

Course 1: Toy Example

Deep learning Basics &

如何选择损失函数/优化器/神经网络/正则化方法
loss optimizer neural network regularization

Predictive models

如何通过数据增强/生成数据来训练更强大的模型
augment generate $\lambda=1$

概念:

1. Supervised learning 监督学习: 通过使用带标签的历史数据训练模型, 使其能够对未知新数据进行预测的机器学习方法。

2. Classification 分类: 监督学习的核心任务之一, 旨在根据输入特征将其预测为两个或多个离散的类别标签。

3. Transfer Learning 迁移学习: 一种机器学习方法, 它从一个任务(源任务)中学到的知识, 应用到另一个不同但相关的任务(目标任务)中, 从而减少对大量新标注数据的依赖并加快模型训练。

4. Unsupervised Learning 无监督学习: 一种机器学习方法, 它在使用无标签数据进行训练时, 自主发现数据内部隐藏的结构、分布或模式(如聚类或降维), 而无需人为提供正确答案。

5. Generative Model: 通过从训练数据中学习隐含的概率分布, 能够生成与原始数据相似的全新样本(如图像、文本或音频)。

6. Generalization 泛化: 模型在训练集上学习后, 能够对未见过的全新数据做出准确预测的能力。

7. Linear Models 线性模型: 假设输入特征与输出之间存在线性关系, 通过特征的加权和进行预测。例如: $y(x; w) = w_0 + w_1 x + w_2 x^2 + \dots + w_m x^m = \sum_{j=0}^m w_j x^j$, x 为特征, w_j 为模型参数(权重), y 为预测输出, 对 w_j 而言是线性。

8. Model Complexity

8. Overfitting 过拟合: 模型过度学习训练数据中的噪声和细节, 导致其在训练集上表现极好, 但在未见过的测试集上表现很差。

在 deep learning 中, 数据量 n 都远大于数据量 $\frac{1}{n} \sum (y - \hat{y})^2$

MSE

$$E(w) = \frac{1}{n} \sum_{n=1}^n (y(x_n, w) - t_n)^2$$

预测值 真实

1. 可导且平滑

2. 若假设 $\epsilon \sim N(0, \sigma^2)$, 最大似然的结果等价于最小化 MSE。

3. 因为误差被平方, 较大的误差会被显著放大。

$$L(\theta) = \prod_{i=1}^n \frac{1}{\sqrt{2\pi}\sigma} \exp\left(-\frac{(y^{(i)} - h_{\theta}(x^{(i)}))^2}{2\sigma^2}\right)$$

$$l(\theta) = n \log \frac{1}{\sqrt{2\pi}\sigma} - \frac{1}{2\sigma^2} \sum (y^{(i)} - h_{\theta}(x^{(i)}))^2$$

~~$\frac{1}{n} \sum (y - \hat{y})$~~ MAE MSE

$$E(w) = \frac{1}{n} \sum_{n=1}^n |y(x_n, w) - t_n|$$

1. 与原始数据单位相同

2. 对异常值鲁棒

3. 在 0 处不可导

$$\sqrt{\frac{1}{n} \sum (y - \hat{y})^2}$$

RMS

$$E_{rms} = \sqrt{\frac{1}{n} \sum_{n=1}^n (y(x_n, w) - t_n)^2}$$

误差的均方根, 单位与原始数据一致。

各个分量
M 越大, w_j 的绝对值也会增大 $\rightarrow$ 抖动复杂。

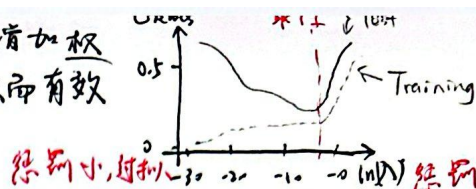
模型拟合数据点的方式是
最小化误差方程。

训练
数据增多数据可以来减小过拟合。

9. Weight Decay 权重衰减: 通过在损失函数中增加权重 w 的平方和惩罚项, 来限制模型参数过大, 从而有效抑制过拟合的常用正则化技术.

$$\mathcal{E}(w) = \frac{1}{2} \sum_{n=1}^N |y(x_n, w) - t_n|^2 + \frac{\lambda}{2} \|w\|^2, \quad \lambda \in (e^{-5}, e^3), \quad \lambda \in (-5, -3) \text{ 正常范围}$$

$\|w\|^2 = w^T w = w_0^2 + w_1^2 + \dots + w_m^2$ 减小 w 的绝对值大小.



