

Coursera Machine Learning: Weeks 5 & 6

Week 5: Backpropagation

Review: Neural Networks

- Given a training set $(x^{(i)}, y^{(i)})$, adjust parameters θ so that $p(h_{\theta}(x^{(i)})) = y^{(i)}$

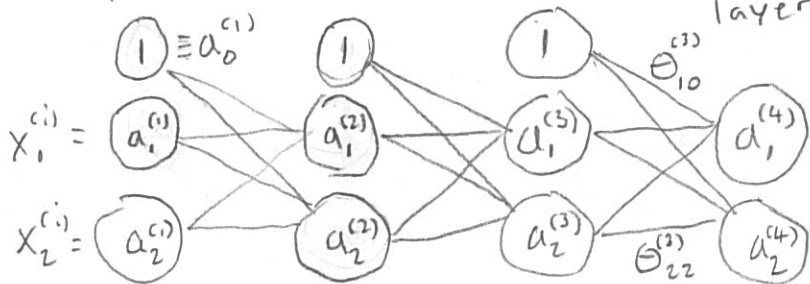
e.g. $x^{(i)}$ is a 20×20 pixel grayscale image: 

$$y^{(i)} = \begin{bmatrix} 0 \\ 0 \\ \vdots \\ 1 \\ 0 \end{bmatrix} \rightarrow \text{labeled as the digit "9"}$$

hidden layers (size 16 for example)

input layer

output layer



weights: $\Theta^{(1)}: 401 \times 16$, $\Theta^{(2)}: 17 \times 16$, $\Theta^{(3)}: 17 \times 10$

$$a^{(4)} = h_{\theta}(x)$$

$$\text{prediction } p = \begin{bmatrix} 0 \\ 0 \\ \vdots \\ 1 \\ 0 \end{bmatrix} \leftarrow \text{index of } \max(h_{\theta}(x))$$

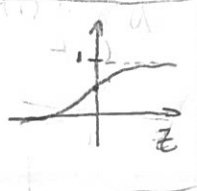
$$x^{(i)} = \begin{bmatrix} x_1 \\ x_2 \\ \vdots \end{bmatrix}, a^{(1)} = \begin{bmatrix} 1 \\ x_1 \\ x_2 \\ \vdots \end{bmatrix} = [1, x^{(i)}]$$

Forward Propagation: $a^{(1)} = [1; x^{(i)}]$, $a^{(2)} = [1; g(\theta^{(1)} a^{(1)})]$,

$$a^{(3)} = [1; g(\theta^{(2)} a^{(2)})]$$

$$a^{(4)} = g(\theta^{(3)} a^{(3)})$$

"sigmoid": $g(z) = \frac{1}{1 + e^{-z}}$



$$\text{Cost: } J(\theta) = -\frac{1}{m} \sum_{i=1}^m \sum_{k=1}^K \left[y_k^{(i)} \log(h_{\theta}(x^{(i)})) + (1 - y_k^{(i)}) \log(1 - h_{\theta}(x^{(i)})) \right]$$

$$= -\frac{1}{m} \sum \sum \left[y_k \log(a_k^{(4)}) + (1 - y_k) \log(1 - a_k^{(4)}) \right]$$

Back propagation : Goal : Adjust $\Theta_{ij}^{(k)}$ to lower cost ($J(\theta)$)
by gradient descent $\Theta = \Theta - \alpha \frac{\partial J}{\partial \Theta}$

Problem: J seems to depend only on output layer $a^{(4)}$

Solution: Propagate cost back, layer by layer

Code :

for n in (# of iterations):

