \partial f_1(\mathbf{x}) / \partial x_m & \dots & \partial f_k(\mathbf{x}) / \partial x_m \end{pmatrix} \in \mathbb{R}^{m \times k}$$

- To make the dimension clear, we use a transpose in the numerator or denominator
- Make $\boldsymbol{\theta}^{(t+1)} = \boldsymbol{\theta}^{(t)} - \alpha \cdot \partial \mathcal{J} / \partial \boldsymbol{\theta}(\boldsymbol{\theta}^{(t)})$ mathematically rigorous ($k = 1$)

Dimension of partial derivatives

1. In advanced math classes, we define derivatives to facilitate linear approximation ($\epsilon \approx 0$)

$$\mathbf{f}(\mathbf{x} + \boldsymbol{\epsilon}) = \begin{pmatrix} f_1(\mathbf{x} + \boldsymbol{\epsilon}) \\ f_2(\mathbf{x} + \boldsymbol{\epsilon}) \\ \vdots \\ f_k(\mathbf{x} + \boldsymbol{\epsilon}) \end{pmatrix} \approx \begin{pmatrix} f_1(\mathbf{x}) \\ f_2(\mathbf{x}) \\ \vdots \\ f_k(\mathbf{x}) \end{pmatrix} + \frac{\partial \mathbf{f}}{\partial \mathbf{x}} \cdot \boldsymbol{\epsilon}$$

- $\partial \mathbf{f} / \partial \mathbf{x}$: Transpose of the one we defined in the previous slide

$$\frac{\partial \mathbf{f}}{\partial \mathbf{x}} = \begin{pmatrix} \partial f_{\mathbf{1}}(\mathbf{x}) / \partial x_{\mathbf{1}} & \dots & \partial f_{\mathbf{1}}(\mathbf{x}) / \partial x_{\mathbf{m}} \\ \partial f_{\mathbf{2}}(\mathbf{x}) / \partial x_{\mathbf{1}} & \dots & \partial f_{\mathbf{2}}(\mathbf{x}) / \partial x_{\mathbf{m}} \\ \vdots & \ddots & \vdots \\ \partial f_{\mathbf{k}}(\mathbf{x}) / \partial x_{\mathbf{1}} & \dots & \partial f_{\mathbf{k}}(\mathbf{x}) / \partial x_{\mathbf{m}} \end{pmatrix} \in \mathbb{R}^{k \times m}$$

Dimension of partial derivatives

1. For the convenience of our discussion, we DO NOT adopt the form of gradient in those advanced mathematics

Examples

1. Example 1: Consider $f(\mathbf{x})$

- $\mathbf{x} = (x_1, \dots, x_m)^T$: Column vector
- $f(\mathbf{x})$: Univariate function

2. Then, we have

$$\frac{\partial f}{\partial \mathbf{x}} = \left(\frac{\partial f}{\partial x_1}, \dots, \frac{\partial f}{\partial x_m} \right)^T$$

3. $\partial f / \partial \mathbf{x}$ has the same dimension with $\mathbf{x}$

Examples

1. Example 2: Consider $f(\boldsymbol{x})$

- $\boldsymbol{x} = (x_1, \dots, x_m) : \text{Row vector}$
- $f(\boldsymbol{x}) : \text{Univariate function}$

2. Then, we have

$$\frac{\partial f}{\partial \boldsymbol{x}} = \left(\frac{\partial f}{\partial x_1}, \dots, \frac{\partial f}{\partial x_m} \right)$$

3. $\partial f / \partial \boldsymbol{x}$ has the same dimension with $\boldsymbol{x}$

Examples

1. Consider $f(\mathbf{x}, \mathbf{y}) = \mathbf{x}^T \cdot \mathbf{y} = \sum_{j=1}^m x_j y_j$

- $\mathbf{x} = (x_1, \dots, x_m)^T$

- $\mathbf{y} = (y_1, \dots, y_m)^T$

2. Then, we have

$$\begin{aligned} \frac{\partial f}{\partial \mathbf{x}} &= \mathbf{y}, & \frac{\partial f}{\partial \mathbf{y}} &= \mathbf{x}, \\ \frac{\partial f}{\partial \mathbf{x}^T} &= \mathbf{y}^T, & \frac{\partial f}{\partial \mathbf{y}^T} &= \mathbf{x}^T \end{aligned}$$

Examples

1. Example 3: Consider $f(\mathbf{X})$

- $\mathbf{X} = (x_{ij}) \in \mathbb{R}^{m \times k} : m \times k$ matrix
- $f(\mathbf{X})$: Univariate function

2. Then, we have

$$\frac{\partial f}{\partial \mathbf{X}} = \left(\frac{\partial f}{\partial x_{ij}} \right)$$