10. Validation set 验证集: 用于评估模型在训练过程中的泛化能力并调整超参数 (如学习率、正则化强度) 的数据子集, 它不与梯度更新, 但会影响模型的选择与配置.

λ, μ 需要通过验证集确定, 以后可以将验证集加入训练集.

11. Test set 测试集: 在模型训练和所有调参工作完全结束后, 用于最终一次性评估模型泛化能力的独立数据.

Why need test and validation sets? If limited data?

验证集用于调参, 测试集用于评估, 并且前期不能给模型接触.

数据少可以采用交叉验证. 训练集分为 k 份, 循环 k 次, 每次用不同一份作为验证集.

Course 2: Linear Classification

Linear function: $y(x, w) = w_0 + w_1 \phi_1(x) + \dots + w_m \phi_m(x) = w_0 + \sum_{i=1}^m w_i \phi_i(x)$

$\uparrow$
learnable parameters 可学习参数
因子/基函数

1. Feature extraction 特征提取: 将原始数据 (如文本、图像、信号) 通过变换或映射, 转换为更具代表性, 信息更浓缩的数值特征向量的过程. eg. PCA

真实 预测

$$t_i = y_i(x_i, w) + \varepsilon_i, \quad \varepsilon_i \sim N(0, \sigma^2) \rightarrow t_i \sim N(y_i(x_i, w), \sigma^2)$$

$$f(t_i | x_i, w) = \frac{1}{\sqrt{2\pi}\sigma} \exp\left(-\frac{(t_i - y_i(x_i, w))^2}{2\sigma^2}\right)$$

$$L(w) = \prod_{i=1}^n \frac{1}{\sqrt{2\pi}\sigma} \exp\left(-\frac{(t_i - y_i(x_i, w))^2}{2\sigma^2}\right)$$

$$\ell(w) = \ln L(w) = \sum_{i=1}^n \ln \left[\frac{1}{\sqrt{2\pi}\sigma} \exp\left(-\frac{(t_i - y_i(x_i, w))^2}{2\sigma^2}\right) \right]$$

$$= n \ln \frac{1}{\sqrt{2\pi}\sigma} - \frac{1}{2\sigma^2} \sum_{i=1}^n (t_i - y_i(x_i, w))^2$$

$$= -n \ln(\sqrt{2\pi}\sigma) - \frac{1}{2\sigma^2} E_0(w)$$

$$E_0 = \frac{1}{2} \sum_{i=1}^n (t_i - w^T \phi(x_i))^2 \text{ 均方误差}$$

$\leftarrow R$ 反函数

$$p(x | \mu, b) = \frac{1}{2b} \exp\left(-\frac{|x - \mu|}{b}\right) \text{ 拉普拉斯分布}$$

$$f(t_i | x_i, w) = \frac{1}{2b} \exp\left(-\frac{|t_i - y_i(x_i, w)|}{b}\right)$$

$$L(w) = \prod_{i=1}^n \frac{1}{2b} \exp\left(-\frac{|t_i - y_i(x_i, w)|}{b}\right)$$

$$\ell(w) = -n \ln(2b) - \frac{1}{b} \sum_{i=1}^n |t_i - y_i(x_i, w)|$$

$$= -n \ln(2b) - \frac{1}{b} E_0(w)$$

$$E_0(w) = \frac{1}{n} \sum_{i=1}^n |t_i - w^T \phi(x_i)|$$

$$\varepsilon \sim N(0, \sigma^2)$$

“最大似然” $\rightarrow$ “最小 MSE”

