

科学版



研究生教学丛书

数理统计

(第三版)

师义民 徐 伟 编著
秦超英 许 勇



科学出版社
www.sciencep.com

科学版研究生教学丛书

数 理 统 计

(第三版)

师义民 徐 伟 编著
秦超英 许 勇



科 学 出 版 社

北 京

内 容 简 介

本书是根据全国工科院校硕士研究生“数理统计”课程的基本要求,在保留第二版的大部分内容和优点的基础上,适当补充和修改而成。全书共分8章,内容包括:统计量与抽样分布、参数估计、统计决策与贝叶斯估计、假设检验、方差分析与试验设计、回归分析、多元分析初步、统计软件SPSS简介。与第二版相比较,加强了数理统计的经典内容和统计软件及应用的介绍,旨在提高工科研究生的统计理论水平和应用能力。书中各章配有适量的习题,书后附有习题答案。

本书可作为工科各专业研究生,数学与应用数学、信息与计算科学、统计学专业本科生的教材,也可供广大工程技术人员参考使用。

图书在版编目(CIP)数据

数理统计/师义民等编著. —3版. —北京:科学出版社,2009

(科学版研究生教学丛书)

ISBN 978-7-03-023176-5

I. 数… II. 师… III. 数理统计-研究生-教材 IV. O212

中国版本图书馆 CIP 数据核字(2008)第 158105 号

责任编辑:赵 靖 滕亚帆/责任校对:鲁 素

责任印制:张克忠/封面设计:陈 敬

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>

源海印刷有限责任公司印刷

科学出版社发行 各地新华书店经销

*

1999年7月西北工业大学出版社第一版

2002年9月第二版 开本:B5(720×1000)

2009年2月第三版 印张:20 1/2

2009年2月第十次印刷 字数:390 000

印数:22 001—26 000

定价:29.00元

(如有印装质量问题,我社负责调换〈新欣〉)

第三版前言

《数理统计》第二版于 2002 年在科学出版社出版, 先后多次印刷. 近年来, 该书一直作为西北工业大学工科研究生及应用数学系本科生的教学用书. 国内各所院校也使用了本教材, 反映良好. 本教材于 2004 年各所获西北工业大学优秀教材二等奖.

为适应研究生“数理统计”课程教学改革的需要, 这次改版是在保留第二版的大部分内容和优点的基础上, 适当补充和修改而成. 将原书第 1 章中有关统计分布、多元正态分布的内容分别归入统计量与抽样分布、多元分析初步中讲述. 本次修订加强了数理统计的经典内容和统计软件及应用的介绍. 增加的新内容有: 经验贝叶斯估计、似然比检验、几类一元非线性回归、统计软件 SPSS 简介及应用实例等. 旨在提高工科研究生的统计理论水平和应用能力.

本书的第 1 章和第 2 章由秦超英编写、修订, 第 3 章和第 7 章由师义民编写、修订, 第 4 章和第 6 章由徐伟编写、修订, 第 5 章和第 8 章由许勇编写、修订. 全书由师义民统稿与整理.

本书的编写得到了西北工业大学概率统计教研室的许多教师和研究生、西北工业大学研究生院及科学出版社的大力支持和帮助, 在此一并致以衷心感谢.

由于编者水平有限, 书中错误之处, 恳请读者指正.

编 者

2008 年 5 月于西北工业大学

第二版前言

本书自 1999 年第一版出版以来,一直在西北工业大学 1999 级、2000 级和 2001 级研究生“数理统计”课程中使用.现根据使用情况并考虑到新世纪研究生“数理统计”课程的教学改革和西北工业大学“三航”高技术各专业对“数理统计”课程的需要,对第一版内容作了进一步调整与修改.本次修订增加了第 1 章“基础知识”和第 4 章“统计决策与贝叶斯估计”;去掉了第一版中的第 8 章“随机模拟”和第 9 章“统计软件 SPSS 简介”;另外将原书中第 4 章“方差分析”和第 6 章“试验设计”合并为现在的第 6 章“方差分析与试验设计”,并在内容上作了适当补充与修改.本次修订重点加强了数理统计的经典内容,如次序统计量及其分布、充分完备统计量、最小方差无偏估计、假设检验的基本概念等,旨在提高工科研究生的统计理论水平.

本书的第 1 章和第 6 章由赵选民编写、修订,第 2 章和第 3 章由秦超英编写、修订,第 4 章和第 8 章由师义民编写、修订,第 5 章和第 7 章由徐伟编写、修订.全书由赵选民同志统稿与整理.陕西师范大学刘新平教授仔细地审阅了全稿,并提出了许多宝贵意见与建议.西北工业大学应用数学系概率统计教研室、西北工业大学研究生院、科学出版社对本书第二版的出版给予了大力支持与帮助,在此,一并致以衷心的感谢.

