



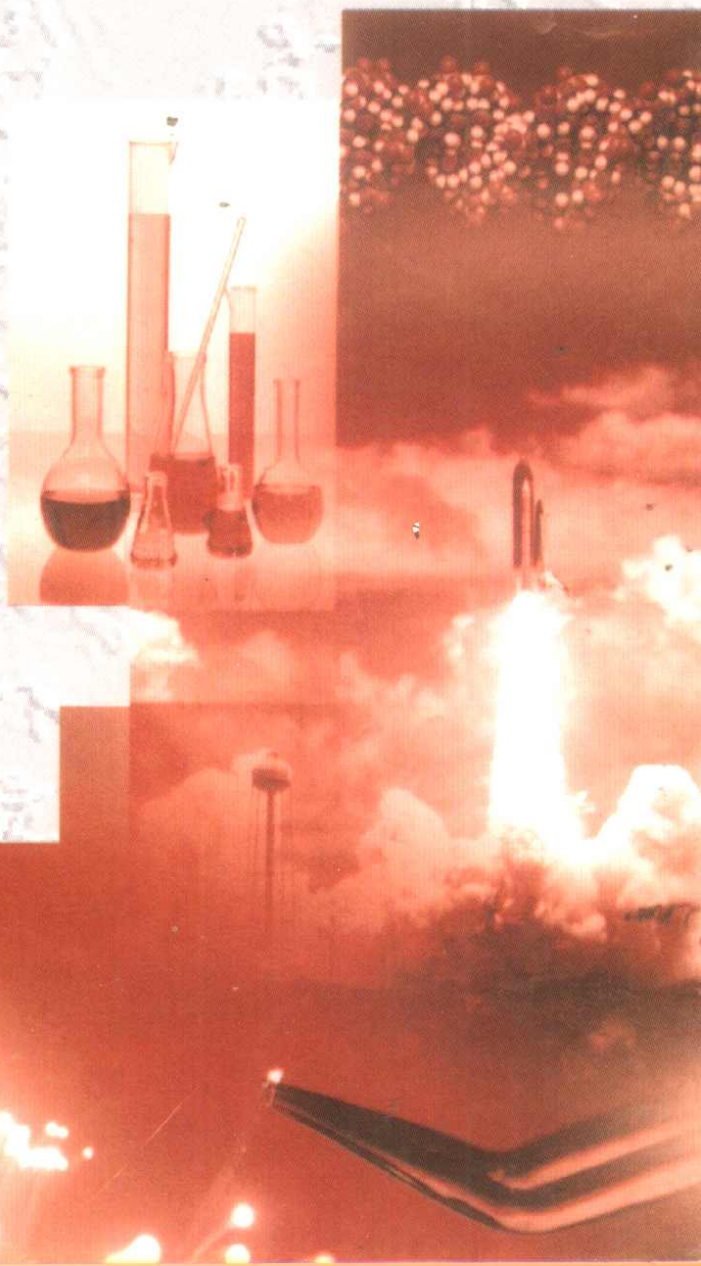
高等学校教材

基础课程系列

数理统计

王新成 编著

*Fundamental
Courses*



-43

工业出版社

661

0212-43
u37

数理统计

王新成 编著

西北工业大学出版社

【内容简介】 本书是依据高等学校工科数学课程教学指导委员会于1995年修订的“概率论与数理统计课程教学的基本要求”而编写的。全书共分为4章,其内容依次为:样本和抽样分布;参数估计;假设检验;回归分析。各章均配有难度适中的习题,且书后附有习题答案,以供读者参考。