Set $\Delta_{ij}^{(1)} = 0$ for all (i, j)

for (x, y) in training data:

forward propagate x (compute $a^{(1)}$ then $a^{(2)}$ etc)

$$\delta^{(4)} = a^{(4)} - y^{(4)} \quad \text{"error in layer 4"}$$

$$\delta^{(3)} = \Theta^{(3)T} \delta^{(4)} * g'(\Theta^{(2)} a^{(2)}) \quad \text{"error in layer 3"}$$

$$S^{(2)} = \Theta^{(2)T} S^{(3)} \cdot g(\Theta^{(1)} a^{(1)})$$

$$\text{end } \Delta_{ij}^{(l)} = \Delta_{ij}^{(l)} + \delta_i^{(l+1)} a_j^{(l)} \quad \text{for all } (l, i, j)$$

$$\frac{\partial J}{\partial \theta_{ij}^{(l)}} \equiv \frac{\Delta_{ij}^{(l)}}{m} \quad \text{for all } (l, i, j)$$

$$\Theta_{ij}^{(l)} = \Theta_{ij}^{(l)} - \lambda \frac{\partial J}{\partial \Theta_{ij}^{(l)}} \quad \text{for all } (l, i, j)$$

end

$$g'(z) = g(z)(1 - g(z))$$

Derivation: (with calculus i.e. the chain rule)

$$\begin{aligned} \frac{\partial J}{\partial \theta^{(3)}} &= \sum_{l=4}^L \frac{\partial J}{\partial a^{(4)}} \frac{\partial a^{(4)}}{\partial \theta^{(3)}} \\ &= \sum_{l=4}^L \left(\frac{a^{(4)} - y}{a^{(4)}(1 - a^{(4)})} \right) \left(a^{(3)} \frac{a^{(4)}}{(1 - a^{(4)})} \right) \\ &= \sum_{l=4}^L \frac{1}{n} \delta^{(4)} a^{(3)} \end{aligned}$$

$$\begin{aligned}
 \frac{\partial J}{\partial \theta^{(2)}} &= \sum_{l=3,4} \frac{\partial J}{\partial a^{(l)}} \frac{\partial a^{(l)}}{\partial a^{(3)}} \frac{\partial a^{(3)}}{\partial \theta^{(2)}} \\
 &= \sum_m \frac{1}{m} \left(\frac{a^{(4)} - y}{a^{(4)}(1-a^{(4)})} \right) \left(\theta^{(3)} a^{(4)} (1-a^{(4)}) \right) \\
 &\quad \times \left(a^{(2)} g'(\theta^{(2)} a^{(2)}) \right) \\
 &= \sum_m \frac{1}{m} \delta^{(4)} \theta^{(3)} a^{(2)} g'(\theta^{(2)} a^{(2)}) \\
 &= \sum_m \frac{1}{m} \delta^{(3)} a^{(2)}
 \end{aligned}$$

$$\begin{aligned}
 \frac{\partial J}{\partial \theta^{(1)}} &= \sum_{l=2,3,4} \frac{\partial J}{\partial a^{(l)}} \frac{\partial a^{(l)}}{\partial a^{(3)}} \frac{\partial a^{(3)}}{\partial a^{(2)}} \frac{\partial a^{(2)}}{\partial \theta^{(1)}} \\
 &= \vdots
 \end{aligned}$$

$$= \sum_m \frac{1}{m} \delta^{(2)} a^{(1)}$$

$$\begin{aligned}
 \frac{\partial a^{(4)}}{\partial a^{(3)}} &= \frac{\partial}{\partial a^{(3)}} g(\theta^{(3)} a^{(3)}) = \theta^{(3)} g'(\theta^{(3)} a^{(3)}) \\
 &= \theta^{(3)} a^{(4)} (1-a^{(4)}) \\
 \frac{\partial a^{(3)}}{\partial \theta^{(2)}} &= \frac{\partial}{\partial \theta^{(2)}} g(\theta^{(2)} a^{(2)}) = a^{(2)} g'(\theta^{(2)} a^{(2)})
 \end{aligned}$$

Useful Resource :

YouTube: 3Blue1Brown "Neural Networks"

$$\begin{aligned}\frac{\partial a_k^{(4)}}{\partial \theta_{ij}^{(3)}} &= \frac{\partial}{\partial \theta_{ij}^{(3)}} \left(g \left(\sum_n \theta_{kn}^{(3)} a_n^{(3)} \right) \right) \\ &= a_j^{(3)} g' \left(\sum_n \theta_{kn}^{(3)} a_n^{(3)} \right) \delta_{ik} \\ &= a_j^{(3)} a_k^{(4)} (1 - a_k^{(4)}) \delta_{ik}\end{aligned}$$