~~ε 均值为 0 的拉普拉斯分布~~

“最大似然” $\rightarrow$ “最小 MAE”

y_i 为真实值

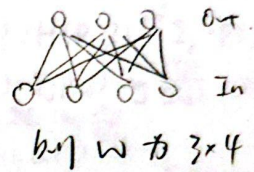
$\hat{y}_i$ 为预测值

2. One-hot encoding 独热编码: 将离散类别特征转换为稀疏二进制向量的方法: 为每个类别创建一个独立的维度, 在样本所属类别的维度上置1, 其余维度置0.

Cat	1	$\{1, 0, 0\}$
Chicken	2	$\{0, 1, 0\}$
Dog	3	$\{0, 0, 1\}$
	$\{1, 2, 3\}$	$\{(1, 0, 0), (0, 1, 0), (0, 0, 1)\}$
X		V

$$\text{softmax}(X_i) = \frac{\exp(x_i)}{\sum_j \exp(x_j)}$$

3. full-connected layer 全连接层: 一种神经网络层, 其中每个输入神经元与每个输出神经元之间均有权重连接, 通过对输入向量进行线性变换 ($y = Wx + b$) 来学习全局特征间的组合关系.



$$Wx + b$$

3x4 4x1 3x1

$P(y=1) \propto \exp(0)$ 归一化 $\rightarrow \hat{y} = \text{softmax}(0), \hat{y}_i = \frac{\exp(o_i)}{\sum_j \exp(o_j)}$ 和为1. 概率特性

$\text{softmax}(X_i) = \frac{\exp(x_i)}{\sum_j \exp(x_j)}$

$\arg \max_j \hat{y}_i = \arg \max_j o_j$ 单调性

向量例: 对每个样本计算 $o = Wx + b$ 太慢

$X \in \mathbb{R}^{n \times d}$, n 为样本数, d 为维度. $O = XW + b$ $\leftarrow$ 会通过广播扩展到与 XW 相同维度

$\hat{y} = \text{softmax}(O)$. 得到每个样本在 q 个类别上的概率分布.

4. Cross-entropy loss 交叉熵损失: 用于量化模型预测的概率分布 $\hat{y}$ 与真实标签分布 y (通常为独热编码) 之间的差异, 其公式为 $-\sum_i y_i \log \hat{y}_i$; 分类任务.

在分类任务中, 它等价于最大化预测正确类别的对数似然, 且当预测完全准确时损失为0. $L(y, \hat{y}) = -\sum_{i=1}^q y_i \log \hat{y}_i \in [0, +\infty)$

正确类别的负对数概率

模型输出 $\hat{y}$ 表概率分布.
真实标签 y 为独热编码
 $L(y, \hat{y}) = -\log \hat{y}_k$.
表示类别间相似度时
用欧氏距离.

$\hat{y} = \text{softmax}(o)$, where $\hat{y}_i = \frac{\exp(o_i)}{\sum_j \exp(o_j)}$, $L(y, \hat{y}) = -\sum_{i=1}^q y_i \log \hat{y}_i$

$$\begin{aligned} L(y, \hat{y}) &= -\sum_{j=1}^q y_j \log \frac{\exp(o_j)}{\sum_{k=1}^q \exp(o_k)} \\ &= \sum_{j=1}^q y_j \log \left(\frac{1}{\sum_{k=1}^q \exp(o_k)} \right) - \sum_{j=1}^q y_j o_j \\ &= \log \left(\frac{1}{\sum_{k=1}^q \exp(o_k)} \right) - \sum_{j=1}^q y_j o_j \end{aligned}$$

$$\begin{aligned} \frac{\partial L(y, \hat{y})}{\partial o_j} &= \frac{\partial}{\partial o_j} \left(\log \left(\frac{1}{\sum_{k=1}^q \exp(o_k)} \right) \right) - \frac{\partial}{\partial o_j} \left(\sum_{i=1}^q y_i o_i \right) = \frac{1}{\sum_{k=1}^q \exp(o_k)} \frac{\partial}{\partial o_j} \left(\sum_{k=1}^q \exp(o_k) \right) - y_j \\ &= \frac{\exp(o_j)}{\sum_{k=1}^q \exp(o_k)} - y_j = \text{softmax}(o_j) - y_j = \hat{y}_j - y_j \end{aligned}$$

交叉熵梯度 = 模型预测概率 - 真实

eg. $q = \{0.2, 0.1, 0.1\}$, $y_1 = 1, y_2 = 0, y_3 = 0$.