本书结构合理,条理清晰,叙述详细,通俗易懂。例题较多,应用实例涉及面较广,适应于高等工科院校不同专业的学生使用,也可作为工程技术人员的参考书。

图书在版编目(CIP)数据

数理统计/王新成编著. —西安:西北工业大学出版社,2002.8
ISBN 7-5612-1518-5

I. 数… II. 王… III. 数理统计-高等学校-教材 IV. 0212

中国版本图书馆CIP数据核字(2002)第052957号

出版发行:西北工业大学出版社

通信地址:西安市友谊西路127号 邮编:710072 电话:029-8493844

网 址: <http://www.nwpup.com>

印刷者:陕西向阳印务有限公司印装

开 本:850 mm×1 168 mm 1/32

印 张:7.125

字 数:179千字

版 次:2002年8月第1版 2002年8月第1次印刷

印 数:1~7 000册

定 价:9.00元

前 言

本书是依据高等学校工科数学课程教学指导委员会于 1995 年修订的“概率论与数理统计课程教学基本要求”，为适应工科数学教学改革的需要，在作者多年实践以及多年讲授数理统计课程的基础上，吸取众多《数理统计》教材的优点编写而成的。

全书共分为 4 章，其内容依次为：样本和抽样分布，参数估计，假设检验，回归分析。各章均配有难度适中的习题，且书后附有习题答案，以供读者参考。

数理统计主要是以概率论为基础，处理带有随机性影响数据的一门数学课程。该门课程中的概念、定义不易理解，方法也不易掌握和应用，学习中可能有一定的困难。但在科学技术研究与生产实践中以及社会的各个方面，却有着极为广泛的应用。为此，本书在有限的篇幅中，尽可能多讲一点统计思想、基本方法、如何应用以及应用中应注意的事项，特别是一些近似方法，实际应用中是非常实用的、有效的。

本书的目的致力于读者确切理解数理统计思想，掌握其常用的处理问题的基本方法，且能用所学的方法解决实际中碰到的问题。

崔荣泉教授在百忙中详细审阅了本书初稿与修改稿，并提出了许多宝贵的意见和建议，在此深表谢意。

鉴于时间仓促，工作繁忙，加之水平有限，错误和疏漏之处在所难免，恳请广大读者批评指正。

作 者

2001 年 6 月于西安建筑科技大学

目 录

绪论	1
第 1 章 样本和抽样分布	3
§ 1.1 基本概念	3
1.1.1 总体及其分布	3
1.1.2 样本(简单随机样本)	4
1.1.3 样本分布	6
1.1.4 统计量与样本矩	7
§ 1.2 抽样分布	12
1.2.1 χ^2 分布	13
1.2.2 t (Student) 分布	18
1.2.3 F 分布	22
1.2.4 正态总体的抽样分布	25
1.2.5 次序统计量的分布与经验分布函数(样本分布函数)	28
习题	34
第 2 章 参数估计	38
§ 2.1 参数的点估计	38
2.1.1 矩估计法	40
2.1.2 极大似然估计法	44
§ 2.2 估计量的评价标准	52
2.2.1 无偏性	53

2.2.2	有效性	55
2.2.3	一致性	58
§ 2.3	区间估计	60
2.3.1	双侧区间估计	60
2.3.2	单侧区间估计	64
§ 2.4	正态总体均值与方差的区间估计	67
2.4.1	单个正态总体均值 μ , 方差 σ^2 的区间估计	67
2.4.2	两个正态总体均值差和方差比的置信区间	73
习题		81

第3章 假设检验 85

§ 3.1	假设检验的基本概念与基本原理	85
3.1.1	假设检验的基本概念	85
3.1.2	假设检验的基本原理	90
3.1.3	两类错误	90
§ 3.2	单个正态总体参数的假设检验	92
3.2.1	单个总体 $N(\mu, \sigma^2)$ 均值 μ 的检验	92
3.2.2	单个总体 $N(\mu, \sigma^2)$ 方差 σ^2 的检验 (χ^2 检验法)	96
§ 3.3	两个正态总体参数的假设检验	99
3.3.1	两个正态总体均值的检验	99
3.3.2	两个正态总体方差相等的检验 (F 检验)	103
§ 3.4	单侧假设检验	105
§ 3.5	分布假设检验	115
3.5.1	分布拟合优度检验	115
3.5.2	χ^2 检验法	116
习题		123

