

第三章 概率密度函数的估计

2010-10-13

最小风险估计

□ 损失函数：把 θ 估计为 $\hat{\theta}$ 所造成的损失： $\lambda(\hat{\theta}, \theta)$

□ 期望风险：

$$\begin{aligned} R &= \int_{E^d} \int_{\Theta} \lambda(\hat{\theta}, \theta) p(\mathbf{x}, \theta) d\theta d\mathbf{x} \\ &= \int_{E^d} \int_{\Theta} \lambda(\hat{\theta}, \theta) p(\theta | \mathbf{x}) p(\mathbf{x}) d\theta d\mathbf{x} \\ &= \int_{E^d} p(\mathbf{x}) \int_{\Theta} \lambda(\hat{\theta}, \theta) p(\theta | \mathbf{x}) d\theta d\mathbf{x} \\ &= \int_{E^d} R(\hat{\theta} | \mathbf{x}) p(\mathbf{x}) d\mathbf{x}; \end{aligned}$$

□ 条件风险：

$$R(\hat{\theta} | \mathbf{x}) = \int_{\Theta} \lambda(\hat{\theta}, \theta) p(\theta | \mathbf{x}) d\theta.$$

最小风险估计

□ 最小化条件风险 $\longrightarrow$ 最小化期望风险。

□ 在有限样本集下，最小化经验风险

$$R(\hat{\theta} | \mathbf{K}) = \int_{\Theta} \lambda(\hat{\theta}, \theta) p(\theta | \mathbf{K}) d\theta.$$

□ 贝叶斯估计量：（在样本集 $\mathbf{K}$ 下）是条件风险（经验风险）最小的估计量 $\hat{\theta}_{\text{BE}}$ ，即

$$\hat{\theta}_{\text{BE}} = \arg \min_{\hat{\theta}} R(\hat{\theta} | \mathbf{K}).$$

最小风险估计

□ 定义平方误差损失函数 $\lambda(\hat{\theta}, \theta) = (\theta - \hat{\theta})^2$.

$$\begin{aligned} R(\hat{\theta} | \mathbf{x}) &= \int_{\Theta} \lambda(\hat{\theta}, \theta) p(\theta | \mathbf{x}) d\theta \\ &= \int_{\Theta} [\theta - E(\theta | \mathbf{x})]^2 p(\theta | \mathbf{x}) d\theta + \int_{\Theta} [E(\theta | \mathbf{x}) - \hat{\theta}]^2 p(\theta | \mathbf{x}) d\theta; \end{aligned}$$

□ 定理：如采用平方损失函数，则有

$$\hat{\theta}_{\text{BE}} = E[\theta | \mathbf{x}] = \int_{\Theta} \theta p(\theta | \mathbf{x}) d\theta;$$

□ 同理，在给定样本集 $\mathbf{K}$ 下， θ 的贝叶斯估计为

$$\hat{\theta}_{\text{BE}} = E[\theta | \mathbf{K}] = \int_{\Theta} \theta p(\theta | \mathbf{K}) d\theta;$$

最小风险估计的求解步骤

1. 确定 θ 的先验分布 $p(\theta)$;
2. 由样本集 $\mathbf{K} = \{\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_N\}$ 求出其联合分布：

$$p(\mathbf{K} | \theta) = \prod_{k=1}^N p(\mathbf{x}_k | \theta);$$

3. 计算 θ 的后验分布

$$p(\theta | \mathbf{K}) = \frac{p(\mathbf{K} | \theta) p(\theta)}{\int_{\Theta} p(\mathbf{K} | \theta) p(\theta) d\theta};$$

4. 计算贝叶斯估计

$$\hat{\theta}_{\text{BE}} = \int_{\Theta} \theta p(\theta | \mathbf{K}) d\theta.$$

一元正态分布的贝叶斯估计

□ 总体分布密度为：

$$p(x | \mu) \sim N(\mu, \sigma^2);$$

□ 均值 μ 未知， μ 的先验分布为：

$$p(\mu) \sim N(\mu_0, \sigma_0^2);$$

□ 样本集： $\mathbf{K} = \{\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_N\}$