~~Softmax~~ Softmax 输出: $\hat{y} = [0.659, 0.242, 0.099]$.

交叉熵损失: $l = -\log(0.659) \approx 0.417$.

梯度 $\frac{\partial l}{\partial \theta} = [0.659 - 1, 0.242 - 0, 0.099 - 0] = [-0.341, 0.242, 0.099]$

$0 \leftarrow 0.099$. 梯度相反方向: $\uparrow 2, 1, 0.1$

Entropy, CE loss and KL divergence.

P 为真实分布, Q 为模型预测分布

熵

$[0, \log 2]$ Entropy 熵 (衡量一个随机分布的不确定性, 分布越集中熵越小): $H[P] = -\sum_j P(j) \log P(j)$

$[0, +\infty)$ Cross-entropy 交叉熵 (衡量两个分布 P 和 Q 的差异, 越小差异越小): $H(P, Q) = -\sum_j P(j) \log Q(j)$

$[0, +\infty)$ KL divergence KL 散度 (衡量从分布 Q 到分布 P 的差异): $D_{KL}(P||Q) = \sum_j P(j) \log \frac{P(j)}{Q(j)}$ 不对称

$H(P, Q) = H[P] + D_{KL}(P||Q)$ 字母顺序

① 当 P 固定时, 最小的交叉熵等价于最小的 KL 散度.

② 若 P 分布不为独热, 则 $H[P] > 0$. 即使 $P = Q$, $H(P, Q) = H[P] > 0$.

③ 只有当 P 分布为独热, $H[P] = 0$. $H(P, Q) = D_{KL}(P||Q)$. 完全正确预测才会使 $H(P, Q) = 0$.

CE Loss 在分类任务中优于 MSE 的根本原因: $\frac{\partial \text{Loss}}{\partial z} = (\hat{y} - y)\hat{y}(1-\hat{y})$. 其中 $\hat{y} \cdot (1-\hat{y})$ 在 $\hat{y}$ 接近 0 或 1 时该值趋近于 0. 梯度小. 而 $\frac{\partial \text{Loss}}{\partial \theta} = \hat{y} - y$ 直接反映预测误差, 在 $[0, 1]$ 上对错误值更新更快.

Course 3: Multi-Layer Perceptron

线性模型中, 多数对结果的方向的影响是单调的

单调不够 $\xrightarrow{\text{举例}}$ 温度对于健康 固定的基函数不够 $\xrightarrow{\text{举例}}$ 不同图像

Multilayer Perceptron 多层感知机: 一种由输入层, 一个或多个隐藏层和输出层构成的前馈神经网络, 通过激活函数引入非线性变换, 能学习复杂的函数映射关系.

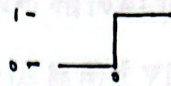
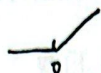
$$H = XW^{(1)} + b^{(1)} \quad O = HW^{(2)} + b^{(2)}$$

$$O = (XW^{(1)} + b^{(1)})W^{(2)} + b^{(2)} = XW^{(1)}W^{(2)} + b^{(1)}W^{(2)} + b^{(2)} = XW + b \text{ 等价于单层}$$

根据通用近似定理, 对于具有线性输出层和至少一个使用“挤压”性质激活函数的隐藏层组成的前馈神经网络, 只要其隐藏层神经元的数量足够, 它可以以任意的精度来近似任何一个定义在实数空间中的有界连续函数.

三种激活函数:

① ReLU $\max(x, 0)$



$\text{ReLU}(x) = \max(x, 0) \in [0, +\infty)$ 梯度 $\in [0, 1]$

简单, 稀疏, 无梯度消失 (但有死神经元), 中间层常用

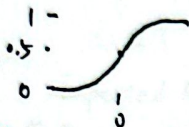
$$\frac{1}{1+e^{-x}}$$

某些神经元始终输出 0

② Sigmoid (在二分类中当 softmax 用)