第 4 章 回归分析	127
§ 4.1 回归分析的基本概念	127
§ 4.2 一元线性回归	129
4.2.1 一元线性回归模型	129
4.2.2 参数 a, b 的最小二乘估计	132
4.2.3 回归方程的显著性检验	137
4.2.4 预测与控制	148
§ 4.3 一元非线性回归的线性化处理	153
§ 4.4 多元线性回归	161
4.4.1 多元线性回归模型	162
4.4.2 参数 $b_0, b_1, \dots, b_p, \sigma^2$ 的最小二乘估计	163
4.4.3 多元回归方程的显著性检验	175
4.4.4 回归系数的显著性检验	180
4.4.5 预测	186
习题	188
附录 I	192
附录 II	215

绪 论

数理统计学是数学的一个分支，它的任务是研究怎样有效地收集和使用带有随机性影响的数据，以对所考察的问题作出推断和预测，直至为采取一定的决策和行动提供依据和建议。

所谓“用有效的方式收集数据”一语中，“有效”一词如何解释，归纳起来有两个方面：一是建立一个在数学上可以处理并尽可能简单方便的模型来描述数据；二是数据中要包含尽可能多的，与所研究的问题有关的信息。用有效的方式收集数据问题的研究，构成了数理统计学中的两个分支，其一称为**抽样理论**，其二称为**试验设计**。

获取数据的目的，是提供与所研究的问题有关的信息，但这种信息并非是一目了然地表现出来，而需要用“有效的方式去集中、提取”，进而用于对所研究的问题作出一定的结论。这种“结论”，在数理统计中称之为“**推断**”（或**统计推断**）。显然推断就是对所提出统计问题的一个回答，也常称其为**统计问题的解**。

所谓“**有效地使用数据**”，就是要使用有效的方法，去集中和提取试验数据中的有关信息，以对所研究的问题作出尽可能精确和可靠的推断。其所以只能做到“尽可能”而非绝对地精确和可靠，是因为数据受到随机性因素的影响，这种影响可以通过统计方法去估计或缩小其干扰作用，但不可能完全消除。为有效地使用数据以进行推断，涉及很多数学知识。需要建立一定的数学模型，并给定某些准则，才有可能去评价和比较种种统计推断方法的优劣。

数据的随机性来源有二：一是问题中涉及的研究对象为数很

大，不可能对其全部逐一研究，而只能用“一定的方式”挑选其一部分（例如抽取简单随机样本之方法）去考察，这一部分被挑出就具有偶然性，从而决定了数据的随机性；另一种来源是试验的随机误差，这是指那种在试验中未加控制、无法控制，甚至不了解的因素所引起的误差。数据必须带有随机性的影响，才能成为数理统计学的研究对象，这也是区别数理统计方法与其他数据处理方法之根本点。

在近代，尤其是出现电子计算机之后，数理统计学发展速度加快，已成为一个成熟的数学分支。数理统计方法众多，例如常用的参数估计、非参数估计、稳健估计、假设检验，拟合优度检验、方差分析、回归分析、非线性回归、相关分析、试验设计、正交设计、抽样理论及多元统计的许多方法等等。

数理统计学应用非常广泛，几乎在人类活动的一切领域中都能程度不同地找到它的应用，也常常用于种种专门领域（物理、化学、工程、生物、经济、社会学等），但只涉及其中带有随机性的数据分析问题，而不是以任何一种专门的知识领域为研究对象。但是，在用数理统计方法分析带随机性的数据时，从统计模型的选择、实验方案的制定、统计方法的正确使用以至所得结论的恰当解释，都离不开所论问题的专门知识。

科学界有人说过这样一句话：不懂得数理统计的工程师只能算半个工程师。由此可见，掌握数理统计的知识与方法对科技人员是何等重要。

第 1 章 样本和抽样分布

数理统计学的内容和方法十分丰富,本章首先介绍总体、随机样本及统计量等这几个常用的基本概念和术语,并着重介绍几个常用的统计量及其抽样分布.

§ 1.1 基本概念

1.1.1 总体及其分布

在数理统计中,我们通常把研究对象的全体元素组成的集合称为**总体或母体**,总体中的每个元素称为**个体**.