由于编者水平有限,书中存在不妥之处,敬请读者指正.

编 者

2001 年 10 月于西北工业大学

第一版前言

本书是根据全国工科院校硕士研究生“数理统计”课程的教学基本要求而编写的。全书共分9章,前5章介绍数理统计的基本理论与基本方法,内容包括:数理统计的基本概念、抽样分布、参数估计、假设检验、方差分析和回归分析。考虑到面向21世纪工科研究生数理统计课程教学改革和实际应用的需要,第6章介绍了正交试验设计、SN比及其试验设计和三次设计等方法,第7章、第8章、第9章分别介绍了多元分析初步、随机模拟方法和常用统计软件。这些方法在工、农业生产,社会、经济、工程技术和自然科学等领域都具有广泛的应用。本书各章均配有适量的习题,书末附有习题答案或提示。

在本书的编写过程中,考虑到工科硕士研究生的数学基础和教学特点,对数理统计学的基础与核心内容,尽量做到循序渐进,由浅入深,叙述严谨,分析透彻。而对应用方法部分,通过对典型实例的分析来介绍方法,培养学生应用所学知识解决工程实际问题的能力。本书的后4章内容基本独立,教师可根据学时和不同的教学要求选讲有关内容,或留给学生自学。

本书可作为工学、经济学硕士研究生“数理统计”课程48~70学时的教材,也可作为数学与应用数学、信息与计算科学、统计学、管理等专业本科生、研究生的教材或教学参考书,亦可供工程技术人员参考。

本书的第1章和第2章由秦超英编写,第3章和第6章由赵选民编写,第4章和第7章由师义民编写,第5章、第8章和第9章由徐伟编写,全书由赵选民统稿整理。本书的初稿曾作为讲义在西北工业大学95级、96级、97级和98级四届研究生教学中使用,并几经修改得以完善。西北工业大学应用数学系概率统计教研室、西北工业大学出版社和研究生院对本书的编写、出版给予了大力支持和帮助,西安交通大学范金城教授仔细地审阅了全稿,并提出了许多宝贵意见与建议,在此,一并致以衷心的感谢。

由于编者水平有限,书中存在的不妥之处,敬请读者指正。

编 者

1998年12月于西北工业大学

目 录

第 1 章 统计量与抽样分布	1
1.1 基本概念	1
1.2 充分统计量与完备统计量	5
1.3 抽样分布	13
1.4 次序统计量及其分布	28
习题 1	33
第 2 章 参数估计	36
2.1 点估计与优良性	36
2.2 点估计量的求法	41
2.3 最小方差无偏估计和有效估计	51
2.4 区间估计	59
习题 2	75
第 3 章 统计决策与贝叶斯估计	79
3.1 统计决策的基本概念	79
3.2 贝叶斯估计	84
3.3 minimax 估计	101
3.4 经验贝叶斯估计	107
习题 3	111
第 4 章 假设检验	113
4.1 假设检验的基本概念	113
4.2 正态总体均值与方差的假设检验	120
4.3 非参数假设检验方法	128
4.4 似然比检验	143
习题 4	146
第 5 章 方差分析与试验设计	151
5.1 单因素方差分析	151
5.2 两因素方差分析	159
5.3 正交试验设计	172
习题 5	186

第 6 章 回归分析	190
6.1 一元线性回归分析	190
6.2 多元线性回归分析	199
6.3 几类一元非线性回归	212
习题 6	214
第 7 章 多元分析初步	217
7.1 多元正态分布的定义及性质	217
7.2 多元正态分布参数的估计与假设检验	221
7.3 判别分析	227
7.4 主成分分析	244
习题 7	250
第 8 章 统计软件 SPSS 简介	253
8.1 引言	253
8.2 SPSS 基本操作	253
8.3 统计图形	257
8.4 统计分析	262
习题 8	281
习题答案	282
参考文献	289
附表	290

第 1 章 统计量与抽样分布

数理统计学是研究随机现象规律性的一门学科,它以概率论为理论基础,研究如何以有效的方式收集、整理和分析受到随机因素影响的数据,并对所考察的问题作出推理和预测,直至为采取某种决策提供依据和建议.数理统计研究的内容非常广泛,概括起来可分为两大类:一是试验设计,即研究如何对随机现象进行观察和试验,以便更合理更有效地获得试验数据;二是统计推断,即研究如何对所获得的有限数据进行整理和加工,并对所考察的对象的某些性质作出尽可能精确可靠的判断.

