

云南大学“211工程”项目资助

GAOJI XINHAO CHULI
YUANLI JI YINGYONG

高级信号处理

原理及应用

柏正尧 主 编

丁洪伟 余 映 副主编



科学出版社

云南大学“211工程”项目资助

高级信号处理原理及应用

柏正尧 主 编

丁洪伟 余 映 副主编



科 学 出 版 社

北 京

内 容 简 介

本书是信号处理的高级教程,详细介绍了信号处理的若干高级专题的基本理论,并提供了丰富的应用实例。全书共分10章,包括贝叶斯推理、隐马尔可夫模型、维纳滤波器、自适应滤波器、功率谱分析、主成分分析、独立成分分析、小波变换、Hilbert-Huang变换、盲解卷积和信道均衡。

本书适合于信号与信息处理、通信与信息系统、电路与系统、控制科学与工程、模式识别与智能系统、生物医学工程等专业的硕士研究生和博士研究生阅读,也可供相关领域的教师 and 广大科技工作者参考。

图书在版编目(CIP)数据

高级信号处理原理及应用 / 柏正尧主编. —北京: 科学出版社, 2013.10

ISBN 978-7-03-038693-9

I. ①高… II. ①柏… III. ①信号处理-研究 IV. ①TN911.7

中国版本图书馆 CIP 数据核字 (2013) 第 227997 号

责任编辑: 杨 岭 黄 桥 / 责任校对: 葛茂香

责任印制: 邱志强 / 封面设计: 墨创文化

科学出版社出版

北京东黄城根北街16号

邮政编码: 100717

<http://www.sciencep.com>

成都创新包装印刷厂印刷

科学出版社发行 各地新华书店经销

*

2013年10月第一版 开本: 787*1092 1/16

2013年10月第一次印刷 印张: 17.25

字数: 410千字

定价: 39.00元

(如有印装质量问题, 我社负责调换)

前 言

随着计算机技术和微电子技术的发展,信号处理技术已经发展为一门主流的技术。信号处理的应用范围非常广泛,包括多媒体技术、音频信号处理、视频信号处理、蜂窝移动通信、雷达系统、模式分析与识别、医学信号处理、人工智能和控制系统等。信号处理的理论与应用涉及信号过程模式与结构的辨识、建模和运用。

为了适应信号与信息处理等专业硕士研究生和博士研究生进一步学习信号处理理论、方法和算法的要求,我们编写了这本信号处理的高级教程。学习本书,一般要求硕士研究生和博士研究生在本科阶段已学习过信号与系统、数字信号处理等课程,具有一定的信号处理基础理论。

本书主要介绍高级信号处理的理论及应用,全书共分为 10 章。其中,第 1~4 章及第 9 章由柏正尧教授编写,第 5、8、10 章由丁洪伟副教授编写,第 6、7 章由余映博士编写。全书由柏正尧教授统稿。各章内容简介如下:

第 1 章介绍贝叶斯理论、经典估计方法、估计最大化方法、最小估计方差的 Cramér-Rao 界、贝叶斯分类、随机信号空间建模。

第 2 章介绍用于非平稳信号建模的隐马尔可夫模型理论、模型训练方法以及模型在解码和分类中的应用。

第 3 章介绍维纳滤波器理论、维纳滤波器的向量空间解释及维纳滤波器的频率响应。

第 4 章介绍自适应滤波器理论,包括卡尔曼滤波器的状态空间方程、非线性卡尔曼滤波器——扩展卡尔曼滤波器和无迹卡尔曼滤波器、递归最小二乘误差滤波器和最速下降法以及最小均方误差滤波器。

第 5 章介绍功率谱估计的方法,包括经典谱估计方法、参数谱估计方法以及非参数谱估计方法。

第 6 章介绍主成分分析方法,包括主成分分析的特征结构及几种常见的主成分分析算法。

第 7 章介绍独立成分分析方法,包括独立成分分析的定义、统计独立性、独立成分分析的估计原理及常见算法。

第 8 章介绍小波变换的理论,包括连续小波变换的定义和性质、离散小波变换、小波框架理论及性质、多分辨分析和 Mallat 算法。

第 9 章介绍 Hilbert-Huang 变换方法,包括瞬时频率的概念、经验模态分解、固有模态函数和 Hilbert 谱。

第 10 章介绍盲解卷积和信道均衡,包括盲解卷积的数学模型、准则及算法,信道均衡的原理。

本书的出版得到了云南大学“211 工程”项目资助,在此表示感谢。

云南大学信息学院原副院长赵东风教授生前对本书的编写给予了关心、支持和鼓励，谨以本书的出版来告慰赵副院长。

在本书编写过程中，研究生万娟、赵变芳、李娟、闫帅辉、尹立国、王峥、郑哲、南京、刘正纲等参与了资料收集、公式编辑、绘图等工作，对他们的工作表示感谢。

在本书编写过程中，作者参考了国内外专家和学者的论文、专著等文献资料，在此一并表示衷心感谢。

感谢科学出版社黄桥编辑及其他编辑同志的辛勤劳动，是他们认真细致的工作，才使得本书能付梓出版。

信号处理的理论、方法、算法仍在不断发展中，应用领域也在不断拓展，尽管编者努力尝试介绍信号处理的高级理论和方法，但书中难免存在不全面、不妥甚至错漏之处，请广大读者批评指正。