$\text{sigmoid}(x) = \frac{1}{1+\exp(-x)} \in [0, 1]$

梯度 $= \frac{\exp(-x)}{(1+\exp(-x))^2} \in [0, \frac{1}{4}]$



$x=0$ 时, 最大 $= \frac{1}{(1+1)^2} = \frac{1}{4}$

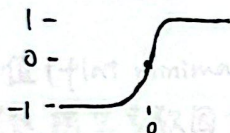
也可写为: $\text{sigmoid}(x)(1-\text{sigmoid}(x))$

平滑, 概率输出 (但梯度消失), 二分类常用

③ Tanh

$\tanh(x) = \frac{1-\exp(-2x)}{1+\exp(-2x)} \in [-1, 1]$

梯度 $= \frac{4e^{-2x}}{(1+e^{-2x})^2} \in [0, 1]$



$x=0$ 时, 最大 $= \frac{4}{(1+1)^2} = 1$

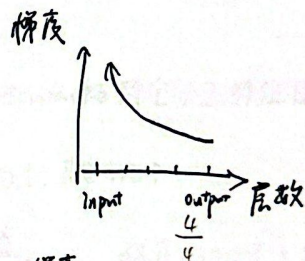
“网”好, 再大会发生梯度爆炸

对称, 范围广 (但优化仍难), 回归问题常用

梯度爆炸的表现: 损失值变为 NaN

原因: 网络过深, 激活函数不当, 学习率过高

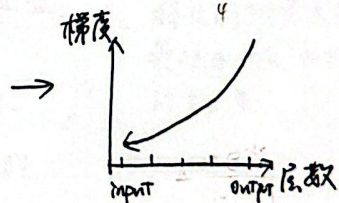
解决方法: 梯度裁剪, 残差连接, 降低学习率



梯度消失的表现: 损失值下降缓慢, 权重几乎不更新

原因: 网络太深, 激活函数不当

解决方法: ReLU, 残差连接



一般而言, 靠近输出层梯度越大, 靠近输入层梯度越小

$$(1-e^{-2x}) \cdot 2e^{-2x} + 2e^{-4x} + 2e^{-2x} - 2e^{-4x}$$

隐藏层 $\rightarrow$ 增加非线性表达 (隐藏层通过多层非线性变换, 使得神经网络能够拟合任意复杂的函数)

激活函数 $\rightarrow$ 增加层仍为线性组合 (如果没有激活函数, 无论网络有多少层, 最终的输出仍然是输入的线性组合)

Sigmoid 缺点 $\rightarrow$ 梯度消失, $(\text{sigmoid})'_{\max} = 0.25 < 1$

$$\frac{1}{(1+e^{-2x})^2}$$

如果所有权重初始化为0? 神经元会学习到相同的特征, 导致网络失去表达能力。

神经网络为什么更深而不是更宽? 更深的网络能用更少的多数指数级地表达更复杂的函数, 多数效率更高, 且深层网络更容易捕捉层次化特征, 泛化能力也更强。

Hinton的Alex Net使用ReLU来加速训练。

Course 4: Optimization

$$\theta^* = \arg \min_{\theta} \sum_{(x,y) \in D} L(f_{\theta}(x), y)$$

↑ 梯度下降
↑ 数据集
↑ 神经网络

Expected Risk

Empirical Risk

真实分布 $P(z)$ 计算期望风险 $R(\theta)$, 训练集 D_n 上计算经验风险 $\hat{R}(\theta, D_n)$.

$R(\theta) \neq \hat{R}(\theta, D_n)$, 两者之间的差距称为泛化误差。(过拟合)

训练集 D_n 与真实数据分布 $P(z)$ 独立同分布 $\rightarrow E(\hat{R}(\theta, D_n)) = R(\theta)$

$$E[\hat{R}(\theta, D_n)] = E\left[\frac{1}{n} \sum_{i=1}^n L(\theta, z_i)\right] = \frac{1}{n} \sum_{i=1}^n E[L(\theta, z_i)] = E[L(\theta, z)] = R(\theta)$$

平坦的最小值 (flat minima of empirical risk) 通常对应更好的泛化能力。
这意味着即使模型多数因为量级误差, 随机初始化或训练集会变动而发生微小偏移, 损失值也不会有太大变化。

可能停止优化: $\nabla L \approx 0$

认为优化收敛: 梯度趋0

可能减慢收敛

Global minimum

Local minimum

Saddle point 损失函数在某些方向↓, 其它方向↑

Hessian 矩阵全为正特征值 & 0.