数理统计是一门应用性很强的数学学科,已被广泛地应用到自然科学和工程技术的各个领域.数理统计方法已成为各学科从事科学研究以及在生产、管理、经济等部门进行有效工作的必不可少的数学工具.本章在回顾数理统计中的一些基本概念,如总体、样本、统计量和经验分布函数的基础上,介绍充分统计量、完备统计量以及一些重要统计量的分布等.

1.1 基本概念

一、总体和样本

1. 总体

在数理统计学中,把所研究对象的全体元素组成的集合称为总体(或称母体),而把组成总体的每个元素称为个体.例如,在考察某批灯泡的质量时,该批灯泡的全体就组成一个总体,而其中每个灯泡就是个体.

但是,在实际应用中,人们所关心的并不是总体中个体的一切方面,而所研究的往往是总体中个体的某一项或某几项数量指标.例如,考察灯泡质量时,我们并不关心灯泡的形状、式样等特征,而只研究灯泡的寿命、亮度等数量指标特征.如果只考察灯泡寿命这一项指标时,由于一批灯泡中每个灯泡都有一个确定的寿命值,因此,自然地把这批灯泡寿命值的全体视为总体,而其中每个灯泡的寿命值就是个体.由于具有不同寿命值的灯泡的比例是按一定规律分布的,即任取一个灯泡其寿命为某一值具有一定概率,因而,这批灯泡的寿命是一个随机变量,也就是说,可以用一个随机变量 X 来表示这批灯泡的寿命这个总体.因此,在数理统计中,任何一个总体都可用一个随机变量来描述.总体的分布及数字特征,即指表示总体的随机

变量的分布及数字特征. 对总体的研究也就归结为对表示总体的随机变量的研究.

2. 样本

为了了解总体 X 的分布规律或某些特征, 必须对总体进行抽样观察, 即从总体 X 中, 随机抽取 n 个个体 $X_1, X_2, \dots, X_n$, 记为 $(X_1, X_2, \dots, X_n)^T$, 并称此为来自总体 X 的容量为 n 的样本. 由于每个 X_i 都是从总体 X 中随机抽取的, 它的取值就在总体 X 的可能取值范围内随机取得, 自然每个 X_i 也是随机变量, 从而样本 $(X_1, X_2, \dots, X_n)^T$ 是一个 n 维随机向量. 在抽样观测后, 它们是 n 个数据 $(x_1, x_2, \dots, x_n)^T$, 称之为样本 $(X_1, X_2, \dots, X_n)^T$ 的一个观测值, 简称样本值. 样本 $(X_1, X_2, \dots, X_n)^T$ 可能取值的全体称为样本空间, 记为 Ω .

我们的目的是依据从总体 X 中抽取的一个样本值 $(x_1, x_2, \dots, x_n)^T$, 对总体 X 的分布或某些特征进行分析推断, 因而要求抽取的样本能很好地反映总体的特征且便于处理, 于是, 提出下面两点要求:

(1) 代表性——要求样本 $X_1, X_2, \dots, X_n$ 同分布且每个 X_i 与总体 X 具有相同的分布.

(2) 独立性——要求样本 $X_1, X_2, \dots, X_n$ 是相互独立的随机变量.

满足上述两条性质的样本称为简单随机样本. 若无特别说明, 今后提到的样本均指简单随机样本. 关于样本 $(X_1, X_2, \dots, X_n)^T$ 的分布有如下定理.

定理 1.1 设总体 X 的分布函数为 $F(x)$ (或分布密度为 $f(x)$ 或分布律为 $P\{X=x^{(i)}\}=p(x^{(i)}), i=1, 2, \dots\}$), 则来自总体 X 的样本 $(X_1, X_2, \dots, X_n)^T$ 的联合分布函数为 $\prod_{i=1}^n F(x_i)$ (或联合分布密度为 $\prod_{i=1}^n f(x_i)$ 或联合分布律为 $\prod_{i=1}^n p(x_i)$).

例 1.1 设总体 X 服从参数为 p 的两点分布, 即

$$P\{X=1\}=p, P\{X=0\}=1-p, \quad 0 < p < 1,$$

试求样本 $(X_1, X_2, \dots, X_n)^T$ 的联合分布律.