3. $\partial f / \partial \mathbf{X}$ has the same dimension with $\mathbf{X}$

Examples

1. Example 4 (chain rule): Consider

$$f(\boldsymbol{x}) = g \circ \boldsymbol{h}(\boldsymbol{x})$$

- $\boldsymbol{x}$: Scalar, column vector, row vector or matrix
- $f(\boldsymbol{x})$: Univariate function
- $\boldsymbol{h}(\boldsymbol{x}) = (h_1(\boldsymbol{x}), \dots, h_m(\boldsymbol{x}))^T$

2. Then, we have

$$\frac{\partial f}{\partial \boldsymbol{x}} = \sum_{j=1}^m \frac{\partial h_j}{\partial \boldsymbol{x}} \cdot \frac{\partial g}{\partial h_j}$$

Examples

1. If $\mathbf{x}$ is a scalar or a column vector, we further have

$$\frac{\partial f}{\partial \mathbf{x}} = \sum_{j=1}^m \frac{\partial h_j}{\partial \mathbf{x}} \cdot \frac{\partial g}{\partial h_j} = \left(\frac{\partial h_1}{\partial \mathbf{x}}, \dots, \frac{\partial h_m}{\partial \mathbf{x}} \right) \cdot \begin{pmatrix} \frac{\partial g}{\partial h_1} \\ \vdots \\ \frac{\partial g}{\partial h_m} \end{pmatrix} = \frac{\partial \mathbf{h}^T}{\partial \mathbf{x}} \cdot \frac{\partial g}{\partial \mathbf{h}}$$

- $\partial g / \partial \mathbf{h} = (\partial g / \partial h_1, \dots, \partial g / \partial h_m)^T$
- $\partial \mathbf{h}^T / \partial \mathbf{x} = (\partial h_1 / \partial \mathbf{x}, \dots, \partial h_m / \partial \mathbf{x})$

2. That is, the calculation can be vectorized

3. However, the essential part is to get $\partial h_j / \partial \mathbf{x}$

4. Vectorization may fail if $\mathbf{x}$ is a matrix

5. **Question:** How about the case when $\mathbf{x}$ is a row vector?

Goal

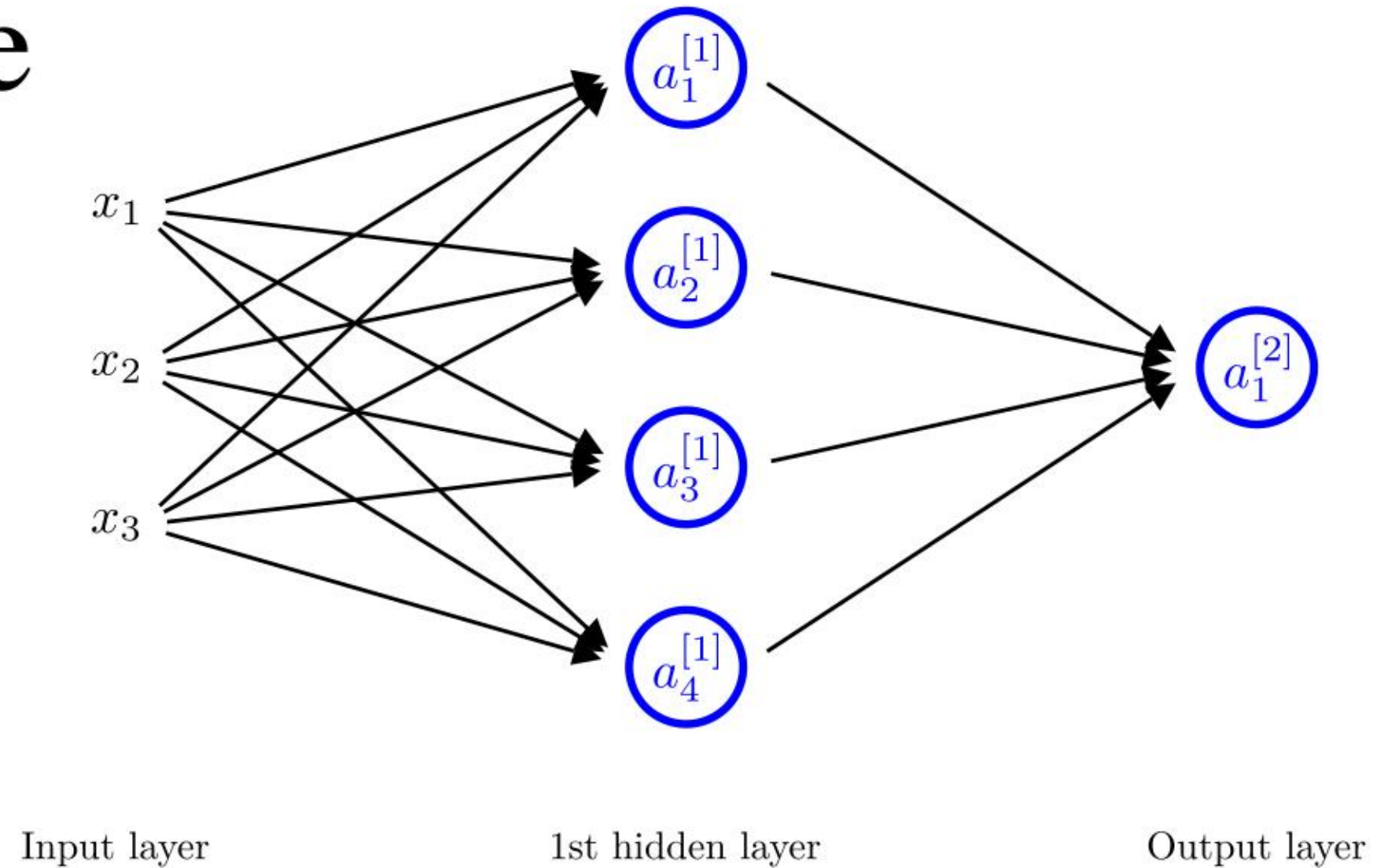
1. To obtain $\partial \mathcal{J}(\boldsymbol{\theta}^{(t)}) / \partial \boldsymbol{\theta}$, we introduce
 - **Forward propagation**: Based on the current parameter, calculate the “activated” values as well as the cost function
 - **Backpropagation**: Based on “those” values, obtain partial derivatives
2. We use a toy neural network to introduce forward propagation and backpropagation