柏正尧

2013年5月于云南大学东陆园

目 录

第 1 章 贝叶斯推理	1
1.1 贝叶斯估计理论	1
1.1.1 贝叶斯定理	2
1.1.2 贝叶斯推理的定义	2
1.1.3 动态概率模型估计方法	2
1.1.4 估计的性能指标	3
1.2 贝叶斯估计	4
1.2.1 最大后验估计	5
1.2.2 最大似然估计	5
1.2.3 最小均方误差估计	7
1.2.4 最小平均绝对值误差估计	8
1.2.5 先验概率密度对估计偏差和方差的影响	9
1.3 最大期望算法	13
1.3.1 似然函数的最大期望	14
1.3.2 最大期望算法的推导及收敛性	14
1.4 最小估计方差的 Cramér-Rao 界	15
1.4.1 随机参数的 Cramér-Rao 界	16
1.4.2 参数向量的 Cramér-Rao 约束	17
1.5 高斯混合模型	17
1.6 贝叶斯分类	19
1.6.1 二元分类	19
1.6.2 分类误差	20
1.6.3 离散值参数的贝叶斯分类器	21
1.6.4 有限状态过程的贝叶斯分类	22
1.7 随机过程空间建模	24
1.7.1 随机过程的向量量化	24
1.7.2 使用集簇高斯模型的向量量化	24
1.7.3 向量量化器设计—— K 均值聚类	25
第 2 章 隐马尔可夫模型	27
2.1 非平稳过程的统计模型	27
2.2 HMM 概述	28
2.2.1 HMM 与贝叶斯模型	29

2.2.2	HMM 的参数	29
2.2.3	状态观测概率模型	30
2.2.4	状态转移概率	31
2.2.5	状态时间网格图	31
2.3	HMM 训练	32
2.3.1	前向-后向概率计算	32
2.3.2	Baum-Welch 模型再估计	33
2.3.3	离散密度 HMM 训练	34
2.3.4	连续密度 HMM	35
2.3.5	高斯混合概率密度 HMM	36
2.4	用 HMM 进行信号解码	37
2.4.1	语音识别应用	37
2.4.2	维特比算法	38
2.5	噪声的 HMM	39
2.5.1	噪声信号的估计	39
2.5.2	信号与噪声模型的组合与分解	40
2.5.3	HMM 组合	41
2.5.4	信号和噪声状态序列的分解	42
2.5.5	基于 HMM 的维纳滤波器	42
2.5.6	噪声特征建模	43
第 3 章	维纳滤波器	45
3.1	最小二乘估计——维纳滤波器	45
3.1.1	维纳滤波方程的推导	45
3.1.2	输入的自相关及输入与期望信号的互相关的计算	47
3.2	维纳滤波器的分块数据公式	48
3.3	最小均方误差方程的 QR 分解	49
3.4	维纳滤波器的向量空间投影描述	50
3.5	最小均方误差信号分析	51
3.6	频域中的维纳滤波器	52
3.7	维纳滤波器的应用	53
3.7.1	维纳滤波器用于加性噪声抑制	53
3.7.2	平方根维纳滤波	54
3.7.3	维纳信道均衡器	55
3.7.4	多信道/多传感器系统中的信号时间校正	56
3.8	维纳滤波器的实现问题	56
3.8.1	维纳滤波器阶数的选择	57
3.8.2	维纳滤波器的改进	58
第 4 章	自适应滤波器	59
4.1	自适应滤波器的概念	59

4.2 状态空间卡尔曼滤波器	60
4.2.1 卡尔曼滤波算法的推导	61
4.2.2 用递归贝叶斯公式表示卡尔曼滤波器	63
4.2.3 卡尔曼滤波器的马尔可夫性	64
4.3 扩展卡尔曼滤波器	65
4.4 无迹卡尔曼滤波器	67
4.5 样本自适应滤波器	69
4.6 递归最小二乘误差自适应滤波器	69
4.6.1 矩阵求逆引理	71
4.6.2 滤波器系数的递归时间更新	71
4.7 最速下降法	73
4.7.1 收敛速度	74
4.7.2 向量值自适应步长	75
4.8 最小均方误差滤波器	75
4.8.1 漏溢最小均方误差算法	76
4.8.2 归一化最小均方误差算法	76
第5章 功率谱分析	79
5.1 概 述	79
5.2 经典谱估计	80
5.2.1 相关图法	80
5.2.2 相关图法功率谱估计——有偏自相关函数估计法	82
5.2.3 周期图法	83
5.2.4 经典谱估计方法改进	87
5.3 参数谱估计	94
5.3.1 信号建模	94
5.3.2 模型参数和自相关函数之间的关系	95
5.3.3 AR 模型谱估计的性质	97
5.3.4 AR 谱估计的方法	100
5.3.5 AR 模型阶次的选择	112
5.4 非参数模型法估计功率谱	114
5.4.1 全向天线波束下倾	114
5.4.2 天线波束方向估计	115
第6章 主成分分析	123
6.1 概 述	124
6.2 主成分分析的特征结构	125
6.3 常 见 算 法	126
6.3.1 基于 Hebb 学习规则的随机梯度算法	127
6.3.2 广义的 Hebb 算法	128
6.3.3 改进的增量计算方法	129