例 1.1 如对一批钢筋的抗拉强度进行某种统计推断,该批钢筋中每根钢筋的抗拉强度(是实数)的全体就是总体,而每一根钢筋的抗拉强度就是该总体的一个个体.

例 1.2 检验某灯泡厂生产的一批灯泡的使用寿命,每个灯泡寿命(使用时数)的全体便是总体,每一个灯泡的寿命便是个体.

由上述两例可见,总体中的元素(或个体)是指研究对象的某个指标值,它一般是一个数量值,因此它们的总体都是一些实数集合. 这些指标值也可以是一些属性指标值,例如学生的学习成绩可用属性指标值:优、良、一般、不及格来表示;产品可用属性指标值:一级品、二级品、三级品、次品来表示. 但这些属性指标值也可数量化,如可用 1,2,3,4 分别表示学生的学习成绩优、良、一般、不及格. 用 0,1,2,3 分别表示次品、一级品、二级品、三级品等.

例 1.3 检验某厂 1000 件产品中的次品率. 总体是由正品

(数量化为 1) 和次品(数量化为 0) 指标值构成. 则每一件产品都是一个个体, 它要么属于正品要么属于次品.

总体据其包含个体数目分为有限总体和无限总体. 在例 1.3 中, 因问题只在于估计这批产品的次品率而不及其他, 故总体就是这批产品(是有限总体), 每件产品就是个体. 但按上述总体定义, 此处略有差异, 在问题明确时, 这点微小差异无关紧要. 例 1.1 和例 1.2 中的总体都是无限总体.

数理统计中, 因为我们主要关心的是研究对象的某项或某几项数量指标 X , 在例 1.1 中 X 表示抗拉强度, 例 1.2 中 X 表示灯泡寿命, 例 1.3 中 X 表示正品(数 1) 和次品(数 0). 要特别指明, 总体中的元素不是指元素本身, 而是指元素的某种数量指标 X . 从数学角度讲, 总体是指研究对象某项数量指标可能取的各种不同数值的集合.

进行统计推断时需要从总体中随机抽取部分个体, 每个个体取值依抽取的结果而定. 例如灯泡寿命的指标值可取 850 h, 900 h, 1 000 h, 1 200 h, 等等, 故而数量指标 X 是随机变量. 由此我们把总体可定义为某个随机变量 X 所有可能取值的集合, 并把总体记为 X . 既然数量指标值 X 是一个随机变量, 必然有其概率分布 $P\{X = x\} = p$ 或者分布函数 $F(x)$. 通常所说的总体分布就是指该数量指标 X (随机变量) 的概率分布 $P\{X = x\} = p$ 或分布函数 $F(x)$.

1.1.2 样本(简单随机样本)

要了解总体的某些特性, 一般有两种途径. 一种途径是进行全面试验, 就是把总体所含个体逐个试验, 以图全面了解总体. 如要了解一批新式导弹的性能, 把该批导弹一一试射, 即可全面了解这批导弹的性能. 但此种破坏性的试验无论如何都不可取. 另一种途径是进行统计推断. 所谓统计推断, 实际上是利用部分个体

的某些数字特征去推断(判断)总体的某些数字特征. 如上例, 用试射几枚新式导弹的特性来推断整批导弹的特性, 才是可行之法.

要用部分个体去推断总体, 首先需要从总体中获得这些部分个体. 从总体中得到部分个体, 谓之抽样. 抽样方式甚多, 此处介绍最常用的抽样方法, 就是所谓简单随机抽样. 简单随机抽样以以下两种方法实现: 对无限总体或当总体所含个体数目较大时, 采取随机不返回式方法抽取部分个体; 当总体所含个体数目不太大时, 则采取随机返回式抽取方法抽取部分个体. 从总体中抽取的部分个体称为总体的一个样本(抽样), 用简单随机抽样方法抽取的样本称为简单随机样本. 样本中的每一个个体称为样品, 样本中所含样品的个数称为样本容量.