Revisit logistic regression

1. The model is too **simple** in practice
2. Why not use more “circles”?

Neural network example

1. Assume $\mathbf{x} = (x_1, x_2, x_3)^T$



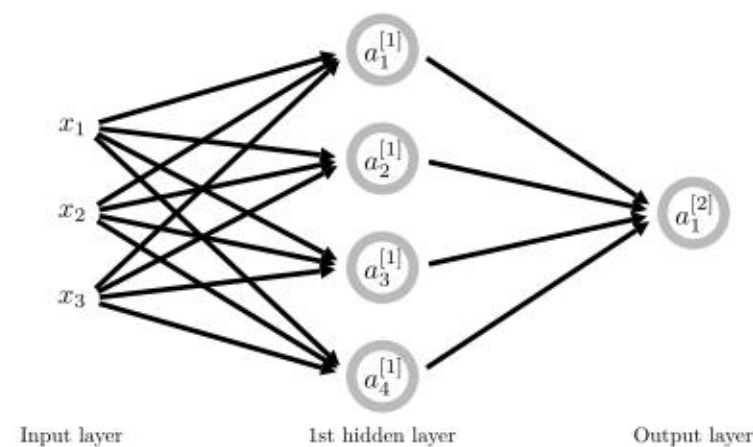
2. Each **circle** represents two operations:

- **Linear** transformation with two model parameters, including a bias and a weight
- **Activation** (nonlinear transformation) without model parameters

Neural network example

1. Remarks

- Different model parameters are used for different “circles” to extract different information
- In other words, we “construct” a set of “new” features in the hidden layer
- For the output layer, the hidden layer can be regarded as its “input” layer
- Compared with a logistic regression model, this neural network uses one hidden layer to extract more information



Calculation

Parameters

$$z_1^{[1]} = b_1^{[1]} + \mathbf{x}^T \cdot \mathbf{w}_1^{[1]}, \quad a_1^{[1]} = \sigma(z_1^{[1]})$$

$$b_1^{[1]} \quad \mathbf{w}_1^{[1]}$$

$$z_2^{[1]} = b_2^{[1]} + \mathbf{x}^T \cdot \mathbf{w}_2^{[1]}, \quad a_2^{[1]} = \sigma(z_2^{[1]})$$

$$b_2^{[1]} \quad \mathbf{w}_2^{[1]}$$

$$z_3^{[1]} = b_3^{[1]} + \mathbf{x}^T \cdot \mathbf{w}_3^{[1]}, \quad a_3^{[1]} = \sigma(z_3^{[1]})$$

$$b_3^{[1]} \quad \mathbf{w}_3^{[1]}$$

$$z_4^{[1]} = b_4^{[1]} + \mathbf{x}^T \cdot \mathbf{w}_4^{[1]}, \quad a_4^{[1]} = \sigma(z_4^{[1]})$$