6.3.4 混合算法	130
6.4 应用举例	130
6.4.1 在图像压缩中的应用	130
6.4.2 在模式识别中的应用	131
6.4.3 产生图像的视觉显著图	132
第7章 独立成分分析	134
7.1 概 述	134
7.1.1 独立成分分析的定义	134
7.1.2 统计独立性	135
7.2 独立成分分析的估计原理	136
7.2.1 非高斯性最大化	136
7.2.2 互信息最小化	139
7.2.3 最大似然估计	140
7.3 独立成分分析的生物意义	141
7.4 常见算法	143
7.4.1 数据的预处理	143
7.4.2 Jutten-Herault 算法	144
7.4.3 Infomax 算法	144
7.4.4 FastICA 算法	145
7.5 应用举例	146
7.5.1 混合信号的分离	146
7.5.2 自然图像的降噪	147
第8章 小波变换	151
8.1 连续小波变换	151
8.1.1 连续小波变换的定义	151
8.1.2 连续小波变换的一些性质	155
8.1.3 小波的基本要求	157
8.1.4 几种常用的基本小波	161
8.1.5 连续小波变换的计算及实现	164
8.2 离散小波变换	166
8.2.1 离散小波及离散小波变换	166
8.2.2 小波框架与离散小波反变换	168
8.3 小波变换的 Mallat 算法	171
8.3.1 多分辨率分析	171
8.3.2 尺度函数和小波函数	173
8.3.3 Mallat 快速算法	179
8.4 小波变换的应用举例	183
8.4.1 小波去噪	183
8.4.2 GPS 信号处理	185

第 9 章 Hilbert-Huang 变换	187
9.1 瞬时频率	187
9.2 固有模态函数	188
9.3 经验模态分解	189
9.4 完备性和正交性	195
9.5 Hilbert 谱	197
9.6 Hilbert 谱验证和校正	202
第 10 章 盲解卷积和信道均衡	211
10.1 盲解卷积	211
10.1.1 盲解卷积简介	211
10.1.2 盲解卷积的数学模型	211
10.1.3 盲解卷积准则	213
10.1.4 算法	215
10.2 信道均衡	252
10.2.1 信道均衡简介	252
10.2.2 基本原理	252
主要参考文献	262

第 1 章 贝叶斯推理

在信号处理中，推理就是从一组观测值中寻求观测信号或参数的过程。由于观测值通常是有噪声或不完全的，所以，从观测值得出的结论，其确定性和准确性依赖于观测的质量和推理方法的有效性。贝叶斯推理是根据一组观测值对信号或参数的未知值进行预测、估计和分类的概率推理，它将观测信号中的证据及其似然函数、随机过程概率分布的先验知识和误差代价函数结合在一起，对一个信号随机过程进行预测或估计。

贝叶斯推理以贝叶斯风险函数最小化为基础，包括在给定的观测条件下未知参数的概率模型和误差代价函数。贝叶斯方法包括最大后验(MAP)估计、最大似然(ML)估计、最小均方误差(MMSE)估计，以及作为其特殊情况的最小平均绝对误差(MAVE)，如图 1.1 所示。

本章首先介绍贝叶斯估计理论的基本概念，然后介绍几种经典估计方法，并用统计学方法对其性能进行评价，最后讨论离散或有限状态信号的贝叶斯分类和 K 均值聚类方法。

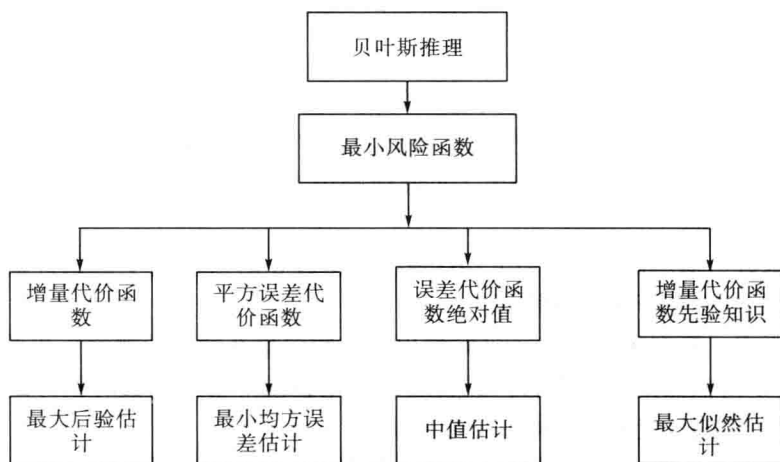


图 1.1 贝叶斯推理

1.1 贝叶斯估计理论

估计理论研究从一个观测信号确定未知参数向量的最优估计，或恢复一个受噪声污染或失真退化的信号。贝叶斯估计理论为统计估计方法的推导提供了一般的理论框架。统计估计法的各种形式可以看做贝叶斯估计的特殊情况。

1.1.1 贝叶斯定理

贝叶斯定理(或贝叶斯准则): 对于给定的观测向量 $\mathbf{y}$ 的随机参数向量 $\boldsymbol{\theta}$ 值的估计, 已知 $\mathbf{y}$ 值时, 参数向量 $\boldsymbol{\theta}$ 的后验概率密度函数可表示为

$$f_{\boldsymbol{\theta}|\mathbf{y}}(\boldsymbol{\theta} | \mathbf{y}) = \frac{1}{f_{\mathbf{y}}(\mathbf{y})} f_{\mathbf{y}|\boldsymbol{\theta}}(\mathbf{y} | \boldsymbol{\theta}) f_{\boldsymbol{\theta}}(\boldsymbol{\theta}) \quad (1.1)$$