从总体中抽取一个个体, 就是对总体 X 进行一次抽样观察(一次试验). 一般情况, 不止进行一次观察, 而要进行 n 次观察, 即抽取 n 个个体, 记为 $(X_1, X_2, \dots, X_n)$. 对总体进行 n 次简单随机抽样(n 次观察), 相当于在相同条件下对总体 X 进行 n 次重复独立的观察. 通过 n 次观察, 得到总体 X (数量指标) 的一组观察值 $(x_1, x_2, \dots, x_n)$, 其中 x_i 为第 i ($i = 1, 2, \dots, n$) 次观察记录的结果. 此处 $(x_1, x_2, \dots, x_n)$ 是完全确定的一组数值, 但它又随每次抽样(n 次观察看作一次抽样) 观察而改变, 因此可看作为随机向量 $(X_1, X_2, \dots, X_n)$ 的一次实现. 由于这个随机向量是总体 X 的一个简单随机抽样, 则我们就把它看作总体的一个容量为 n 的样本.

特别注意, 由于抽样的随机性, 从母体抽取容量为 n 的样本, 在抽取之前无法预知其各分量所取之值, 所以样本 $(X_1, X_2, \dots, X_n)$ 是一个 n 维随机向量. 又因为 $X_1, X_2, \dots, X_n$ 是对总体 X 的 n 次观察, 且各次观察是在相同条件下独立进行的, 故有足够的理由确认 $X_1, X_2, \dots, X_n$ 是相互独立的随机变量. 样本 $(X_1, X_2, \dots, X_n)$ 所有可能取值的全体称为样本空间, 一个样本观察值 $(x_1, x_2,$

$\cdots, x_n$) 就是该样本空间中的一个点, 简称样本值(或样本点).

综上所述, 给出如下定义:

定义 1.1 设 X 是具有分布函数 $F(x)$ 的随机变量, 若 $X_1, X_2, \cdots, X_n$ 是服从同一分布函数 $F(x)$ 且相互独立的随机变量, 则称随机向量 $(X_1, X_2, \cdots, X_n)$ 为来自总体 X (或总体 $F(x)$) 的容量为 n 的简单随机样本, 简称样本. 它的观察值 $(x_1, x_2, \cdots, x_n)$ 称为样本值, 又称为总体 X 的 n 个独立观察值.

值得注意的是, 样本具有二重性. 一方面样本本身是 n 维随机向量, 具有随机性. 另一方面, 一经抽取便得到一组确定的数值, 因此又具有确定性. 在不同的场合, 要加以严格区别, 且不可混为一谈. 通常所说的样本是指 n 维随机向量 $(X_1, X_2, \cdots, X_n)$, 而样本观察值 $(x_1, x_2, \cdots, x_n)$ 才是一组确定的数值, 在不致混淆的场合, 也简称二者为样本. 本书此后如不特作说明, 所说样本均指简单随机样本.

1.1.3 样本分布

由上所述, 样本是一 n 维随机向量, 依概率论知, 它必有其概率分布函数. 所谓样本分布即指样本的概率分布函数. 对简单随机样本 $(X_1, X_2, \cdots, X_n)$, 据概率论有关论述知, 它的分布函数(联合分布函数)可由总体 X 的分布函数 $F(x)$ 完全决定, 其表达式为

$$F_n(x_1, x_2, \cdots, x_n) = F(x_1) \cdot F(x_2) \cdot \cdots \cdot F(x_n) = \prod_{i=1}^n F(x_i)$$

特别, 如果总体 X 是连续型的随机变量, 且具有概率密度函数 $f(x)$, 则样本 $(X_1, X_2, \cdots, X_n)$ 的概率密度函数(联合分布密度)为

$$f_n(x_1, x_2, \cdots, x_n) = f(x_1) \cdot f(x_2) \cdot \cdots \cdot f(x_n) =$$

$$\prod_{i=1}^n f(x_i)$$

如果总体 X 是离散型随机变量, 则其分布律为

$$P\{X_1 = x_1, X_2 = x_2, \dots, X_n = x_n\} =$$

$$\prod_{i=1}^n P\{X_i = x_i\} = \prod_{i=1}^n p(x_i)$$

其中 $P\{X = x\} = p(x)$ 为离散型随机变量 X 的概率分布.