$$b_4^{[1]} \quad \mathbf{w}_4^{[1]}$$

$$z_1^{[2]} = b_1^{[2]} + (\mathbf{a}^{[1]})^T \cdot \mathbf{w}_1^{[2]}, \quad a_1^{[2]} = \sigma(z_1^{[2]})$$

$$b_1^{[2]} \quad \mathbf{w}_1^{[2]}$$

Vectorization

1. Denote

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

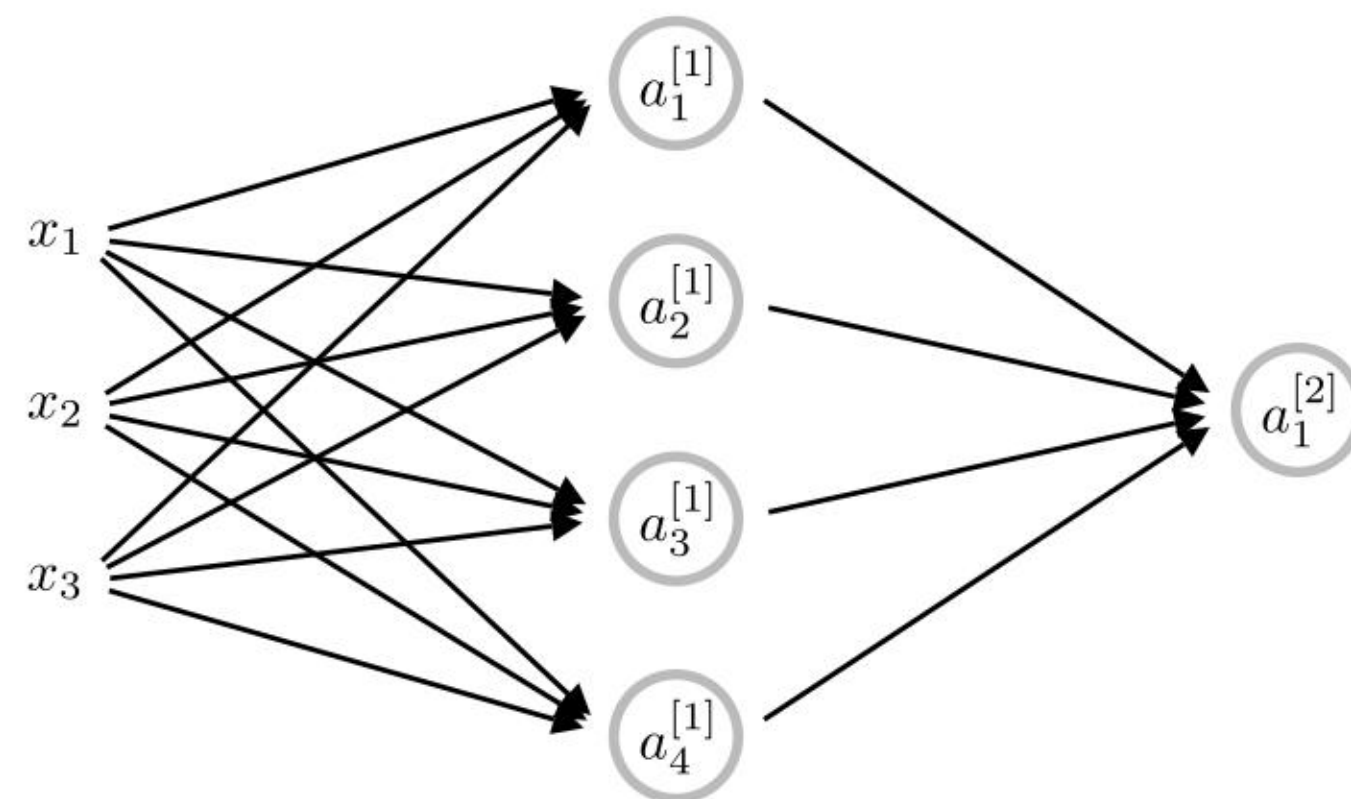
2. Model parameters

- $\{(\mathbf{b}^{[l]}, \mathbf{W}^{[l]}) : l = 1, \dots, L\}$

Vectorization

1. For the previous example, we have

- $d^{[0]} = 3, d^{[1]} = 4, d^{[2]} = 1$
- $\mathbf{W}^{[1]} = (\mathbf{w}_1^{[1]}, \mathbf{w}_2^{[1]}, \mathbf{w}_3^{[1]}, \mathbf{w}_4^{[1]})^T \in \mathbb{R}^{4 \times 3}, \mathbf{W}^{[2]} = (\mathbf{w}_1^{[2]})^T \in \mathbb{R}^{1 \times 4}$
- $\mathbf{b}^{[1]} = (b_1^{[1]}, b_2^{[1]}, b_3^{[1]}, b_4^{[1]})^T \in \mathbb{R}^{4 \times 1}, b^{[2]} \in \mathbb{R}$