Hessian 矩阵同时有正和负的特征值

梯度的方向是函数值上升最快的方向, 大小是该方向的变率。

$$\theta_{t+1} = \theta_t + \eta g_t = \theta_t - \eta \nabla L(\theta_t)$$

Gradient Descent 梯度下降 GD

一种迭代算法, 通过沿与梯度相反的方向移动来最小化函数。

$$\theta_{t+1} = \theta_t - \eta \nabla L(\theta_t)$$

θ 为参数, η 为学习率, $\nabla L(\theta_t)$ 为损失函数 L 的梯度。

优: ① 对于凸函数收敛到全局最小; ② 实现简单。

缺: ① 计算成本高, 每次迭代需使用整个数据集; ② 对大数据集, 训练缓慢;

③ 对于非凸函数, 会卡在局部最小。

$$\theta_{t+1} = \theta_t - \eta \nabla L(\theta_t; x_i, y_i)$$

随机初始化 θ 。

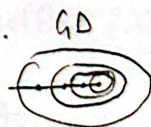
Stochastic Gradient Descent 随机梯度下降 SGD

GD 的一种变体, 使用单个随机样本的梯度来更新参数。

$$\theta_{t+1} = \theta_t - \eta \nabla L(\theta_t; x_i, y_i)$$

优: ① 效率高, 每次迭代只用一个样本;
② 波动性有助于逃离局部最小。

缺: ① 更新不稳定, 导致收敛慢; ② 可能无法达到精确最小值。



Mini-batch SGD 小批量SGD

$$\theta_{t+1} = \theta_t - \eta \nabla L(\theta_t; X_b, Y_b)$$



GD和SGD之间的折衷, 每次迭代使用小批量样本.

$$\theta_{t+1} = \theta_t - \eta \nabla L(\theta_t; X_b, Y_b) \quad (X_b, Y_b) \text{ 是样本的一个小批量}$$

优: ①平衡了效率与稳定性; ②可利用矩阵计算 (GPU) 加速; ③很适合深度学习 & 大数据集.

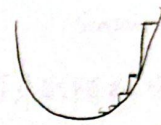
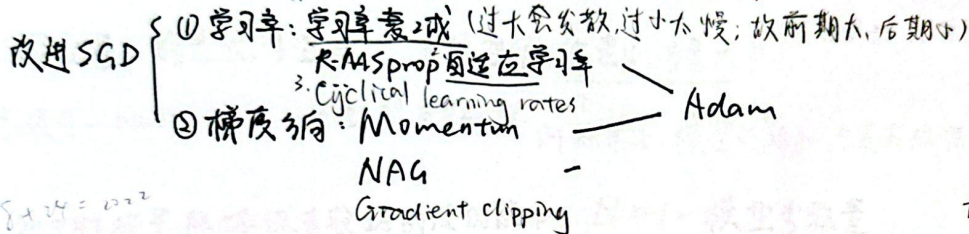
缺: ④需调整批量大小 (eg. 32, 64, 256); ⑤仍容易陷入局部最小.

Batch Size & Learning Rate

批量越大, 梯度方差越小 & 噪声越小, 用大学习率

批量越小, 梯度方差越大 & 噪声越大, 用小学习率.

(学习率调度) Learning Rate Schedule 可以动态调整学习率.



Cyclical Learning Rates 循环学习率, 帮助模型跳出坏区域.

$$\Delta \theta_t = -\alpha g_t$$

Momentum Method 动量法: $\Delta \theta_t = \rho \Delta \theta_{t-1} - \alpha g_t = -\alpha \sum_{r=0}^t \rho^{t-r} g_r$ 的衰减速度.

