



人工智能与深度学习

3-3 批量归一化

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. 简介

2. 前向传播

3. 后向传播

4. 关于 dropout 和批量归一化的讨论

简介

1. 2015 年以前，深度学习模型的训练存在严重的优化困难

- 后向传播的不稳定
- 网络对参数初始值极端敏感

2. 例如，一个常被讨论的问题是内部协变量漂移（internal covariate shift）

- 较浅层模型参数的微小更新，可能会对深层输入的分布，产生较大的影响
- 较深层的模型参数，始终在“追逐移动目标”

3. 这种不稳定，主要表现为下面的两个问题

- 梯度消失：当线性变换结果，漂移到激活函数的饱和区时，局部梯度趋近于 0
- 梯度爆炸：权重初始化过大，则激活的方差会随层数放大，导致代价函数发散

简介

1. 对训练集进行“白化（whitening）”，是经典统计学习模型的常规操作

- 将输入数据变为零均值、单位方差，且特征间互不相关
- 算法收敛可能会更快

2. 上面的白化操作，对于深度学习而言不可行

- 训练集的规模往往较大
- 我们不能确保每一层都能实现白化操作
- 白化操作的计算成本高

简介

1. Ioffe and Szegedy (2015) 提出了批量归一化 (Batch normalization)

- 逐特征独立处理：对每个（神经元输出的）特征独立进行归一化，而不考虑特征之间的协方差
- 小批量统计估计：使用当前 mini-batch，得到所需的均值和方差

2. 在**线性变换后**，但在**激活运算前**，实施批量归一化

3. 批量归一化可实现如下效果

- 可用于解决内部协变量漂移的问题
- 基于两个额外模型参数，调整特征之间的异质性

前向传播

1. 对于小批量梯度下降法（批量大小为 m ），第 l 层的计算为：

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

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

3. 线性变换后，我们考虑下面的新运算：

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

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

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

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

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

前向传播

1. 考虑红色计算的向量化，暂时不考虑蓝色部分的计算细节

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

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

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

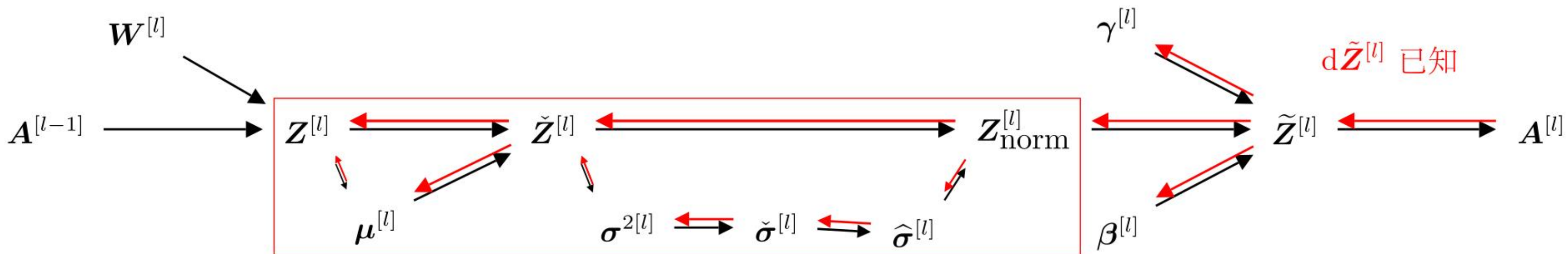
$$\boldsymbol{\sigma}^{2[l]} = m^{-1} \sum_{i=1}^m \check{\mathbf{z}}_i^{[l]} \circ \check{\mathbf{z}}_i^{[l]} \in \mathbb{R}^{d^{[l]} \times 1},$$