□ 用贝叶斯估计方法求 μ 的估计量

一元正态分布的贝叶斯估计

□ 计算 μ 的后验分布

$$p(\mu | \mathbf{K}) = \frac{p(\mathbf{K} | \mu)p(\mu)}{p(\mathbf{K})}$$

$$= \alpha \prod_{k=1}^N p(x_k | \mu)p(\mu) \sim N(\mu_N, \sigma_N^2)$$

$$\mu_N = \frac{N\sigma_0^2}{N\sigma_0^2 + \sigma^2} m_N + \frac{\sigma^2}{N\sigma_0^2 + \sigma^2} \mu_0, \quad \sigma_N^2 = \frac{\sigma_0^2 \sigma^2}{N\sigma_0^2 + \sigma^2};$$

其中 $m_N = \frac{1}{N} \sum_{k=1}^N \mathbf{x}_k$ 为样本均值;

一元正态分布的贝叶斯估计

□ 计算 μ 的贝叶斯估计

$$\hat{\mu} = \int \mu p(\mu | \mathbf{K}) d\mu = \mu_N;$$

- 当 $N=0$ 时, $\hat{\mu}_{BE} = \mu_0$;
- 当 $N \rightarrow \infty$ 时, $\hat{\mu}_{BE} \rightarrow m_N$;
- 如 $\sigma_0^2 = 0$, 则 $\hat{\mu}_{BE} \equiv \mu_0$; 即先验知识可靠, 样本不起作用。
- 如 $\sigma_0^2 \gg \sigma^2$, 则 $\hat{\mu}_{BE} = m_N$; 即先验知识十分不确定, 完全依靠样本信息。

贝叶斯学习

基本思想

□ 利用 θ 的先验分布 $p(\theta)$ 及训练样本提供的信息 $p(\mathbf{K} | \theta)$, 求 θ 的后验分布 $p(\theta | \mathbf{K})$; 然后直接求解总体分布。

$$p(\mathbf{x} | \mathbf{K}) = \int p(\mathbf{x}, \theta | \mathbf{K}) d\theta$$

$$= \int p(\mathbf{x} | \theta) p(\theta | \mathbf{K}) d\theta;$$

- 将类条件概率密度 (总体分布) $p(\mathbf{x} | \mathbf{K})$ 和未知参数的后验概率密度 $p(\theta | \mathbf{K})$ 联系起来;
- 贝叶斯学习的结果与最大似然估计的结果近似:

$$p(\mathbf{x} | \mathbf{K}) \approx p(\mathbf{x} | \hat{\theta}_{ML}).$$

递推 (序贯) 后验概率

□ 考虑 $N > 1$ 个学习样本, 记样本集 $\mathbf{K}^N = \{\mathbf{x}_1, \dots, \mathbf{x}_N\}$;

$$p(\mathbf{K}^N | \theta) = p(\mathbf{x}_N | \theta) p(\mathbf{K}^{N-1} | \theta),$$

$$p(\theta | \mathbf{K}^N) = \frac{p(\mathbf{K}^N | \theta) p(\theta)}{\int_{\Theta} p(\mathbf{K}^N | \theta) p(\theta) d\theta}$$

$$= \frac{p(\mathbf{x}_N | \theta) p(\mathbf{K}^{N-1} | \theta) p(\theta) \cdot \frac{1}{p(\mathbf{K}^{N-1})}}{\int_{\Theta} p(\mathbf{x}_N | \theta) p(\mathbf{K}^{N-1} | \theta) p(\theta) d\theta \cdot \frac{1}{p(\mathbf{K}^{N-1})}}$$

$$= \frac{p(\mathbf{x}_N | \theta) p(\theta | \mathbf{K}^{N-1})}{\int_{\Theta} p(\mathbf{x}_N | \theta) p(\theta | \mathbf{K}^{N-1}) d\theta}.$$

