



AI and Deep Learning

3-3 Batch Normalization

Zhonglei Wang

WISE, SOE, CHOW

XMU

(Version v1)

Contents

1. Introduction

2. Forward propagation

3. Backpropagation

4. Discussion for dropout and batch normalization

Introduction

1. Before 2015, there existed optimization issues for training deep neural networks
 - Instability for backpropagation
 - Severe sensitivity to parameter initialization
2. For example, there is an issue, called internal covariate shift
 - Minor modification of parameters in shallower layers may change the distribution of inputs for deep layers
 - Parameters of deeper layers focus on “a moving target”
3. Such instability may lead to the following problems
 - Gradient vanishing
 - Gradient exploding

Introduction

1. Whitening is popular in classical statistical learning
 - Use linear transformation to achieve mean 0 and identity covariance matrix
 - Algorithm may converge fast
2. Such whitening is impossible for deep learning
 - Size of the training dataset is usually large
 - We cannot guarantee whitening for each layer
 - Computation cost for whitening is high

Introduction

1. Ioffe and Szegedy (2015) proposed batch normalization (BN)
 - Implemented independently for each feature (associated with a certain neuron):
Does not consider covariance among features
 - Mean and standard error is obtained by the current mini-batch
2. Normalize values **after** linear transformation but **before** activation for EACH neuron
3. BN has the following properties
 - Tries to solve the interclass covariate shift problem
 - Introduce two more parameters to allow for heterogeneity

Forward propagation

1. For a mini-batch with m training examples, the forward propagation for the l th layer is

$$\mathbf{Z}^{[l]} = \mathbf{A}^{[l-1]}(\mathbf{W}^{[l]})^T + (\mathbf{b}^{[l]})^T \in \mathbb{R}^{m \times d^{[l]}}, \quad \mathbf{A}^{[l]} = \sigma^{[l]}(\mathbf{Z}^{[l]}) \in \mathbb{R}^{m \times d^{[l]}}$$

2. Denote $\mathbf{Z}^{[l]} = (\mathbf{z}_1^{[l]}, \dots, \mathbf{z}_m^{[l]})^T$

3. After linear transformation, we consider the following **new** calculations:

$$\boldsymbol{\mu}^{[l]} = m^{-1}(\mathbf{Z}^{[l]})^T \mathbf{1} \in \mathbb{R}^{d^{[l]} \times 1},$$

$$\boldsymbol{\sigma}^{2[l]} = m^{-1} \sum_{i=1}^m \left(\mathbf{z}_i^{[l]} - \boldsymbol{\mu}^{[l]} \right) \circ \left(\mathbf{z}_i^{[l]} - \boldsymbol{\mu}^{[l]} \right) \in \mathbb{R}^{d^{[l]} \times 1},$$

$$\tilde{\mathbf{z}}_{i,norm}^{[l]} = \frac{\mathbf{z}_i^{[l]} - \boldsymbol{\mu}^{[l]}}{\sqrt{\boldsymbol{\sigma}^{2[l]} + \epsilon}} \in \mathbb{R}^{d^{[l]} \times 1},$$

$$\tilde{\mathbf{z}}_i^{[l]} = \gamma^{[l]} \circ \tilde{\mathbf{z}}_{i,norm}^{[l]} + \boldsymbol{\beta}^{[l]} \in \mathbb{R}^{d^{[l]} \times 1},$$

$$\tilde{\mathbf{Z}}^{[l]} = (\tilde{\mathbf{z}}_1^{[l]}, \dots, \tilde{\mathbf{z}}_m^{[l]})^T \in \mathbb{R}^{m \times d^{[l]}}$$

Forward propagation

1. Vectorization for the **red** calculations, but leave the **blue** parts alone

$$\mathbf{Z}^{[l]} = \mathbf{A}^{[l-1]}(\mathbf{W}^{[l]})^T \in \mathbb{R}^{m \times d^{[l]}},$$

$$\boldsymbol{\mu}^{[l]} = m^{-1}(\mathbf{Z}^{[l]})^T \mathbf{1} \in \mathbb{R}^{d^{[l]} \times 1},$$

$$\check{\mathbf{Z}}^{[l]} = \mathbf{Z}^{[l]} - \mathbf{1}(\boldsymbol{\mu}^{[l]})^T \in \mathbb{R}^{m \times d^{[l]}},$$

$$\boldsymbol{\sigma}^{2[l]} = m^{-1} \sum_{i=1}^m \check{\mathbf{z}}_i^{[l]} \circ \check{\mathbf{z}}_i^{[l]} \in \mathbb{R}^{d^{[l]} \times 1},$$

$$\check{\boldsymbol{\sigma}}^{[l]} = \sqrt{\boldsymbol{\sigma}^{2[l]} + \epsilon} \in \mathbb{R}^{d^{[l]} \times 1},$$

$$\hat{\boldsymbol{\sigma}}^{[l]} = (\check{\boldsymbol{\sigma}}^{[l]})^{-1} \in \mathbb{R}^{d^{[l]} \times 1},$$

$$\mathbf{Z}_{\text{norm}}^{[l]} = \check{\mathbf{Z}}^{[l]} \circ \{\mathbf{1}(\hat{\boldsymbol{\sigma}}^{[l]})^T\} \in \mathbb{R}^{m \times d^{[l]}}$$

2. Notice that the previous bias term $\mathbf{b}^{[l]}$ is useless for batch normalization



$$d\mathbf{Z}_2^{[l]} = m^{-1} \mathbf{1} (d\boldsymbol{\mu}^{[l]})^T$$

Remarks about BN

1. Disadvantage

- Since we normalize “inputs” before activation, many different “inputs” may result in same “activations”
- Besides, batch normalization also introduces more model parameters

2. Advantage

- Stabilize forward propagation (The variance can be controlled by γ 's)
- Higher learning rates (Batch normalization can make loss and its gradients more smooth)
- Regularization
 - ▷ We inject noises from other training examples through mean and variance
 - ▷ Thus, batch normalization may improve the generality of the network

Remarks about BN

1. When applied to validation or test, we usually use the sample mean and variance from the training dataset, instead
 - The sample size may be limited
2. Check Santurkar (2018) and others for more discussion
3. We may introduce other normalization methods in the future sections
4. Why not do batch normalization after activation?

Remarks about dropout and BN

1. BN normalizes each feature to achieve better convergence rate and generality
2. Dropout improves robustness through randomly removing neurons
3. If BN is applied after dropout, we may have problems ([Pitfall 1](#))
 - During training, dropout may change the variance of neurons
 - However, BN stabilizes variance for each mini-batch
 - Moreover, dropout is not conducted for the validation and test data
 - It may cause the shift of statistics, leading to instability (Li et al., 2019)
4. Recommended order:
Linear $\rightarrow$ BN $\rightarrow$ Activation $\rightarrow$ Dropout

Remarks about dropout and BN

1. The implementation of both BN and dropout may lead to over regularization ([Pitfall 2](#))
 - BN is a regularization method
 - Implementation of both BN and dropout may result in underfitting
2. Recommendation
 - Use BN + dropout with smaller removal probability
 - Use BN without dropout

Summary

1. One obstacle for training deep neural networks is internal covariate shift
2. Based on the idea of whitening, BN normalizes each feature individually
3. BN stabilizes training, enables larger learning rates, and provides mild regularization