Forward propagation

1. Forward propagation obtains “activated” values using CURRENT parameters
2. Assume the **current model parameters** are $\mathbf{b}^{[1]}, \mathbf{W}^{[1]}, b^{[2]}, \mathbf{W}^{[2]}$
3. The calculation can be simplified as

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

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

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

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

$$\mathbf{z}^{[1]} = \begin{pmatrix} b_1^{[1]} \\ b_2^{[1]} \\ b_3^{[1]} \\ b_4^{[1]} \end{pmatrix} + \begin{pmatrix} (\mathbf{w}_1^{[1]})^T \\ (\mathbf{w}_2^{[1]})^T \\ (\mathbf{w}_3^{[1]})^T \\ (\mathbf{w}_4^{[1]})^T \end{pmatrix} \cdot \mathbf{x} = \begin{pmatrix} b_1^{[1]} + (\mathbf{w}_1^{[1]})^T \cdot \mathbf{x} \\ b_2^{[1]} + (\mathbf{w}_2^{[1]})^T \cdot \mathbf{x} \\ b_3^{[1]} + (\mathbf{w}_3^{[1]})^T \cdot \mathbf{x} \\ b_4^{[1]} + (\mathbf{w}_4^{[1]})^T \cdot \mathbf{x} \end{pmatrix}$$

Forward propagation

1. Forward propagation obtains “activated” values using CURRENT parameters
2. Assume the **current model parameters** are $\mathbf{b}^{[1]}, \mathbf{W}^{[1]}, b^{[2]}, \mathbf{W}^{[2]}$
3. The calculation can be simplified as

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

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

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

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

$$\mathbf{a}^{[1]} = \begin{pmatrix} \sigma(z_1^{[1]}) \\ \sigma(z_2^{[1]}) \\ \sigma(z_3^{[1]}) \\ \sigma(z_4^{[1]}) \end{pmatrix}$$

Forward propagation

1. Consider a binary response ($y \in \{0, 1\}$)
2. $a^{[2]}$ is the estimated value for $P(y = 1 \mid \mathbf{x})$

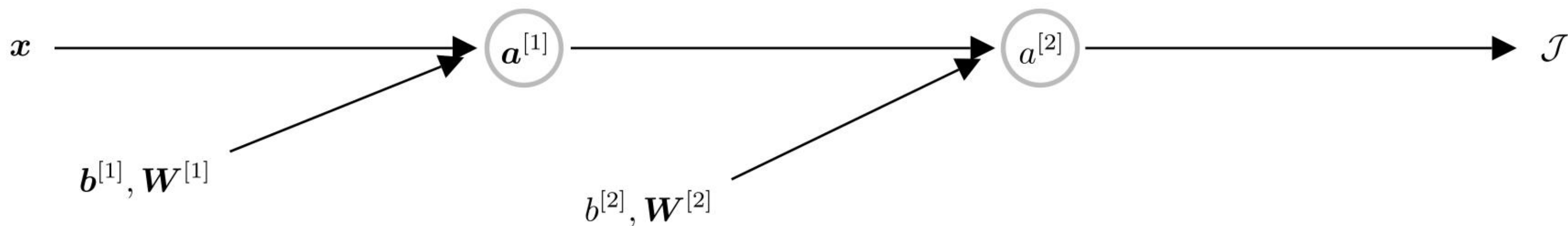
3. Thus, we consider a loss function

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

4. **Question:** What would $a^{[2]}$ and $\mathcal{L}$ be if $y \in \mathbb{R}$?
5. **Question:** What about the case when the range of labels is nonnegative?

Forward propagation

1. Visualize the calculation details



Forward propagation

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

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

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

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

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

Remarks

1. When obtaining $\mathbf{a}^{[1]}$ in the first layer, $\mathbf{x}$ is the input
2. When obtaining $\mathbf{a}^{[2]}$ in the second layer, $\mathbf{a}^{[1]}$ can be treated as the input
3. That is, when obtaining activated values in the l th layer, those in the previous layer can be treated as the input
4. This is the general case for forward propagation, so forward propagation can be coded by a for-loop