式中, $f_{\mathbf{y}|\boldsymbol{\theta}}(\mathbf{y} | \boldsymbol{\theta})$ 为似然函数; $f_{\boldsymbol{\theta}}(\boldsymbol{\theta})$ 为参数向量的先验概率密度函数。

对于给定的观测值, $f_{\mathbf{y}}(\mathbf{y}) = \int f_{\mathbf{y}|\boldsymbol{\theta}}(\mathbf{y} | \boldsymbol{\theta}) f_{\boldsymbol{\theta}}(\boldsymbol{\theta}) d\boldsymbol{\theta}$ 是一个常数。贝叶斯定理说明, 后验概率与似然函数和先验概率的乘积成正比。

似然函数和先验概率密度函数对后验概率密度函数的影响程度取决于它们的函数形状。一般来说, 概率密度函数的形状越尖, 它对估计结果的影响就越大。而均匀分布的先验概率密度函数除那些超出概率密度函数范围的值外, 对后验概率密度函数没有什么影响。

1.1.2 贝叶斯推理的定义

贝叶斯推理可以由风险函数最小化和未知参数向量值的误差代价平均定义, 即

$$\text{Risk}(\hat{\boldsymbol{\theta}}) = \iint_{\boldsymbol{\theta}, \mathbf{y}} \text{Cost}(\hat{\boldsymbol{\theta}}, \boldsymbol{\theta}) f_{\mathbf{y}|\boldsymbol{\theta}}(\mathbf{y} | \boldsymbol{\theta}) f_{\boldsymbol{\theta}}(\boldsymbol{\theta}) d\mathbf{y} d\boldsymbol{\theta} \quad (1.2)$$

贝叶斯推理包括以下几个要素:

(1) 贝叶斯风险。估计未知参数向量 $\boldsymbol{\theta}$ 的估计值 $\hat{\boldsymbol{\theta}}$ 的风险与误差代价函数的概率有关。由于参数向量 $\boldsymbol{\theta}$ 的真实值是未知的, 因此需要对 $\boldsymbol{\theta}$ 的所有可能取值的代价求平均。

(2) 代价函数。误差代价函数决定了贝叶斯解决方案的类型。最常用的一个代价函数是最小均方误差代价。

(3) 似然函数。似然函数 $f_{\mathbf{y}|\boldsymbol{\theta}}(\mathbf{y} | \boldsymbol{\theta})$ 是给定参数 $\boldsymbol{\theta}$ 时观测值 $\mathbf{y}$ 的条件概率, 它给出了由参数 $\boldsymbol{\theta}$ 决定观测信号 $\mathbf{y}$ 的可能性。

(4) 先验概率。概率密度函数 $f_{\boldsymbol{\theta}}(\boldsymbol{\theta})$ 给出了参数向量 $\boldsymbol{\theta}$ 的先验概率。先验概率函数降低了参数向量 $\boldsymbol{\theta}$ 对似然函数的影响。似然函数和先验概率密度函数对估计的影响程度取决于信噪比、观测值长度和先验概率函数形状等。

(5) 后验概率。后验概率与似然函数和先验概率密度函数的乘积成正比。

1.1.3 动态概率模型估计方法

优化估计算法, 如卡尔曼滤波器使用了观测信号的动态概率模型。动态预测模型刻画了信号的相关结构, 是对信号当前值和未来值与信号过去的变化轨迹及输入依赖关系的建模。概率统计模型用统计特性, 如均值和方差, 刻画一个随机信号的不同实现构成的空间。条件概率模型除对随机波动建模外, 也可以对信号与其过去值和其他相关参数值或过程的关系进行建模。有限状态模型和卡尔曼滤波器是动态概率模型的两个例子。

对受噪声污染的观测向量 $\mathbf{y} = [y(0), y(1), \dots, y(N-1)]$, 其 P 维参数向量 $\boldsymbol{\theta} = [\theta_0, \theta_1, \dots, \theta_{P-1}]$ 的估计可以用如下模型表示:

$$\mathbf{y} = \mathbf{x} + \mathbf{n} = h(\boldsymbol{\theta}, \mathbf{e}) + \mathbf{n} \quad (1.3)$$

如图 1.2 所示, 假设纯净信号 $\mathbf{x}$ 是预测模型 $h(\cdot)$ 的输出, $\mathbf{e}$ 为预测模型的随机输入, $\boldsymbol{\theta}$ 为其参数向量, $\mathbf{n}$ 是加性随机噪声向量, 则随机输入 $\mathbf{e}$ 、参数向量 $\boldsymbol{\theta}$ 和加性随机噪声向量 $\mathbf{n}$ 的概率密度函数分别为 $f_E(\cdot)$ 、 $f_{\boldsymbol{\theta}}(\cdot)$ 和 $f_N(\cdot)$ 。概率密度函数模型常采用高斯模型。

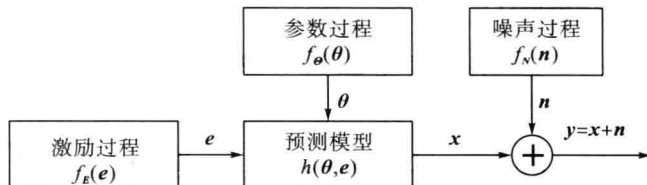


图 1.2 用预测模型 $h(\cdot)$ 表示随机过程 $\mathbf{y}$

随机过程的预测统计模型决定了未知参数的估计, 这些未知参数与模型参数和观测值的先验分布是最一致的。一般来说, 如果模型准确地刻画了观测值和参数过程, 则估计过程中可利用的信息越多, 估计结果就越可靠。反之, 如果模型不准确, 则估计结果不可靠。