1.1.4 统计量与样本矩

上节中, 依总体分布得到了样本分布. 由于对总体的某些特征(特性)未知, 总体分布中常含有一些未知参数, 故而样本分布中也含有相应的未知参数. 这些参数在一定范围内变化, 因此把这种含有未知参数的样本分布称为样本分布族. 如

$$\left\{ \frac{1}{2^{\frac{n}{2}} \pi^{\frac{n}{2}} \sigma^n} \exp \left\{ -\frac{\sum_{i=1}^n (x_i - \mu)^2}{2\sigma^2} \right\} \mid -\infty < \mu < +\infty, 0 < \sigma < +\infty \right\}$$

即为一广泛应用的正态分布族, 其中 μ, σ^2 均为未知参数. 此类样本分布族包含了所抽样本用于推断总体特性的全部信息, 鉴于此因, 也称样本分布族为统计模型. 因而对总体特性(未知参数)的统计推断问题, 就转化为利用样本对统计模型中所含未知参数(数字特征)进行统计推断.

样本是统计推断的依据, 但实际应用时, 一般并非直接利用样本本身, 而是针对不同问题对样本进行“加工”和“提炼”, 把样本中我们最关心、最有用、最利于所论问题的信息尽可能集中起来, 以便对所论问题进行更为有效的统计推断. 在数理统计中, 这种对样本的“加工”和“提炼”, 往往是通过构造一个合适的只依赖于样本的函数(不含任何未知参数)——样本统计量来表达的.

定义 1.2 设 $(X_1, X_2, \dots, X_n)$ 是来自总体 X 的一个样本, 如

果 $g(x_1, x_2, \dots, x_n)$ 是 $(x_1, x_2, \dots, x_n)$ 的一个实值函数, 且 $g(x_1, x_2, \dots, x_n)$ 中不含任何未知参数, 则称 $T = g(X_1, X_2, \dots, X_n)$ 是一个统计量. 若 $(x_1, x_2, \dots, x_n)$ 是样本 $(X_1, X_2, \dots, X_n)$ 的一组观察值, 则称 $t = g(x_1, x_2, \dots, x_n)$ 是统计量 $T = g(X_1, X_2, \dots, X_n)$ 的一个观察值(或统计值).

例 1.4 设总体 $X \sim N(\mu, \sigma^2)$, μ 已知, 但 σ^2 未知, $(X_1, X_2, \dots, X_n)$ 为 X 的一个样本. 则

$$\frac{1}{n} \sum_{i=1}^n X_i, \quad \frac{1}{n} \sum_{i=1}^n (X_i - \mu)^2, \quad \frac{2X_1 + 3X_2}{4}$$

都是统计量, 而

$$\sum_{i=1}^n \left(\frac{X_i - \mu}{\sigma} \right), \quad \frac{1}{n} \sum_{i=1}^n \left(\frac{X_i}{\sigma} \right)^2, \quad \frac{X_5 - X_6}{\sigma}$$

都不是统计量, 这是由于后三者均包含未知参数 σ .

例 1.5 设 (X_1, X_2) 是从泊松总体 $X \sim P(\lambda)$ 中抽取的一个二维样本, 其中 λ 未知. 则 $X_1, X_2 + 3, X_1^2 + X_2^2$ 都是统计量, 但 $(2X_1 + 3X_2) - \lambda, (X_2 - X_1) + \lambda$ 都不是统计量, 其原因是后二者都含有未知参数 λ .

下面给出几个常用的统计量.