Backpropagation

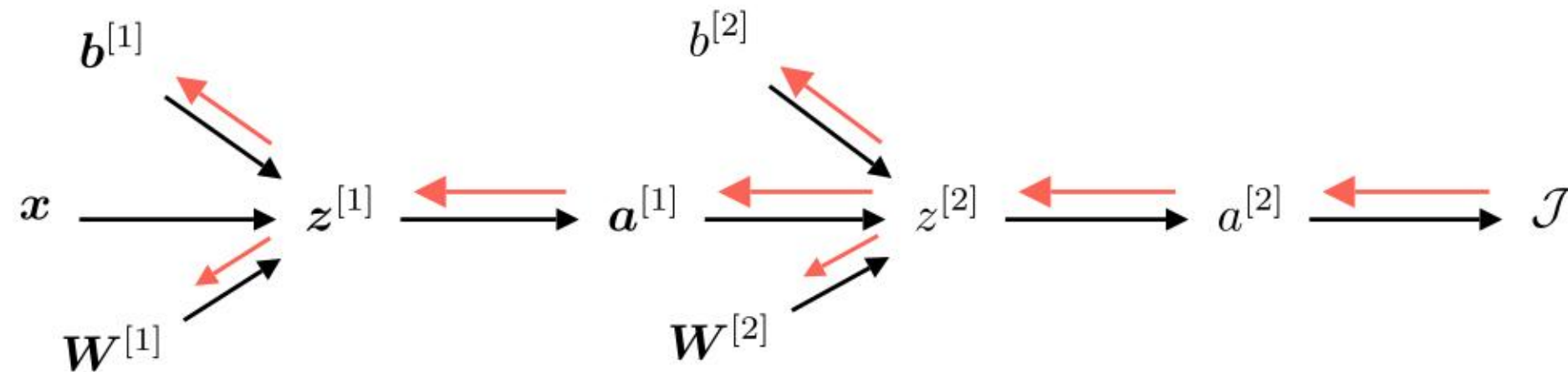
1. Based on current model parameters, backpropagation is to obtain

$$\frac{\partial \mathcal{J}}{\partial \mathbf{b}^{[l]}}, \quad \frac{\partial \mathcal{J}}{\partial \mathbf{W}^{[l]}} \quad (l = 1, \dots, L)$$

2. Core technique: **Chain rule**

- Denote $f(\mathbf{x}) = g \circ h(\mathbf{x})$
- Then, we have

$$\frac{\partial f}{\partial \mathbf{x}} = \frac{\partial h}{\partial \mathbf{x}} \cdot \frac{\partial g}{\partial h}$$



$$\frac{\partial \mathcal{J}}{\partial b^{[2]}} = a^{[2]} - y$$

$$\frac{\partial \mathcal{J}}{\partial \mathbf{W}^{[2]}} = (a^{[2]} - y) \cdot (\mathbf{a}^{[1]})^T$$

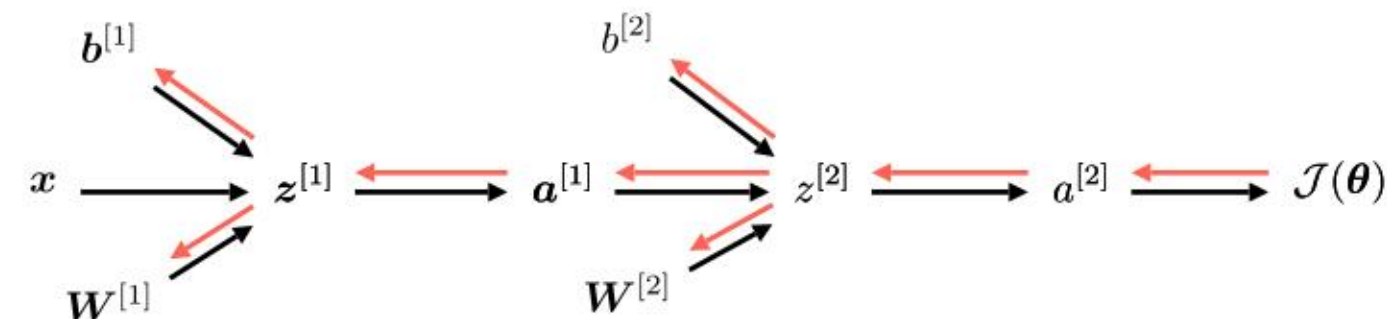
$$\frac{\partial \mathcal{J}}{\partial \mathbf{b}^{[1]}} = (a^{[2]} - y) \cdot \mathbf{D} \cdot (\mathbf{W}^{[2]})^T$$

$$\frac{\partial \mathcal{J}}{\partial \mathbf{W}^{[1]}} = (a^{[2]} - y) \cdot \mathbf{D} \cdot (\mathbf{W}^{[2]})^T \cdot \mathbf{x}^T$$