1.1.4 估计的性能指标

通过 N 个观测样本 $\mathbf{y}$ 估计参数向量 $\boldsymbol{\theta}$ 时, 需要对不同估计值的性能进行量化和比较。一般来说, 参数向量是观测向量的函数, 参数估计的结果取决于估计方法、观测值和先验知识。常用估计误差的均值和方差来表征估计的性能。评价估计性能最常用的指标包括: ①估计期望值 $E[\hat{\boldsymbol{\theta}}]$; ②估计偏差 $E[\hat{\boldsymbol{\theta}} - \boldsymbol{\theta}] = E[\hat{\boldsymbol{\theta}}] - \boldsymbol{\theta}$; ③估计协方差 $\text{Cov}(\hat{\boldsymbol{\theta}}) = E[(\hat{\boldsymbol{\theta}} - E[\hat{\boldsymbol{\theta}}])(\hat{\boldsymbol{\theta}} - E[\hat{\boldsymbol{\theta}}])^T]$

一个最优估计的偏差为零, 估计误差的协方差最小, 它应该具有如下优良特性:

(1) 无偏估计。如果估计期望等于真实参数值, 则 $\boldsymbol{\theta}$ 的估计值是无偏的, 即

$$E[\hat{\boldsymbol{\theta}}] = \boldsymbol{\theta} \quad (1.4)$$

如果随着观测长度 N 增加有

$$\lim_{N \rightarrow \infty} E[\hat{\boldsymbol{\theta}}] = \boldsymbol{\theta} \quad (1.5)$$

则估计结果是渐近无偏的。

(2) 有效估计。如果与所有其他无偏估计 $\hat{\boldsymbol{\theta}}$ 相比具有最小协方差矩阵, 则 $\boldsymbol{\theta}$ 的无偏估计是有效估计, 即

$$\text{Cov}[\hat{\boldsymbol{\theta}}_{\text{Efficient}}] \leq \text{Cov}[\hat{\boldsymbol{\theta}}] \quad (1.6)$$

其中, $\hat{\boldsymbol{\theta}}$ 是 $\boldsymbol{\theta}$ 的其他任何估计量。

(3) 一致估计。如果观测样本数 N 增加能改善估计值, 当 N 变为无穷大时估计值 $\hat{\boldsymbol{\theta}}$ 依概率收敛于真实值 $\boldsymbol{\theta}$, 则估计是一致的, 即

$$\lim_{N \rightarrow \infty} P[|\hat{\boldsymbol{\theta}} - \boldsymbol{\theta}| > \epsilon] = 0 \quad (1.7)$$

其中, ϵ 是任意小的正数。

例 1.1 信号均值和方差的估计。遍历随机过程 y 的 N 个观测样本为 $[y(0), \dots, y(N-1)]$, 均值 μ_y 和方差 σ_y^2 的时间平均估计偏差可以计算如下:

$$\hat{\mu}_y = \frac{1}{N} \sum_{m=0}^{N-1} y(m) \quad (1.8)$$

$$\hat{\sigma}_y^2 = \frac{1}{N} \sum_{m=0}^{N-1} [y(m) - \hat{\mu}_y]^2 \quad (1.9)$$

容易证明, $\hat{\mu}_y$ 是一个无偏估计, 即

$$E[\hat{\mu}_y] = \frac{1}{N} \sum_{m=0}^{N-1} E[y(m)] = \mu_y \quad (1.10)$$

方差 $\hat{\sigma}_y^2$ 估计的期望为

$$\begin{aligned} E[\hat{\sigma}_y^2] &= E\left[\frac{1}{N} \sum_{m=0}^{N-1} \left(y(m) - \frac{1}{N} \sum_{k=0}^{N-1} y(k)\right)^2\right] \\ &= \sigma_y^2 - \frac{2}{N} \sigma_y^2 + \frac{1}{N} \sigma_y^2 \\ &= \sigma_y^2 - \frac{1}{N} \sigma_y^2 \end{aligned} \quad (1.11)$$

从式(1.11)可以看出, 方差估计的偏差反比于观测样本数 N , 当 N 趋于无穷大时偏差消失, 因此估计是渐近无偏的。在一般情况下, 估计的偏差和方差随着观测样本数 N 的增加和模型的改进而降低。图 1.3 说明了渐进无偏估计的方差和偏差分布与观测样本数 N 的关系。