α 为学习率, g_t 为当前步的梯度, $\Delta \theta_t$ 为第 t 步的多数更新量, ρ 为动量因子, 控制历史梯度的衰减速度.

$$\Delta \theta_t = \rho \Delta \theta_{t-1} - \alpha g_t$$

0.9 / 0.99 历史情况的重要

Nesterov accelerate gradient (NAG): $\Delta \theta_t = \rho \Delta \theta_{t-1} - \alpha g_t(\theta_{t-1} + \rho \Delta \theta_{t-1})$ 在这个点计算梯度
 $\rho \Delta \theta_{t-1}$ 为历史更新方向, $(\theta_{t-1} + \rho \Delta \theta_{t-1})$ 为向前看一步的位置, α 为学习率.

优: "预见"能力, 比动量法震荡更小, 收敛更快.

$$\Delta \theta_t = \rho \Delta \theta_{t-1} - \alpha g_t(\theta_{t-1} + \rho \Delta \theta_{t-1})$$

Adam 的本质是 Momentum (动量) + RMSprop (自适应学习率)

Gradient Clipping

- by value: $g_t = \max(\min(g_t, b), -b)$
- by norm: $g_t = \frac{b}{\|g_t\|} g_t$

$$\Delta \theta_t = \rho \Delta \theta_{t-1} - \alpha g_t(\theta_{t-1} + \rho \Delta \theta_{t-1})$$

$$\Delta \theta_t = \rho \Delta \theta_{t-1} - \alpha g_t$$

Parameter Initialization

如果初始多数全为 0 会导致对称权重 $\rightarrow$ 随机初始化 ($N(0, \epsilon)$ 或 $U[-\epsilon, \epsilon]$).

数据归一化使得梯度从指向最小值方向不直接 $\rightarrow$ 指向最小值方向直接.

同时, 减小震荡, 要用更大学习率, 对学习率敏感度更低.

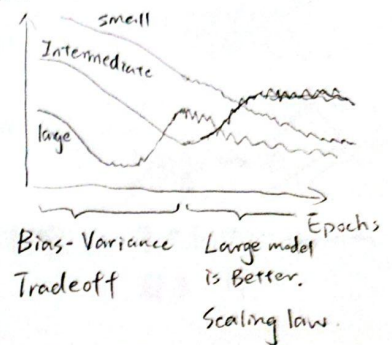
可

Course 5: Regularization

过拟合原因: ①训练数据不足; ②参数过多(过参数化); ③数据存在噪声

double decent, 且后期下降时震荡大 $\rightarrow$ overfitting 不是单调的

Test error
大模型(过多参数)通常需要更少的 epoch 就会开始过拟合。



① 偏差的平方 $bias^2 = (\bar{g}(x) - f(x))^2$: $(\bar{g}(x) - f(x))^2$

平均模型 $\bar{g}$ 与真实目标函数 f 之间的差距。衡量模型集的表达能力。表达能力越高, 偏差越小。

② 方差 $Var = E_0[(g^{(D)}(x) - \bar{g}(x))^2]$: $E_0[(g^{(D)}(x) - \bar{g}(x))^2]$

不同数据集训练出的模型 $g^{(D)}$ 与平均模型 $\bar{g}$ 之间的差异。衡量模型对数据扰动的敏感度。

小模型集: 偏差大, 方差小。大模型集: 偏差小, 方差大。模型开始记住训练集的随机噪声。

总误差 = $bias^2 + variance + \text{数据本身的噪声}$ 。训练集上, 模型开始拟合真实规律。

多少数据量能确保有较好的泛化能力: $N \geq 10 \cdot \text{模型参数量}$

1. Regularization 正则化: 通过对模型复杂度施加惩罚(如限制参数大小或稀疏性)来抑制过拟合, 提升泛化能力的技术。

正则化 { ① 在优化过程中添加约束条件: L1/L2 正则化, 数据增强
② 干扰优化过程: Weight Decay 权重衰减, SGD, Early stopping 早停