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

2. We only need to cache $\mathbf{a}^{[1]}$ and $a^{[2]}$

Remarks

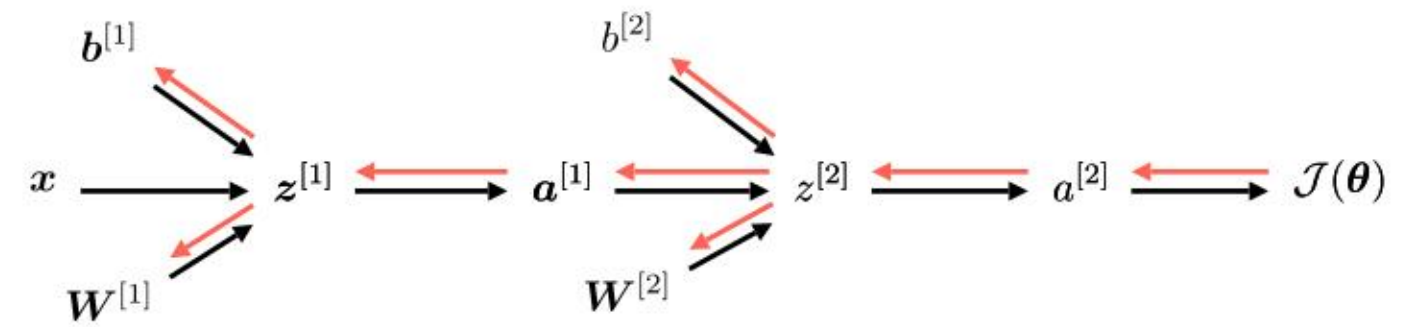


1. The form of $\partial \mathcal{J} / \partial a^{[2]}$ is determined by the form of the cost function $\mathcal{J}$
 - We have shown the result for **binary classification problems**
 - Derive $\partial \mathcal{J} / \partial a^{[2]}$ if $\mathcal{J} = (y - a^{[2]})^2 / 2$ for **regression problems** by yourself
2. Once we obtain $\partial \mathcal{J} / \partial a^{[2]}$, we use a chain rule to get

$$\frac{\partial \mathcal{J}}{\partial z^{[2]}} = \frac{\partial a^{[2]}}{\partial z^{[2]}} \cdot \frac{\partial \mathcal{J}}{\partial a^{[2]}}$$

- Result for $\partial a^{[2]} / \partial z^{[2]}$ is determined by the form of activation function to get $a^{[2]}$
- $\partial \mathcal{J} / \partial z^{[2]} = a^{[2]} - y$ for binary classification problems

Remarks



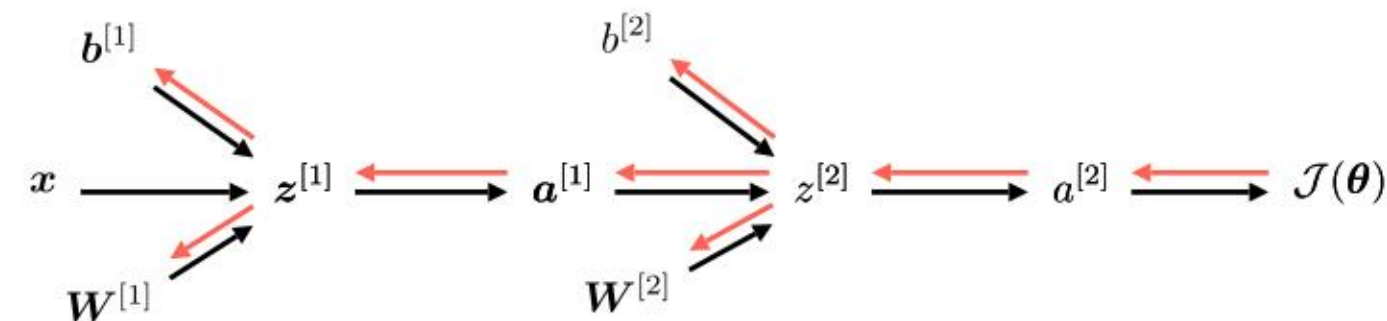
1. Once we obtain $\partial \mathcal{J} / \partial z^{[2]}$, we can easily get

$$\frac{\partial \mathcal{J}}{\partial b^{[2]}} = \frac{\partial z^{[2]}}{\partial b^{[2]}} \cdot \frac{\partial \mathcal{J}}{\partial z^{[2]}} = \frac{\partial \mathcal{J}}{\partial z^{[2]}}$$

$$\frac{\partial \mathcal{J}}{\partial \mathbf{W}^{[2]}} = \frac{\partial z^{[2]}}{\partial \mathbf{W}^{[2]}} \cdot \frac{\partial \mathcal{J}}{\partial z^{[2]}} = (\mathbf{a}^{[1]})^T \cdot \frac{\partial \mathcal{J}}{\partial z^{[2]}}$$