递推贝叶斯学习

□ 设 $p(\theta | \mathbf{K}^0) = p(\theta)$, 当样本数目增多, 可得到后验概率密度函数序列:

$$p(\theta), p(\theta | \mathbf{x}_1), p(\theta | \mathbf{x}_1, \mathbf{x}_2), \dots$$

□ 如果此序列收敛于以真实数值为中心的 δ 函数, 则称样本分布具有 **贝叶斯学习 (Bayesian Learning) 性质**:

$$p(\theta | \mathbf{K}^{N \rightarrow \infty}) = \delta(\theta - \theta_0);$$

$$p(\mathbf{x} | \mathbf{K}^{N \rightarrow \infty}) = p(\mathbf{x} | \hat{\theta} = \theta_0) = p(\mathbf{x}).$$

一元正态分布的贝叶斯学习

□ 计算 μ 的后验分布

$$p(\mu | \mathbf{K}^N) = \frac{p(\mathbf{K}^N | \mu) p(\mu)}{p(\mathbf{K}^N)}$$

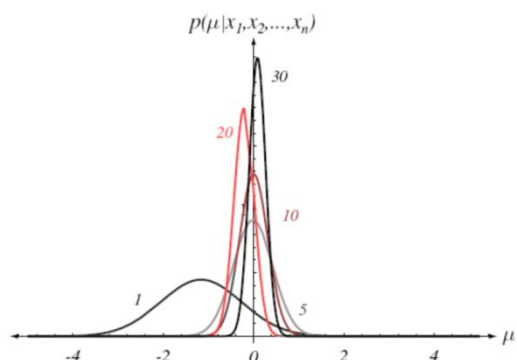
$$= \alpha \prod_{k=1}^N p(x_k | \mu) p(\mu) \sim N(\mu_N, \sigma_N^2);$$

$$\mu_N = \frac{N\sigma_0^2}{N\sigma_0^2 + \sigma^2} m_N + \frac{\sigma^2}{N\sigma_0^2 + \sigma^2} \mu_0, \quad \sigma_N^2 = \frac{\sigma_0^2 \sigma^2}{N\sigma_0^2 + \sigma^2};$$

其中 $m_N = \frac{1}{N} \sum_{k=1}^N \mathbf{x}_k$ 为样本均值;

当 $N \rightarrow \infty$ 时, $\sigma_N^2 \rightarrow 0$, $p(\mu | \mathbf{K}) \rightarrow \delta$ 函数。

一元正态分布的贝叶斯学习



一元正态分布的贝叶斯学习

□ 直接计算总体密度:

$$p(x | \mathbf{K}^N) = \int p(x | \mu) p(\mu | \mathbf{K}^N) d\mu$$

$$= \frac{1}{\sqrt{2\pi} \sqrt{\sigma^2 + \sigma_N^2}} \exp \left\{ -\frac{1}{2} \left(\frac{x - \mu_N}{\sqrt{\sigma^2 + \sigma_N^2}} \right)^2 \right\}$$

$$\sim N(\mu_N, \sigma^2 + \sigma_N^2).$$

均值 μ_N ,

方差由 σ^2 增为 $\sigma^2 + \sigma_N^2$ ----- 由于用了 μ 的估计值而不确定性增加

非监督参数估计

简介最大似然法

假设条件

1. 样本集 $\mathbf{K} = \{\mathbf{x}_1, \dots, \mathbf{x}_N\}$ 中的样本分属于 c 个类别, 但未知各样本所属类别;
2. 已知各类先验概率 $P(\omega_i)$, $i=1, \dots, c$; (有时也可未知, 一起估计)
3. 已知类条件概率密度形式 $p(\mathbf{x} | \omega_i, \theta_i)$, $i=1, \dots, c$;
4. 需估计未知的 c 个参数向量 $\theta_1, \theta_2, \dots, \theta_c$ 。

似然函数

□ 混合密度函数: 分量密度的线性组合

$$p(\mathbf{x} | \theta) = \sum_{i=1}^c \underbrace{p(\mathbf{x} | \omega_i, \theta_i)}_{\text{分量密度}} \underbrace{P(\omega_i)}_{\text{混合参数}}$$