解 由于总体 X 的分布律可以写成

$$p(x) = P\{X=x\} = p^x(1-p)^{1-x}, \quad x=0, 1,$$

故由定理 1.1, 样本 $(X_1, X_2, \dots, X_n)^T$ 的联合分布律为

$$\prod_{i=1}^n p(x_i) = \prod_{i=1}^n p^{x_i}(1-p)^{1-x_i} = p^{\sum_{i=1}^n x_i} (1-p)^{n-\sum_{i=1}^n x_i}.$$

例 1.2 设总体 X 服从正态分布 $N(\mu, \sigma^2)$, 试求样本 $(X_1, X_2, \dots, X_n)^T$ 的联合分布密度.

解 总体 X 的分布密度为

$$f(x) = \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{(x-\mu)^2}{2\sigma^2}}, \quad -\infty < x < \infty,$$

故样本 $(X_1, X_2, \dots, X_n)^T$ 的联合分布密度为

$$\prod_{i=1}^n f(x_i) = \prod_{i=1}^n \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{(x_i-\mu)^2}{2\sigma^2}} = \frac{1}{(\sqrt{2\pi}\sigma)^n} e^{-\frac{1}{2\sigma^2} \sum_{i=1}^n (x_i-\mu)^2}.$$

二、统计量和样本矩

1. 统计量

样本是总体的代表和反映,但在抽取样本之后,并不能直接利用样本进行推断,而需要对样本进行“加工”和“提炼”,把样本中关于总体的信息集中起来,这便是针对不同的问题构造出样本的某种函数.为此,引进统计量的概念.

定义 1.1 设 $(X_1, X_2, \dots, X_n)^T$ 为总体 X 的一个样本,若 $f(X_1, X_2, \dots, X_n)$ 为一个函数,且 f 中不含任何未知参数,则称 $f(X_1, X_2, \dots, X_n)$ 为一个统计量.

由于样本 $X_1, X_2, \dots, X_n$ 是随机变量,统计量 $f(X_1, X_2, \dots, X_n)$ 也是随机变量,它们应有确定的分布,其分布称之为抽样分布.

2. 常用统计量——样本矩

定义 1.2 设 $(X_1, X_2, \dots, X_n)^T$ 是从总体 X 中抽取的样本,称统计量:

$$\bar{X} = \frac{1}{n} \sum_{i=1}^n X_i \text{ 为样本均值;}$$

$$S_n^2 = \frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^2 = \frac{1}{n} \sum_{i=1}^n X_i^2 - \bar{X}^2 \text{ 为样本方差;}$$

$$S_n^{*2} = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2 \text{ 为修正样本方差(简称样本方差);}$$

$$S_n = \sqrt{\frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^2} \text{ 为样本标准差;}$$

$$A_k = \frac{1}{n} \sum_{i=1}^n X_i^k \quad (k=1,2,\dots) \text{ 为样本 } k \text{ 阶原点矩;}$$

$$B_k = \frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^k \quad (k=1,2,\dots) \text{ 为样本 } k \text{ 阶中心矩.}$$

由定义 1.2 可见, $A_1 = \bar{X}$, $B_2 = S_n^2$, $S_n^{*2} = \frac{n}{n-1} S_n^2$.

用 $\bar{x}, s_n^2, a_k, b_k$ 分别表示 $\bar{X}, S_n^2, A_k, B_k$ 的观测值,此时只要把定义 1.2 中 X_i 改为 x_i 即可.

由大数定律可以证明,只要总体 X 的 k 阶矩存在,则样本的 k 阶矩依概率收

敛于总体 X 的 k 阶矩. 即对任意 $\epsilon > 0$, 有

$$\lim_{n \rightarrow \infty} P\{|\bar{X} - \mu| < \epsilon\} = 1,$$

$$\lim_{n \rightarrow \infty} P\{|S_n^2 - \sigma^2| < \epsilon\} = 1,$$

式中 $\mu = EX, \sigma^2 = DX$. 此结论表明, n 很大时可用一次抽样后所得的样本均值 $\bar{X}$ 和样本方差 S_n^2 分别作为总体 X 的均值 μ 和方差 σ^2 的近似值.