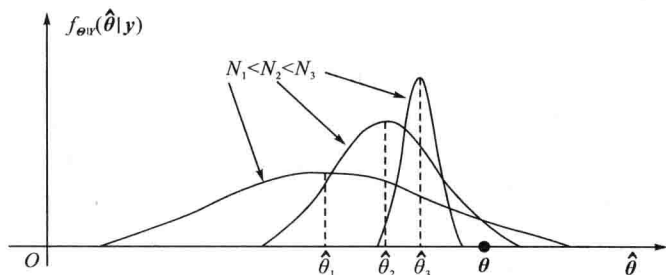


图 1.3 参数 θ 的渐进无偏估计的偏差和方差随观测样本数的变化

1.2 贝叶斯估计

根据观测向量 y 对参数向量 θ 进行贝叶斯估计, 可以用贝叶斯风险最小平均代价的误差函数来定义:

$$\begin{aligned} R(\hat{\theta}) &= E[C(\hat{\theta}, \theta)] = \iint_{\theta, y} C(\hat{\theta}, \theta) f_{Y, \theta}(y, \theta) dy d\theta \\ &= \iint_{\theta, y} C(\hat{\theta}, \theta) f_{Y|\theta}(y|\theta) f_{\theta}(\theta) dy d\theta \end{aligned} \quad (1.12)$$

其中, 误差代价函数 $C(\hat{\theta}, \theta)$ 可对各种结果进行适当加权以得到所期望的估计。贝叶斯风险函数是在参数向量 θ 和观测值 y 所有取值空间上进行平均。

对于给定的观测向量 y , $f_{\theta|Y}(\theta|y)$ 是一个常数, 对风险最小化过程没有影响。因

此, 将 $f_{\theta, Y}(\theta, y) = f_{\theta|Y}(\theta|y)f_Y(y)$ 代入式(1.12), 可以得到条件风险函数为

$$R(\hat{\theta}|y) = \int_{\theta} C(\hat{\theta}, \theta) f_{\theta|Y}(\theta|y) d\theta \quad (1.13)$$

由此可得到 θ 的贝叶斯估计 $\hat{\theta}$ (即最小风险参数向量) 为

$$\hat{\theta}_{\text{Bayes}} = \arg \min_{\hat{\theta}} R(\hat{\theta}|y) = \arg \min_{\hat{\theta}} \left[\int_{\theta} C(\hat{\theta}, \theta) f_{\theta|Y}(\theta|y) d\theta \right] \quad (1.14)$$

利用贝叶斯准则, 式(1.14)可写成:

$$\hat{\theta}_{\text{Bayes}} = \arg \min_{\hat{\theta}} \left[\int_{\theta} C(\hat{\theta}, \theta) f_{Y|\theta}(y|\theta) f_{\theta}(\theta) d\theta \right] \quad (1.15)$$

如果风险函数是可微的, 并有一个定义明确的最小值, 则可以得到贝叶斯估计为

$$\hat{\theta}_{\text{Bayes}} = \arg \text{zero} \frac{\partial R(\hat{\theta}|y)}{\partial \hat{\theta}} = \arg \text{zero} \left[\frac{\partial}{\partial \hat{\theta}} \int_{\theta} C(\hat{\theta}, \theta) f_{Y|\theta}(y|\theta) f_{\theta}(\theta) d\theta \right] \quad (1.16)$$

1.2.1 最大后验估计

最大后验估计 $\hat{\theta}_{\text{MAP}}$ 是由参数向量后验概率密度函数的最大化得到的。它是具有均匀代价函数的贝叶斯估计, 其代价函数是如图 1.4 所示的陷波形状, 定义为

$$C(\hat{\theta}, \theta) = 1 - \delta(\hat{\theta} - \theta) \quad (1.17)$$

其中, $\delta(\hat{\theta}, \theta)$ 是克罗内克 δ 函数。

将代价函数代入贝叶斯风险公式中得到:

$$R_{\text{MAP}}(\hat{\theta}|y) = \int_{\theta} [1 - \delta(\hat{\theta} - \theta)] f_{\theta|Y}(\theta|y) d\theta = 1 - f_{\theta|Y}(\hat{\theta}|y) \quad (1.18)$$

可以看出, 参数向量 θ 的最大后验估计是由贝叶斯风险最小化或后验函数的最大化得到的, 即

$$\hat{\theta}_{\text{MAP}} = \arg \max_{\hat{\theta}} f_{\theta|Y}(\hat{\theta}|y) = \arg \max_{\hat{\theta}} [f_{Y|\theta}(y|\hat{\theta}) f_{\theta}(\hat{\theta})] \quad (1.19)$$

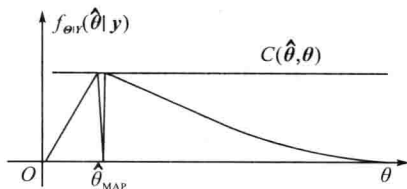


图 1.4 最大后验估计的贝叶斯代价函数

1.2.2 最大似然估计

最大似然估计 $\hat{\theta}_{\text{ML}}$ 是使似然函数 $f_{Y|\theta}(y|\theta)$ 最大的参数向量。其代价函数为陷波形状, 参数先验概率密度函数为均匀分布时的贝叶斯估计:

$$R_{\text{ML}}(\hat{\theta}|y) = \int_{\theta} [1 - \delta(\hat{\theta} - \theta)] f_{Y|\theta}(y|\theta) f_{\theta}(\theta) d\theta = \text{const.} [1 - f_{Y|\theta}(y|\hat{\theta})] \quad (1.20)$$

其中, 先验概率密度函数 $f_{\theta}(\theta) = \text{常数}$ 。

最大似然估计和最大后验估计的主要区别在于,前者假定 θ 的先验概率密度函数是均匀分布的。先验概率均匀分布除可用于真正均匀概率密度函数的建模外,还可用于先验概率密度函数的参数未知或参数是一个未知的常数时的建模。