定义 1.3 设 $(X_1, X_2, \dots, X_n)$ 是从总体 X 抽取的一个样本, $(x_1, x_2, \dots, x_n)$ 是该样本观察值, 则称

$$(1) \quad \bar{X} = \frac{1}{n} \sum_{i=1}^n X_i$$

为样本均值;

$$(2) \quad S^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2$$

为样本方差;

$$(3) \quad S = \sqrt{S^2} = \sqrt{\frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2}$$

为样本标准差；

$$(4) \quad A_k = \frac{1}{n} \sum_{i=1}^n X_i^k \quad (k = 1, 2, \dots)$$

为样本 k 阶原点矩；

$$(5) \quad B_k = \frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^k \quad (k = 1, 2, \dots)$$

为样本 k 阶中心矩，且它们相应的观察值依次为

$$(1) \quad \bar{x} = \frac{1}{n} \sum_{i=1}^n x_i$$

$$(2) \quad s^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2$$

$$(3) \quad s = \sqrt{\frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2}$$

$$(4) \quad a_k = \frac{1}{n} \sum_{i=1}^n x_i^k \quad (k = 1, 2, \dots)$$

$$(5) \quad b_k = \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^k \quad (k = 1, 2, \dots)$$

略微细心一点就可发现，样本均值和样本方差与样本 k 阶原点矩有一些关系，例如

$$(1) \quad \bar{X} = \frac{1}{n} \sum_{i=1}^n X_i = A_1$$

$$(2) \quad S^2 = \frac{n}{n-1} A_2 - \frac{n}{n-1} A_1^2 = \frac{1}{n-1} \sum_{i=1}^n X_i^2 - \frac{n}{n-1} \bar{X}^2$$

(1) 式是显然的；(2) 式按上述定义有

$$\begin{aligned} S^2 &= \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2 = \\ &= \frac{n}{n-1} \left[\frac{1}{n} \sum_{i=1}^n (X_i^2 - 2\bar{X}X_i + \bar{X}^2) \right] = \end{aligned}$$

$$\begin{aligned}
& \frac{n}{n-1} \left[\frac{1}{n} \sum_{i=1}^n X_i^2 - 2\bar{X} \frac{1}{n} \sum_{i=1}^n X_i + \frac{1}{n} \sum_{i=1}^n \bar{X}^2 \right] = \\
& \frac{n}{n-1} \left[\frac{1}{n} \sum_{i=1}^n X_i^2 - 2\bar{X}^2 + \bar{X}^2 \right] = \\
& \frac{n}{n-1} \left[\frac{1}{n} \sum_{i=1}^n X_i^2 - \bar{X}^2 \right] = \\
& \frac{1}{n-1} \sum_{i=1}^n X_i^2 - \frac{n}{n-1} \bar{X}^2 = \frac{n}{n-1} A_2 - \frac{n}{n-1} A_1^2
\end{aligned}$$

其一般关系在此不再细论,有兴趣的读者可参看其他数理统计书籍. 但要注意,这里等式(2)中的简单推证方法我们是要熟悉、掌握的.

样本 k 阶原点矩的一个重要应用是将其作为总体未知参数点估计的矩估计量,为说明其合理性,在此回顾一下概率论中的有关大数定律.

定理 1.1 (辛钦大数定律) 设 $X_1, X_2, \dots$ 是一列独立同分布的随机变量,且数学期望存在

$$E(X_i) = \mu \quad (i = 1, 2, \dots)$$

则对任意的 $\epsilon > 0$, 有

$$\lim_{n \rightarrow \infty} P \left\{ \left| \frac{1}{n} \sum_{i=1}^n X_i - \mu \right| < \epsilon \right\} = 1 \quad (1.1)$$

成立. 借用数学分析中大家所熟悉的“收敛”、“极限”这些术语,把式(1.1)所表示的关系记为

$$\frac{1}{n} \sum_{i=1}^n X_i \xrightarrow{P} \mu \quad (n \rightarrow \infty)$$

并且称 $\frac{1}{n} \sum_{i=1}^n X_i$ 依概率收敛于 μ .

定理 1.2 设 $X_1, X_2, \dots, X_n$ 是来自总体 X 的一个样本,且总体 X 的 k 阶原点矩 $E(X^k) = \mu_k$ 和 $2k$ 阶原点矩 $E(X^{2k}) = \mu_{2k}$ ($k =$