定理 1.2 设总体 X 具有 $2k$ 阶矩, 则来自总体 X 的样本 k 阶原点矩 A_k 的数学期望和方差分别为

$$EA_k = \alpha_k,$$

$$DA_k = \frac{\alpha_{2k} - \alpha_k^2}{n},$$

其中 $\alpha_k = EX^k (k=1, 2, \dots)$ 表示总体 X 的 k 阶原点矩.

$$\text{证明} \quad EA_k = E\left(\frac{1}{n} \sum_{i=1}^n X_i^k\right) = \frac{1}{n} \sum_{i=1}^n EX_i^k = \frac{1}{n} \sum_{i=1}^n EX^k = \alpha_k,$$

$$\begin{aligned} DA_k &= D\left(\frac{1}{n} \sum_{i=1}^n X_i^k\right) = \frac{1}{n^2} \sum_{i=1}^n DX_i^k = \frac{1}{n^2} \sum_{i=1}^n DX^k \\ &= \frac{1}{n} (EX^{2k} - (EX^k)^2) = \frac{\alpha_{2k} - \alpha_k^2}{n}. \end{aligned}$$

$$\text{推论} \quad E\bar{X} = EX, D\bar{X} = \frac{1}{n} DX, ES_n^2 = \frac{n-1}{n} DX, ES_n^{*2} = DX.$$

三、经验分布函数

根据样本来估计和推断总体 X 的分布函数 $F(x)$, 是数理统计要解决的一个重要问题. 为此, 引进经验分布函数的概念, 并介绍它的性质.

设总体 X 的分布函数为 $F(x)$, 现对 X 进行 n 次重复独立观测, 即对总体作 n 次简单随机抽样, 以 $v_n(x)$ 表示随机事件 $\{X \leq x\}$ 在这 n 次重复独立观测中出现的次数, 即 n 个观测值 $x_1, x_2, \dots, x_n$ 中小于等于 x 的个数.

对 X 每进行了 n 次重复独立观测, 便得到总体 X 的样本 $(X_1, X_2, \dots, X_n)^T$ 的一个观测值 $(x_1, x_2, \dots, x_n)^T$, 从而对固定的 $x \in (-\infty, +\infty)$ 可以确定 $v_n(x)$ 所取的数值, 这个数值就是 $x_1, x_2, \dots, x_n$ 的 n 个数中小于等于 x 的个数. 若重复进行 n 次抽样时, 对于同一个 x , $v_n(x)$ 可能将取不同数值, 即 $v_n(x)$ 随样本取不同样本值而取不同值. 因此, $v_n(x)$ 实际上是一个统计量, 从而也是随机变量. $v_n(x)$ 通常称为经验频数. 由于在 n 重独立试验中, 某事件出现的次数服从二项分布, 故有 $v_n(x)$ 服从二项分布

$$\begin{aligned} P\{v_n(x) = k\} &= C_n^k (P\{X \leq x\})^k (1 - P\{X \leq x\})^{n-k} \\ &= C_n^k [F(x)]^k [1 - F(x)]^{n-k}, \end{aligned}$$

其中 $k=0, 1, 2, \dots, n$, 即

$$v_n(x) \sim B(n, F(x)).$$

定义 1.3 称函数

$$F_n(x) = \frac{v_n(x)}{n}, \quad -\infty < x < +\infty \quad (1.1)$$

为总体 X 的经验分布函数.

设 $(X_1, X_2, \dots, X_n)^T$ 是来自总体 X 的样本, 其样本值为 $(x_1, x_2, \dots, x_n)^T$. 将 $x_1, x_2, \dots, x_n$ 按由小到大的顺序排列并重新编号为

$$x_{(1)} \leq x_{(2)} \leq \dots \leq x_{(n)},$$

则总体 X 的经验分布函数可表示为

$$F_n(x) = \frac{v_n(x)}{n} = \begin{cases} 0, & x < x_{(1)}, \\ \frac{k}{n}, & x_{(k)} \leq x < x_{(k+1)}, \quad k = 1, 2, \dots, n-1, \\ 1, & x \geq x_{(n)}. \end{cases} \quad (1.2)$$

经验分布函数 $F_n(x)$ 的性质:

(1) 当给定样本值 $(x_1, x_2, \dots, x_n)^T$ 时, $F_n(x)$ 是一个分布函数, 即具有以下性质:

- (i) $0 \leq F_n(x) \leq 1$;
- (ii) $F_n(-\infty) = 0, F_n(+\infty) = 1$;
- (iii) $F_n(x)$ 非减且右连续.

(2) $F_n(x)$ 是随机变量, 且 $nF_n(x) = v_n(x) \sim B(n, F(x))$, 进而

$$E[F_n(x)] = F(x), \quad D[F_n(x)] = \frac{1}{n}F(x)[1 - F(x)].$$

(3) 当 $n \rightarrow \infty$ 时, $F_n(x)$ 依概率收敛于总体 X 的分布函数 $F(x)$, 即对任意 $\epsilon > 0$ 及 $x \in (-\infty, +\infty)$, 有

$$\lim_{n \rightarrow \infty} P\{|F_n(x) - F(x)| < \epsilon\} = 1.$$

实际上, $F_n(x)$ 还依概率 1 一致地收敛于 $F(x)$, 所谓格里文科定理指出了这一更深刻的结论, 即

$$P\left\{\lim_{n \rightarrow \infty} \sup_{-\infty < x < +\infty} |F_n(x) - F(x)| = 0\right\} = 1.$$

这一性质表明, 当 n 很大时, 由一个样本值得到的经验分布函数 $F_n(x)$ 是总体分布函数 $F(x)$ 的一个优良的估计.