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

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

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

2. 注意：对于批量归一化而言，偏置项 $\mathbf{b}^{[l]}$ 被均值中心化抵消



红色计算部分的后向传播过程为：

$$d\mathbf{Z}^{[l]} = d\mathbf{Z}_1^{[l]} + d\mathbf{Z}_2^{[l]} \quad d\beta^{[l]} = \left(d\tilde{\mathbf{Z}}^{[l]}\right)^T \mathbf{1} \quad d\gamma^{[l]} = \left(d\tilde{\mathbf{Z}}^{[l]} \circ \mathbf{Z}_{\text{norm}}^{[l]}\right)^T \mathbf{1}$$

$$d\check{\mathbf{Z}}^{[l]} = d\check{\mathbf{Z}}_1^{[l]} + d\check{\mathbf{Z}}_2^{[l]} \quad d\check{\sigma}^{[l]} = -d\hat{\sigma}^{[l]} \circ \hat{\sigma}^{[l]} \circ \hat{\sigma}^{[l]} \quad d\mathbf{Z}_{\text{norm}}^{[l]} = d\tilde{\mathbf{Z}}^{[l]} \circ \left\{ \mathbf{1} (\gamma^{[l]})^T \right\}$$

$$d\mathbf{Z}_1^{[l]} = d\check{\mathbf{Z}}^{[l]} \quad d\sigma^{2[l]} = d\check{\sigma}^{[l]} \circ \hat{\sigma}^{[l]} / 2 \quad d\check{\mathbf{Z}}_1^{[l]} = d\mathbf{Z}_{\text{norm}}^{[l]} \circ \left\{ \mathbf{1} (\hat{\sigma}^{[l]})^T \right\}$$

$$d\mu^{[l]} = -\left(d\check{\mathbf{Z}}^{[l]}\right)^T \mathbf{1} \quad d\check{\mathbf{Z}}_2^{[l]} = 2m^{-1} \check{\mathbf{Z}}^{[l]} \circ \left\{ \mathbf{1} (d\sigma^{2[l]})^T \right\} \quad d\hat{\sigma}^{[l]} = \left(d\mathbf{Z}_{\text{norm}}^{[l]} \circ \check{\mathbf{Z}}^{[l]}\right)^T \mathbf{1}$$

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

批量归一化的备注

1. 缺点

- 由于我们在激活前使用批量归一化，不同的“输入特征”可能有相同的“激活值”
- 此外，批量归一化对应着更多的模型参数

2. 优点：

- 稳定化前向传播（批量归一化的方差被 γ 控制）
- 可以接受更大的学习率（批量归一化可使代价函数及其梯度变得更平滑）
- 正则化
 - ▷ 批量归一化可被认为，将其他训练样本的随机性，加入到每个训练样本
 - ▷ 因此，批量归一化可以提升神经网络模型的泛化能力

批量归一化的备注

1. 在验证和测试阶段，我们往往利用训练集上累计的均值和方差运行估计，即其对应的 `runtime estimators`，进行批量归一化处理
 - 在这两个阶段，样本点的数量可能较少，甚至只有一个
2. 更多的讨论，请参考 Santurkar (2018) 等
3. 我们将在后续课程中，介绍其他的归一化方法
4. 为什么不在激活运算之后，进行批量归一化？

Dropout 与批量归一化的备注

1. 批量归一化通过更改每个特征对应的均值和方差，加快算法收敛，提高模型的泛化能力
2. Dropout 通过随机删除隐藏层的神经元，来提升模型的泛化能力
3. 当我们在 dropout 后实施批量归一化时，可能会遇到麻烦（陷阱一）
 - 在训练阶段，由于 dropout 随机丢弃神经元，其可能会改变未丢弃神经元的方差
 - 然而，批量归一化则在努力稳定神经元的方差
 - 此外，dropout 在模型的验证和测试阶段不再实施
 - 这可能会造成统计量偏移，导致训练阶段与推理阶段的统计量不匹配，进而降低推理性能（Li 等）
4. 建议的实施顺序：
线性变换 $\rightarrow$ 批量归一化 $\rightarrow$ 激活 $\rightarrow$ Dropout

Dropout 与批量归一化的备注

1. 当我们同时实施批量归一化和 dropout 时，可能会造成过度正则化的风险（[陷阱二](#)）

- 批量归一化具有一些正则化效果
- 同时实施批量归一化和一个具有较高移除概率的 dropout，可能会导致欠拟合

2. 建议

- 批量归一化 + 一个带有较小移除概率的 dropout
- 只实施批量归一化

课程总结

1. 深层网络训练困难的一大原因，是内部层输入分布，在训练中不断变化
2. 基于白化操作的思想，批量归一化对每一个特征单独操作
3. 批量归一化能稳定训练、支持更大学习率，并带来轻度正则化