由式(1.20)可知,风险函数最小化也可以由似然函数最大化实现,即

$$\hat{\theta}_{\text{ML}} = \arg \max_{\theta} f_{\mathbf{y}|\theta}(\mathbf{y}|\theta) \quad (1.21)$$

在实际应用中,利用对数似然函数的最大化计算最大似然估计更方便,即

$$\hat{\theta}_{\text{ML}} = \arg \max_{\theta} \log f_{\mathbf{y}|\theta}(\mathbf{y}|\theta) \quad (1.22)$$

对数似然函数具有良好的性质,如对数似然函数具有与似然函数相同的拐点,多个独立变量的联合对数似然函数等于各个变量似然函数之和,对数似然函数的动态范围不像似然函数那样会导致计算值溢出。

例 1.2 高斯过程的均值和方差的最大似然估计。设 P 维高斯向量过程为 $y(m)$,观测向量为 $[y(0), \dots, y(N-1)]$,现对该过程的均值向量和协方差矩阵进行最大似然估计。假设观测向量之间不相关,则观测向量序列的概率密度函数为

$$f_{\mathbf{y}}(y(0), \dots, y(N-1)) = \prod_{m=0}^{N-1} \frac{1}{(2\pi)^{P/2} |\Sigma_{yy}|^{1/2}} \exp \left\{ -\frac{1}{2} [y(m) - \mu_y]^T \Sigma_{yy}^{-1} [y(m) - \mu_y] \right\} \quad (1.23)$$

相应的对数似然函数为

$$\ln f_{\mathbf{y}}(y(0), \dots, y(N-1)) = \sum_{m=0}^{N-1} \left\{ -\frac{P}{2} \ln(2\pi) - \frac{1}{2} \ln |\Sigma_{yy}| - \frac{1}{2} [y(m) - \mu_y]^T \Sigma_{yy}^{-1} [y(m) - \mu_y] \right\} \quad (1.24)$$

求对数似然函数对均值向量的偏导数并令其等于零:

$$\frac{\partial \ln f_{\mathbf{y}}(y(0), \dots, y(N-1))}{\partial \mu_y} = \sum_{m=0}^{N-1} [-2 \Sigma_{yy}^{-1} y(m) + 2 \Sigma_{yy}^{-1} \mu_y] = 0 \quad (1.25)$$

由此可以得到均值的估计:

$$\hat{\mu}_y = \frac{1}{N} \sum_{m=0}^{N-1} y(m) \quad (1.26)$$

为了获得协方差矩阵的最大似然估计,取对数似然函数关于协方差矩阵的偏导数并令其等于零:

$$\frac{\partial \ln f_{\mathbf{y}}(y(0), \dots, y(N-1))}{\partial \Sigma_{yy}^{-1}} = \sum_{m=0}^{N-1} \left\{ \frac{1}{2} \Sigma_{yy} - \frac{1}{2} [y(m) - \mu_y][y(m) - \mu_y]^T \right\} = 0 \quad (1.27)$$

得到协方差矩阵的估计值:

$$\hat{\Sigma}_{yy} = \frac{1}{N} \sum_{m=0}^{N-1} [y(m) - \hat{\mu}_y][y(m) - \hat{\mu}_y]^T \quad (1.28)$$

例 1.3 一个随机高斯参数的最大似然估计和最大后验估计。利用 N 维观测向量 $\mathbf{y}$ 估计 P 维随机参数向量 θ 。假设信号向量 $\mathbf{y}$ 和参数向量 θ 的关系是一个线性模型,则

$$\mathbf{y} = \mathbf{G}\theta + \mathbf{e} \quad (1.29)$$

其中, $\mathbf{e}$ 是一个随机激励信号输入; $\mathbf{G}$ 是一个数据矩阵。

例如,对于自回归过程, $\mathbf{G}$ 由 $\mathbf{y}$ 过去的取值组成。在已知观测向量 $\mathbf{y}$ 的情况下,参数向量 θ 的概率密度函数可利用贝叶斯准则描述为

$$f_{\boldsymbol{\theta}|\mathbf{y}}(\boldsymbol{\theta}|\mathbf{y}) = \frac{1}{f_{\mathbf{y}}(\mathbf{y})} f_{\mathbf{y}|\boldsymbol{\theta}}(\mathbf{y}|\boldsymbol{\theta}) f_{\boldsymbol{\theta}}(\boldsymbol{\theta}) \quad (1.30)$$

假设式(1.29)中的 $\mathbf{G}$ 是已知的, 则在给定参数向量 $\boldsymbol{\theta}$ 的条件下信号 $\mathbf{y}$ 的似然函数是随机向量的概率密度分布函数:

$$f_{\mathbf{y}|\boldsymbol{\theta}}(\mathbf{y}|\boldsymbol{\theta}) = f_{\mathbf{E}}(\mathbf{e} = \mathbf{y} - \mathbf{G}\boldsymbol{\theta}) \quad (1.31)$$