1.2 充分统计量与完备统计量

一、充分统计量

在数理统计中, 由样本来推断总体的前提是: 样本中包含了总体分布的信息.

简单随机样本满足了这一前提条件. 样本中所包含的关于总体分布的信息可分为两部分, 其一是关于总体结构的信息, 即反映总体分布的结构(类型). 例如, 假定总体分布是正态分布, 则来自该总体的样本也是相互独立、相同的正态分布. 因此, 在样本中包含了总体分布是正态分布的信息. 其二是关于总体分布中未知参数的信息, 这是由于样本的分布中包含了总体分布中的未知参数. 现在, 我们把目标集中在后一部分的信息, 即要推断总体分布的未知参数, 为此构造一个合适的统计量, 对样本进行加工, 以便把样本中关于未知参数的信息“提炼”出来. 譬如, 为了估计总体的均值 μ , 人们把样本 $(X_1, X_2, \dots, X_n)^T$ 加工成样本均值 $\bar{X}$, 为了估计总体方差 σ^2 , 把样本加工成样本方差 S_n^2 , 然后用 $\bar{X}$ 和 S_n^2 分别去估计总体均值 μ 和方差 σ^2 . 那么试问: 统计量 $\bar{X}$ 或 S_n^2 与样本 $(X_1, X_2, \dots, X_n)^T$ 中所含 μ 或 σ^2 的信息是否一样多? 换言之, 统计量 $\bar{X}$ 和 S_n^2 是否把样本 $(X_1, X_2, \dots, X_n)^T$ 中关于 μ 和 σ^2 的信息全部提炼出来, 而没有任何信息损失. 显然, 一个“好”的统计量, 应该能够将样本中所包含的关于未知参数的信息全部提取出来. 如何将这样一个直观想法用数学形式来表示呢? 英国著名统计学家费希尔(R. A. Fisher)在 1922 年提出了一个重要的概念——充分统计量. 粗略地说, 充分统计量就是“不损失信息”的统计量, 它的精确表述如下:

定义 1.4 设 $(X_1, X_2, \dots, X_n)^T$ 是来自总体 X 具有分布函数 $F(x; \theta)$ 的一个样本, $T = T(X_1, X_2, \dots, X_n)$ 为一个(一维或多维的)统计量, 当给定 $T = t$ 时, 若样本 $(X_1, X_2, \dots, X_n)^T$ 的条件分布(离散总体为条件概率, 连续总体为条件密度)与参数 θ 无关, 则称 T 是 θ 的充分统计量.

充分统计量的含义可以这样来解释: 样本中包含关于总体分布中未知参数 θ 的信息, 是因为样本的联合分布与 θ 有关. 对统计量 T , 如果已经知道它的值以后, 样本的条件分布与 θ 无关, 就意味着样本的剩余部分中不再包含关于 θ 的信息. 换言之, 在 T 中包含了关于 θ 的全部信息, 因此, 要做关于 θ 的统计推断只需从 T 出发即可. 这就是“充分统计量”这个概念中“充分”这个词的含义.

例 1.3 设总体 X 服从两点分布 $B(1, p)$, 即

$$P\{X = x\} = p^x(1-p)^{1-x}, \quad x = 0, 1,$$

其中 $0 < p < 1$, $(X_1, X_2, \dots, X_n)^T$ 是来自总体 X 的一个样本, 试证 $\bar{X} = \frac{1}{n} \sum_{i=1}^n X_i$ 是参数 p 的充分统计量.