2. That is, to obtain derivatives with respect to the model parameters in the second layer, it is desirable to obtain $\partial \mathcal{J} / \partial z^{[2]}$
3. Similar argument applies to the case for the first layer

Remarks



1. Once we have $\partial \mathcal{J} / \partial \mathbf{a}^{[1]}$, we can easily get

$$\frac{\partial \mathcal{J}}{\partial \mathbf{z}^{[1]}} = \frac{\partial (\mathbf{a}^{[1]})^T}{\partial \mathbf{z}^{[1]}} \cdot \frac{\partial \mathcal{J}}{\partial \mathbf{a}^{[1]}} = \mathbf{D} \cdot \frac{\partial \mathcal{J}}{\partial \mathbf{a}^{[1]}},$$

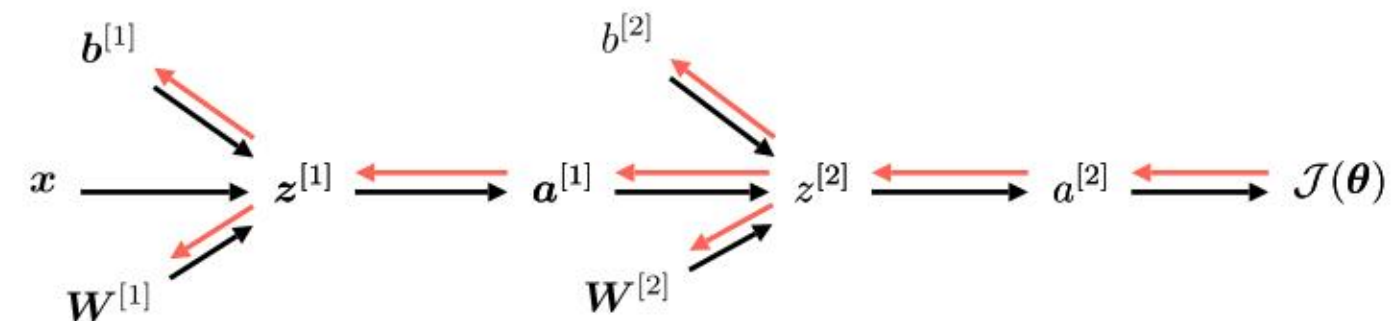
$$\frac{\partial \mathcal{J}}{\partial \mathbf{b}^{[1]}} = \frac{\partial (\mathbf{z}^{[1]})^T}{\partial \mathbf{b}^{[1]}} \cdot \frac{\partial \mathcal{J}}{\partial \mathbf{z}^{[1]}} = \frac{\partial \mathcal{J}}{\partial \mathbf{z}^{[1]}},$$

$$\frac{\partial \mathcal{J}}{\partial \mathbf{W}^{[1]}} = \sum_{k=1}^{d^{[1]}} \frac{\partial z_k^{[1]}}{\partial \mathbf{W}^{[1]}} \cdot \frac{\partial \mathcal{J}}{\partial z_k^{[1]}} = \frac{\partial \mathcal{J}}{\partial \mathbf{z}^{[1]}} \cdot \mathbf{x}^T$$

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

2. It remains to show how to obtain $\partial \mathcal{J} / \partial \mathbf{a}^{[1]}$ from $\partial \mathcal{J} / \partial \mathbf{z}^{[2]}$

Remarks



1. By the computation graph, we only need to consider ONE more step:

$$\frac{\partial \mathcal{J}}{\partial \mathbf{a}^{[1]}} = \frac{\partial z^{[2]}}{\partial \mathbf{a}^{[1]}} \cdot \frac{\partial \mathcal{J}}{\partial z^{[2]}} = (\mathbf{W}^{[2]})^T \cdot (a^{[2]} - y)$$

Remarks

1. To update parameters in the l th layer, it is enough to know $\partial \mathcal{J} / \partial \mathbf{a}^{[l]}$
2. To obtain $\partial \mathcal{J} / \partial \mathbf{a}^{[l-1]}$ from $\partial \mathcal{J} / \partial \mathbf{z}^{[l]}$, we only need to consider the linear transformation in the l th layer
3. Thus, backpropagation can also be coded up using for-loops

(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. Generalized logistic regression to neural networks
2. Introduced how to obtain gradients for the model parameters
 - Forward propagation: Calculation of output and cost function based on the **current** model parameters
 - Backpropagation: Obtain gradients by a chain rule
3. Gradient descent algorithm