与 L2 区别: 在自适应优化器中, L2 正则化效果会扭曲, weight decay 正则化力度与梯度大小解耦。

2. Data augmentation 数据增强: 对现有训练数据进行微小的随机变换(如旋转, 裁剪, 翻转)生成更多变体, 以低成本扩充数据集, 从而缓解过拟合。不能改变 label

提高泛化能力的方法: $\|w\|_1 = \sum |w_i|$

$$\|w\|_2 = \sqrt{\sum w_i^2}$$

① L1/L2 正则化: $\theta^* = \arg \min_{\theta} \frac{1}{N} \sum_{i=1}^N L(y^{(i)}, f(x^{(i)}, \theta)) + \lambda L_p(\theta)$, $p \in \{1, 2\}$, λ 是正则化的强度

$$L1: \|w\|_1 = \sum_{j=1}^m |w_j|, L2: \|w\|_2 = \sqrt{\sum_{j=1}^m w_j^2}$$

稀疏化, 用于特征选择(公司裁员, 裁低工资) 削弱所有权重但仍保留, 对更大的权重影响更大, 也叫权重衰减(公司降工资, 先降发薪高的)

② Weight Decay 权重衰减: 在每次迭代时引入一个衰减系数 w .

$$\theta_t \leftarrow (1-w)\theta_{t-1} - \alpha g_t$$

在标准 SGD 中, 权重衰减与 L2 正则化有相同效果

在 Adam 中, 因为自适应学习率, 二者不同, 应用 Weight Decay.

$$L2: \text{Loss function} =$$

$$L(\theta) + \frac{\lambda}{2} \|\theta\|_2^2$$

$$\bar{\theta} = \theta_t + \lambda \theta_{t-1}$$

$$\theta_t \leftarrow \theta_{t-1} - \alpha(g_t + \lambda \theta_{t-1}) \leftarrow (1-\alpha\lambda)\theta_{t-1} - \alpha g_t$$

③ Early Stopping 早停

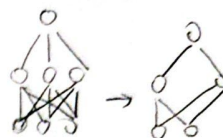
早停法通过在验证集误差不再下降时提前终止训练，防止模型在训练集上过度拟合，从而提升泛化能力。

④ Dropout

对于神经网络层 $y = f(Wx + b)$ ，引入函数 $d(\cdot)$ ，使得 $y = f(Wd(x) + b)$ 。

$$d(x) = \begin{cases} m \odot x & \text{训练阶段} \\ x & \text{测试阶段} \end{cases}$$

$$y = f(Wx + b) \rightarrow y = f(Wd(x) + b)$$



其中 $m \in \{0, 1\}^n$ 是 dropout 掩码，以概率 p 按伯努利分布随机生成。 p 为“1”的概率，也就是保留的概率。

每次训练的 iteration 都不一样

有效的
原因

① 集成学习：每次 dropout，相当于采样一个子网络； n 个神经元的网络可以采样出 2^n 个子网络。

② 贝叶斯学习： $E_{q(\theta)}[y] = \int_{\theta} f(x, \theta) q(\theta) d\theta \approx \frac{1}{M} \sum_{m=1}^M f(x, \theta_m)$ 经验采样的近似
↑
第 m 次 dropout 后的网络

⑤ Data augmentation

⑥ Label Smoothing 标签平滑

将硬标签（如 one-hot 中的 1）替换为软标签（ $1 \rightarrow 1 - \epsilon$ ，其余类别均分 ϵ ）来防止模型过度自信，从而提升泛化能力。

$$y = [0, \dots, 0, 1, 0, \dots, 0]^T \rightarrow \tilde{y} = [\frac{\epsilon}{K-1}, \dots, \frac{\epsilon}{K-1}, 1 - \epsilon, \frac{\epsilon}{K-1}, \dots, \frac{\epsilon}{K-1}]^T$$

⑦ Batch Normalization 批归一化

逐层归一化 $\xrightarrow{\text{目标}}$ ① 更好的尺度不变性 B/N ② 减少内部协变量偏移 ③ 更平滑的优化地形 对每个中间层进行归一化。
Batch Normalization 批归一化
Layer 层归一化
Weight 权重
Local Response

BN: 1. 在批次维度上进行归一化。

2. 依赖批次大小。

3. 常用于 CNN

4. 稳定训练过程并降低 CNN 对初始参数的敏感性。

LN: 1. 在单个样本维度进行归一化。

2. 不依赖于批量大小。

3. 常用于 RNNs, Transformer, 或批量大小变化的任务。

循环神经网络

4. 在序列模型，批次估计不可靠时稳健。