证明 由于 $X_i \sim B(1, p)$, 易知 $n\bar{X} = \sum_{i=1}^n X_i \sim B(n, p)$, 即有

$$P\{n\bar{X} = k\} = C_n^k p^k (1-p)^{n-k}, \quad k = 0, 1, \dots, n.$$

设 $(x_1, x_2, \dots, x_n)^T$ 为样本值, 其中 $x_i = 0$ 或 1 . 当已知 $\sum_{i=1}^n x_i = k$, 即 $\bar{X} = \frac{k}{n}$

时,样本 $(X_1, X_2, \dots, X_n)^T$ 的条件概率

$$\begin{aligned}
 & P\left\{X_1 = x_1, X_2 = x_2, \dots, X_n = x_n \mid \bar{X} = \frac{k}{n}\right\} \\
 &= \frac{P\left\{X_1 = x_1, X_2 = x_2, \dots, X_n = x_n, \bar{X} = \frac{k}{n}\right\}}{P\left\{\bar{X} = \frac{k}{n}\right\}} \\
 &= \begin{cases} \frac{P\{X_1 = x_1, X_2 = x_2, \dots, X_n = x_n\}}{P\{n\bar{X} = k\}}, & \sum_{i=1}^n x_i = k, \\ 0, & \sum_{i=1}^n x_i \neq k \end{cases} \\
 &= \begin{cases} \frac{p^{\sum_{i=1}^n x_i} (1-p)^{n-\sum_{i=1}^n x_i}}{C_n^k p^k (1-p)^{n-k}}, & \sum_{i=1}^n x_i = k, \\ 0, & \sum_{i=1}^n x_i \neq k \end{cases} \\
 &= \begin{cases} \frac{1}{C_n^k}, & \sum_{i=1}^n x_i = k, \\ 0, & \sum_{i=1}^n x_i \neq k \end{cases}
 \end{aligned}$$

与 p 无关,所以 $\bar{X}$ 是 p 的充分统计量.

二、因子分解定理

根据充分统计量的含义,在对总体未知参数进行推断时,应在可能的情况下尽量找出关于未知参数的充分统计量.但从定义出发来判别一个统计量是否是充分统计量是很麻烦的.为此,需要一个简单的判别准则.下面给出一个定理——因子分解定理,运用这个定理,判别甚至寻找一个充分统计量有时会很方便.

定理 1.3(因子分解定理) (1) 连续型情况:设总体 X 具有分布密度 $f(x; \theta)$, $(X_1, X_2, \dots, X_n)^T$ 是一样本, $T(X_1, X_2, \dots, X_n)$ 是一个统计量,则 T 为 θ 的充分统计量的充要条件是:样本的联合分布密度函数可以分解为

$$L(\theta) \triangleq \prod_{i=1}^n f(x_i; \theta) = h(x_1, x_2, \dots, x_n) g(T(x_1, x_2, \dots, x_n); \theta), \quad (1.3)$$

其中 h 是 $x_1, x_2, \dots, x_n$ 的非负函数且与 θ 无关, g 仅通过 T 依赖于 $x_1, x_2, \dots, x_n$.

(2) 离散型情况:设总体 X 的分布律为 $P\{X = x^{(i)}\} = p(x^{(i)}; \theta)$ ($i = 1, 2, \dots$), $(X_1, X_2, \dots, X_n)^T$ 是一样本, $T(X_1, X_2, \dots, X_n)$ 是一个统计量,则 T 是 θ 的充分统

计量的充要条件是:样本的联合分布律可表示为

$$\begin{aligned} P\{X_1 = x_1, X_2 = x_2, \dots, X_n = x_n\} &\triangleq \prod_{i=1}^n P\{X = x_i\} \\ &= h(x_1, x_2, \dots, x_n)g(T(x_1, x_2, \dots, x_n); \theta), \end{aligned} \quad (1.4)$$

其中 h 是 $x_1, x_2, \dots, x_n$ 的非负函数且与 θ 无关, g 仅通过 T 依赖于 $x_1, x_2, \dots, x_n$.

定理 1.3 又称为费希尔-奈曼准则, 由于证明涉及较多的测度论知识, 故从略.

如果 θ 是参数向量, 如正态总体 $N(\mu, \sigma^2)$ 中, μ, σ^2 都未知, 则记 $\theta = (\mu, \sigma^2)^T$; 统计量 T 是随机向量, 且式 (1.3) 或式 (1.4) 成立, 则称 T 关于 θ 是联合充分的.

必须指出, 如果 θ 和 T 的维数相等, 我们不能由 T 关于 θ 的充分性而推出 T 的第 i 个分量关于 θ 的第 i 个分量是充分的.

例 1.4 根据因子分解定理证明例 1.3.