□ 似然函数和对数似然函数:

$$l(\theta) = p(\mathbf{K} | \theta) = \prod_{i=1}^N p(\mathbf{x}_i | \theta)$$

$$H(\theta) = \ln[l(\theta)] = \sum_{i=1}^N \ln p(\mathbf{x}_i | \theta)$$

最大似然估计

19

$$\hat{\theta} = \arg \max_{\theta \in \Theta} \prod_{k=1}^N p(x_k | \theta) = \arg \max_{\theta \in \Theta} \sum_{k=1}^N \ln p(x_k | \theta);$$

可识别性

- 设 $\theta \neq \theta'$, 如对混合分布中每个 $\mathbf{x}$ 都有 $p(\mathbf{x}|\theta) \neq p(\mathbf{x}|\theta')$, 则称密度 $p(\mathbf{x}|\theta)$ 是可识别的;
- 大部分常见连续随机变量的分布密度函数都是可识别的; 离散随机变量的混合概率函数往往是不可识别的。

最大似然估计

20

计算问题: 求解微分方程组

$$\begin{aligned} \nabla_{\theta_i} H(\theta) &= \sum_{k=1}^N \frac{1}{p(x_k | \theta)} \nabla_{\theta_i} \left[\sum_{j=1}^c p(x_k | \omega_j, \theta_j) P(\omega_j) \right] \\ &= \sum_{k=1}^N \frac{1}{p(x_k | \theta)} \nabla_{\theta_i} [p(x_k | \omega_i, \theta_i) P(\omega_i)] \quad (\text{设 } \theta_i, \theta_j \text{ 独立}) \\ &= \sum_{k=1}^N P(\omega_i | x_k, \theta_i) \nabla_{\theta_i} \ln p(x_k | \omega_i, \theta_i) \end{aligned}$$

$$\text{其中后验概率} \quad P(\omega_i | x_k, \theta_i) = \frac{p(x_k | \omega_i, \theta_i) P(\omega_i)}{p(x_k | \theta)}$$

正态分布的非监督最大似然估计

21

均值向量 μ_i 未知, Σ_i , $P(\omega_i)$, c 已知

- 最大似然估计满足方程组

$$\sum_{k=1}^N \hat{P}(\omega_i | x_k, \hat{\theta}_i) \nabla_{\theta_i} \ln p(x_k | \omega_i, \hat{\theta}_i) = 0, \quad i = 1, \dots, c$$

代入正态分布

$$\sum_{k=1}^N P(\omega_i | x_k, \hat{\mu}_i) \hat{\mu}_i(j+1) = \frac{\sum_{k=1}^N P(\omega_i | x_k, \hat{\mu}_i(j)) x_k}{\sum_{k=1}^N P(\omega_i | x_k, \hat{\mu}_i(j))} \quad (\omega_i | x_k, \hat{\mu}_i)$$

参数估计总结

22

最大似然估计

$$\hat{\theta}_{\text{ML}} = \arg \max_{\theta} p(\mathbf{K} | \theta) = \arg \max_{\theta} \prod_{k=1}^N p(\mathbf{x}_k | \theta);$$

最大后验概率估计

$$\hat{\theta}_{\text{MAP}} = \arg \max_{\theta} p(\mathbf{K} | \theta) p(\theta) = \arg \max_{\theta} \sum_{k=1}^N \ln p(\mathbf{x}_k | \theta) + \ln p(\theta);$$

贝叶斯估计

$$p(\theta | \mathbf{K}) = \frac{p(\mathbf{K} | \theta) p(\theta)}{\int p(\mathbf{K} | \theta) p(\theta) d\theta}, \quad \hat{\theta}_{\text{BE}} = E(\theta | \mathbf{K}) = \int_{\Theta} \theta p(\theta | \mathbf{K}) d\theta;$$

贝叶斯学习

$$p(\mathbf{x} | \mathbf{K}) = \int p(\mathbf{x} | \theta) p(\theta | \mathbf{K}) d\theta.$$