$$\begin{aligned}\frac{\partial J}{\partial \theta_{ij}^{(3)}} &= \sum_k \frac{\partial J}{\partial a_k^{(4)}} \frac{\partial a_k^{(4)}}{\partial \theta_{ij}^{(3)}} \\ &= \sum_k \frac{(a_k - y_k)}{a_k(1-a_k)} a_k(1-a_k) \delta_{ik} a_j^{(3)} \\ &= (a_i - y_i) a_j^{(3)} \\ &= \delta_i^{(4)} a_j^{(3)}\end{aligned}$$

$$\begin{aligned}\frac{\partial J}{\partial \theta_{ij}^{(2)}} &= \sum_k \sum_p \frac{\partial J}{\partial a_k^{(4)}} \frac{\partial a_k^{(4)}}{\partial a_p^{(3)}} \frac{\partial a_p^{(3)}}{\partial \theta_{ij}^{(2)}} \\ &= \sum_k \sum_p \frac{(a_k - y_k)}{a_k(1-a_k)} \theta_{kp}^{(3)} a_k^{(4)} (1 - a_k^{(4)}) a_j^{(2)} \underbrace{a_p^{(3)} (1 - a_p^{(3)})}_{g'(\sum_n \theta_{kn}^{(2)} a_n^{(2)})} \delta_{ip} \\ &= \sum_k \delta_{ik}^{(4)} \theta_{ki}^{(3)} g' \left(\sum_n \theta_{in}^{(2)} a_n^{(2)} \right) a_j^{(2)} \\ &= a_j^{(2)} \delta_i^{(3)}\end{aligned}$$

$$\delta_i^{(3)} = \sum_k \theta_{ki}^{(3)} \delta_k^{(4)} g' \left(\sum_n \theta_{in}^{(2)} a_n^{(2)} \right)$$

$$\vec{\delta}^{(3)} = \Theta^{(3)T} \vec{\delta}^{(4)} * g'(\Theta^{(2)} \vec{a}^{(2)})$$

$$\begin{aligned}\frac{\partial J}{\partial a_k^{(4)}} &= -y_k + \frac{(1-y_k)(-1)}{1-a_k} \\ &= -y_k(1-a_k) + (1-y_k)a_k \\ &= \frac{a_k - y_k}{a_k(1-a_k)}\end{aligned}$$

$$\begin{aligned}\frac{\partial a_k^{(4)}}{\partial a_p^{(3)}} &= \frac{\partial}{\partial a_p^{(3)}} g \left(\sum_n \theta_{kn}^{(3)} a_n^{(3)} \right) \\ &= \theta_{kp}^{(3)} g' \left(\sum_n \theta_{kn}^{(3)} a_n^{(3)} \right) \\ \frac{\partial a_p^{(3)}}{\partial \theta_{ij}^{(2)}} &= a_j^{(2)} a_p^{(3)} (1 - a_p^{(3)}) \delta_{ip}\end{aligned}$$

Week 6: Advice for Applying Machine Learning

Options to Improve an Algorithm

- increase # of training examples (m)
- reduce # of features $x_1, \dots, x_{100} \rightarrow x_1, \dots, x_{50}$
- increase # of features
- add polynomial features (increased d) $(x_1, x_1^2, \dots, x_1^d)$
- decrease λ (regularization parameter)
- increase λ $\Theta = \Theta - \alpha \frac{\partial J}{\partial \Theta} - \lambda \Theta$
- add more layers

Fixes:	
Underfitting (high bias)	Overfitting (high variance)
	✓
	✓
✓	
✓	
✓	
	✓
✓	

Model Selection and Testing

Split data set:

e.g.

$\begin{Bmatrix} x^{(1)}, y^{(1)} \\ \vdots \\ x^{(6)}, y^{(6)} \end{Bmatrix}$ training set $\sim 60\%$

$\begin{Bmatrix} x^{(7)} \\ x^{(8)} \end{Bmatrix}$ cross validation set $\sim 20\%$

$\begin{Bmatrix} x^{(9)} \\ x^{(10)} \end{Bmatrix}$ test set $\sim 20\%$

- use to train algorithm

- algorithm will usually be most successful on this set

- use to choose model parameters (d, λ , etc)

- check if model generalizes

Learning Curves