证明 样本的联合分布律为

$$\begin{aligned} P\{X_1 = x_1, X_2 = x_2, \dots, X_n = x_n\} \\ = p^{\sum_{i=1}^n x_i} (1-p)^{n-\sum_{i=1}^n x_i} = (1-p)^n \left(\frac{p}{1-p}\right)^{\sum_{i=1}^n x_i}. \end{aligned}$$

若取

$$\begin{aligned} T(x_1, x_2, \dots, x_n) &= \frac{1}{n} \sum_{i=1}^n x_i, \\ h(x_1, x_2, \dots, x_n) &= 1, \\ g(T(x_1, x_2, \dots, x_n); p) &= (1-p)^n \left(\frac{p}{1-p}\right)^{nT}, \end{aligned}$$

则有

$$P\{X_1 = x_1, X_2 = x_2, \dots, X_n = x_n\} = h(x_1, x_2, \dots, x_n)g(T(x_1, x_2, \dots, x_n); p),$$

由因子分解定理知, $T(X_1, X_2, \dots, X_n) = \frac{1}{n} \sum_{i=1}^n X_i = \bar{X}$ 是 p 的充分统计量.

例 1.5 设 $(X_1, X_2, \dots, X_n)^T$ 是来自泊松分布 $P(\lambda)$ 的一个样本, 试证明样本均值 $\bar{X}$ 是 λ 的充分统计量.

证明 样本 $(X_1, X_2, \dots, X_n)^T$ 的联合分布律为

$$P\{X_1 = x_1, X_2 = x_2, \dots, X_n = x_n\} = \frac{\lambda^{\sum_{i=1}^n x_i}}{\prod_{i=1}^n x_i!} e^{-n\lambda},$$

若取

$$T(x_1, x_2, \dots, x_n) = \frac{1}{n} \sum_{i=1}^n x_i,$$

$$h(x_1, x_2, \dots, x_n) = \frac{1}{\prod_{i=1}^n x_i!},$$

$$g(T(x_1, x_2, \dots, x_n); \lambda) = \lambda^{nT} e^{-n\lambda},$$

则

$$P\{X_1 = x_1, X_2 = x_2, \dots, X_n = x_n\} = h(x_1, x_2, \dots, x_n)g(T(x_1, x_2, \dots, x_n); \lambda).$$

由因子分解定理知, $T(X_1, X_2, \dots, X_n) = \bar{X}$ 是 λ 的充分统计量.

例 1.6 设 $(X_1, X_2, \dots, X_n)^T$ 是来自正态总体 $N(\mu, 1)$ 的一个样本, 试证样本均值 $\bar{X}$ 是 μ 的充分统计量.

证明 样本 $(X_1, X_2, \dots, X_n)^T$ 的联合分布密度为

$$\begin{aligned} L(\mu) &= \frac{1}{(\sqrt{2\pi})^n} \exp\left\{-\frac{1}{2} \sum_{i=1}^n (x_i - \mu)^2\right\} \\ &= \frac{1}{(\sqrt{2\pi})^n} \exp\left\{-\frac{1}{2} \sum_{i=1}^n (x_i - \bar{x})^2 - \frac{n}{2} (\mu - \bar{x})^2\right\} \\ &= \frac{1}{(\sqrt{2\pi})^n} \exp\left\{-\frac{1}{2} \sum_{i=1}^n (x_i - \bar{x})^2\right\} \exp\left\{-\frac{n}{2} (\mu - \bar{x})^2\right\}. \end{aligned}$$

若取

$$T(x_1, x_2, \dots, x_n) = \frac{1}{n} \sum_{i=1}^n x_i = \bar{x},$$

$$h(x_1, x_2, \dots, x_n) = \exp\left\{-\frac{1}{2} \sum_{i=1}^n (x_i - \bar{x})^2\right\},$$

$$g(T(x_1, x_2, \dots, x_n); \mu) = \frac{1}{(\sqrt{2\pi})^n} \exp\left\{-\frac{n}{2} (\mu - T)^2\right\},$$

则

$$L(x_1, x_2, \dots, x_n; \mu) = h(x_1, x_2, \dots, x_n)g(T(x_1, x_2, \dots, x_n); \mu).$$

由因子分解定理知, $T(X_1, X_2, \dots, X_n) = \bar{X}$ 是 μ 的充分统计量.

例 1.7 设 $(X_1, X_2, \dots, X_n)^T$ 是来自正态总体 $N(\mu, \sigma^2)$ 的一个样本, 试证 $T(X_1, X_2, \dots, X_n) = \left(\bar{X}, \sum_{i=1}^n X_i^2\right)^T$ 是关于 $\theta = (\mu, \sigma^2)^T$ 的联合充分统计量.

证明 样本的联合分布密度为

$$L(\theta) = \frac{1}{(\sqrt{2\pi}\sigma)^n} \exp\left\{-\frac{1}{2\sigma^2} \sum_{i=1}^n (x_i - \mu)^2\right\}$$