进一步假设输入信号向量 $\mathbf{e}$ 是具有零均值和对角协方差矩阵的高斯分布随机过程, 参数向量 $\boldsymbol{\theta}$ 也是一个均值为 $\mu_{\boldsymbol{\theta}}$, 协方差矩阵为 $\Sigma_{\boldsymbol{\theta}\boldsymbol{\theta}}$ 的高斯过程。因此, 似然函数可以写成:

$$f_{\mathbf{y}|\boldsymbol{\theta}}(\mathbf{y}|\boldsymbol{\theta}) = f_{\mathbf{E}}(\mathbf{e}) = \frac{1}{(2\pi\sigma_e^2)^{N/2}} \exp\left[-\frac{1}{2\sigma_e^2}(\mathbf{y} - \mathbf{G}\boldsymbol{\theta})^T(\mathbf{y} - \mathbf{G}\boldsymbol{\theta})\right] \quad (1.32)$$

$$f_{\boldsymbol{\theta}}(\boldsymbol{\theta}) = \frac{1}{(2\pi)^{P/2} |\Sigma_{\boldsymbol{\theta}\boldsymbol{\theta}}|^{1/2}} \exp\left[-\frac{1}{2}(\boldsymbol{\theta} - \mu_{\boldsymbol{\theta}})^T \Sigma_{\boldsymbol{\theta}\boldsymbol{\theta}}^{-1}(\boldsymbol{\theta} - \mu_{\boldsymbol{\theta}})\right] \quad (1.33)$$

对数似然函数 $\ln[f_{\mathbf{y}|\boldsymbol{\theta}}(\mathbf{y}|\boldsymbol{\theta})]$ 取得最大值时的 $\boldsymbol{\theta}$ 值就是最大似然估计值:

$$\hat{\boldsymbol{\theta}}_{\text{ML}}(\mathbf{y}) = (\mathbf{G}^T \mathbf{G})^{-1} \mathbf{G}^T \mathbf{y} \quad (1.34)$$

为了获得最大后验估计值, 首先将式(1.32)和式(1.33)代入式(1.30), 得到后验概率密度函数:

$$\begin{aligned} f_{\boldsymbol{\theta}|\mathbf{y}}(\boldsymbol{\theta}|\mathbf{y}) &= \frac{1}{f_{\mathbf{y}}(\mathbf{y})} \frac{1}{(2\pi\sigma_e^2)^{N/2}} \frac{1}{(2\pi)^{P/2} |\Sigma_{\boldsymbol{\theta}\boldsymbol{\theta}}|^{1/2}} \\ &\times \exp\left(-\frac{1}{2\sigma_e^2}(\mathbf{y} - \mathbf{G}\boldsymbol{\theta})^T(\mathbf{y} - \mathbf{G}\boldsymbol{\theta}) - \frac{1}{2}(\boldsymbol{\theta} - \mu_{\boldsymbol{\theta}})^T \Sigma_{\boldsymbol{\theta}\boldsymbol{\theta}}^{-1}(\boldsymbol{\theta} - \mu_{\boldsymbol{\theta}})\right) \end{aligned} \quad (1.35)$$

将对数似然函数 $\ln f_{\boldsymbol{\theta}|\mathbf{y}}(\boldsymbol{\theta}|\mathbf{y})$ 求偏导数并令其为零, 即可得到最大后验估计值:

$$\hat{\boldsymbol{\theta}}_{\text{MAP}}(\mathbf{y}) = (\mathbf{G}^T \mathbf{G} + \sigma_e^2 \Sigma_{\boldsymbol{\theta}\boldsymbol{\theta}}^{-1})^{-1} (\mathbf{G}^T \mathbf{y} + \sigma_e^2 \Sigma_{\boldsymbol{\theta}\boldsymbol{\theta}}^{-1} \mu_{\boldsymbol{\theta}}) \quad (1.36)$$

可以看到, 随着高斯分布参数的协方差的增加, 即协方差矩阵的逆矩阵 $\Sigma_{\boldsymbol{\theta}\boldsymbol{\theta}}^{-1} \rightarrow 0$, 高斯先验概率分布趋于均匀的先验概率分布, 最大后验估计值趋近于最大似然估计值。反过来, 若参数向量 $\boldsymbol{\theta}$ 的概率密度函数越尖锐, 即 $\Sigma_{\boldsymbol{\theta}\boldsymbol{\theta}} \rightarrow 0$, 则最大后验估计值趋近于先验概率密度函数的均值 $\mu_{\boldsymbol{\theta}}$ 。

1.2.3 最小均方误差估计

均方误差代价函数定义为

$$R_{\text{MSE}}(\hat{\boldsymbol{\theta}}|\mathbf{y}) = E[(\hat{\boldsymbol{\theta}} - \boldsymbol{\theta})^2|\mathbf{y}] = \int_{\boldsymbol{\theta}} (\hat{\boldsymbol{\theta}} - \boldsymbol{\theta})^2 f_{\boldsymbol{\theta}|\mathbf{y}}(\boldsymbol{\theta}|\mathbf{y}) d\boldsymbol{\theta} \quad (1.37)$$

贝叶斯最小均方误差估计是使 $R_{\text{MSE}}(\hat{\boldsymbol{\theta}}|\mathbf{y})$ 最小化的参数向量, 如图 1.5 所示。下面将证明, 贝叶斯最小均方误差估计是后验概率密度函数的条件均值。假设均方误差风险函数是可微的, 且有确定的最小值, 则令均方误差风险函数的梯度值为零, 可求得最小均方误差估计值, 即

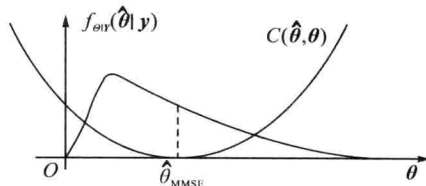


图 1.5 均方误差代价函数和估计