

高等学校试用教材

概 率 论

第二册 数理统计

第一分册

复旦大学 编

人 民 教 育 出 版 社

1

42

高等学校试用教材
概 率 论

第二册 数理统计

第一分册

复旦大学 编

人民教育出版社 出版

新华书店北京发行所发行

北京印刷一厂印装

开本 $787 \times 1092 \frac{1}{32}$ 印张 $4 \frac{10}{16}$ 字数 110,000

1979年5月第1版 1980年3月第2次印刷

印数 28,001—58,500

书号 13012·0352 定价 0.36 元

目 录

第二册 数理统计

第一章 引论	1
§ 1. 统计学	1
§ 2. 数理统计的基本概念	3
§ 3. 求估计量的两种常用方法	10
习题	20
第二章 抽样分布	22
§ 1. 正态母体子样的线性函数的分布	23
§ 2. χ^2 -分布	25
§ 3. t -分布和 F -分布	31
§ 4. 正态母体子样均值和方差的分布	38
习题	41
第三章 假设检验 (I)	43
§ 1. 引言	43
§ 2. 正态母体参数的检验	49
§ 3. 正态母体参数的置信区间	65
§ 4. 多项分布的 χ^2 -检验	68
§ 5. 广义似然比检验	82
习题	86
第四章 线性统计推断	89
§ 1. 最小二乘法	89
§ 2. 回归分析	113
§ 3. 方差分析	123
习题	140

第一章 引 论

§1. 统 计 学

统计学是一门关于数据资料的收集、整理、分析和推断的科学。但人们常常将统计这一概念误解为大量数据资料的收集以及对这些数据作一些简单的运算(如求和、求平均值、求百分比等)或用图表、表格等形式把它们表示出来。其实这些工作仅是统计学工作的非主要部分。统计学还包括怎样设计试验、采集数据以及怎样对获得的数据进行分析、推断等其它许多工作。

随着研究随机现象规律性的科学——概率论的发展,应用概率论的结果更深入地分析研究统计资料,通过对某些现象的频率的观察来发现该现象的内在规律性,并作出一定精确程度的判断和预测;将这些研究的某些结果加以归纳整理,逐渐形成一定的数学概型。这些组成了数理统计的内容。

数理统计的方法及考虑的问题不同于一般的资料统计,它更侧重于应用随机现象本身的规律性来考虑资料的收集、整理和分析,从而找出相应的随机变量的分布律或它的数字特征。由于大量的随机试验必能呈现出它的规律性,因而从理论上讲,只要对随机现象进行足够多次观察,被研究的随机现象的规律性一定能清楚地呈现出来,但是实际上所允许的观察永远只能是有限的,有时甚至是少量的。因此我们所关心的问题是怎样有效地利用有限的资料,便能去掉那些由于资料不足所引起的随机干扰,而把那些实质性的东西找出来。一个好的统计方法就在于能有效地利用所获得的资料,尽可能作出精确而可靠的结论。

在数理统计里,不是对所研究的全部对象进行观察,而是抽取其中的部分进行观察、获得数据(即采样),并通过这些数据来对所研究的全体进行推断。由于推断是基于采样数据,而采样数据又不能包含研究对象的全部信息,因此由此所获得的结论必然会包含不定性。概率是这种不定性的度量。造成不定性的原因可分为两类:(1)由于采样数据的随机性所引起的不定性;(2)由于我们对系统真实状态的“无知”造成的不定性。数理统计工作者的任务就是要分辨这两种不定性。下面举一例来说明。

某元件厂生产了一批三极管,共一百万只,每十只装成一盒,共得十万盒。现有某仪器厂需向该元件厂购买此三极管一百盒,准备安装在某种仪表上。每台仪表需用此三极管十只,恰好是一台仪表需一盒三极管,一百盒可供装一百台。但是该仪表对三极管有一定的质量要求,要求十只中至少有八只是一级品,其余的可以是二级品,否则仪表将不能稳定工作。此时仪器厂对这批三极管就面临两种不定性需要分辨:(1)元件厂生产的十万盒三极管,对仪器厂来说是满意的(即一盒中至少有八只是一级品)盒子所占比例 p 是多少?(2)由于有十万盒三极管,现在仅购买其中的一百盒,因而就面临着另一种不定性。即使已知此十万盒中,满意的盒子所占的比例为 p ,又怎样能确定买来的一百盒中,满意的占多少比例呢?例如 $p=0.99$,即十万盒中大约有九万九千盒是满意的,这个比例对仪器厂来说应该是很好了,但也有可能发生这样的意外,即仪器厂所购买的一百盒全部落在不满意的大约一千盒之中。

第一种不定性是不知道 p ,是我们对系统真实状态的“无知”;而后一种不定性是由于所谓“随机性”造成的。为了改善这不定性,仪器厂可要求元件厂对这批三极管的质量进行测试,也就是要求抽取部分三极管进行测试,通过这部分中一级品所占的比例(概

率)来对 p 的真实值进行推断。当然我们不能完全精确地决定 p , 但是我们可以希望获得一个(在某种意义上)比较好的判断。这就涉及到怎样设计试验, 决定观察的数目, 和怎样利用试验观察的结果作出一个“好的”推断等, 这些都是数理统计所要研究的问题。至于在已知 p 的条件下, 第二种不定性的程度已在概率论基础部分作过讨论。

数理统计研究的内容随着科学技术和生产的不断发展而逐步扩大, 但概括地说可以分为两大类: (1) 试验的设计和研究的, 即研究如何更合理更有效地获得观察资料的方法; (2) 统计推断, 即研究如何利用一定的资料对所关心的问题作出尽可能精确、可靠的结论。当然这两部分有密切联系, 在实际应用中更应前后兼顾。但因限于篇幅, 本书将只讨论统计推断。

数理统计方法提供了在不确定性占优势的情况下进行判断的工具, 它为许多科学研究部门和生产、经济部门的工作者所采用。它不仅为提高产量、质量起直接的推动作用, 而且也提供了从大量现象中发现某些事物发展规律的方法。

§ 2. 数理统计的基本概念

为了更好地介绍数理统计所研究的问题, 先引入一些常用的概念和术语。

一、母体和子样

我们今后所讨论的统计问题主要属于下面这种类型: 从一个集合中选取一部分元素, 对这部分元素的某些数量指标进行测量, 根据测量获得的这些数据来推断这集合中全部元素的这些数量指标的分布情况。在统计学中, 我们把所研究的全部元素组成的集合称为母体, 或总体。而把组成母体的每个元素称为个体。例如在研究某批灯泡的平均寿命时, 该批灯泡的全体就组成了母体, 而

其中每个灯泡就是个体。在研究上海市男大学生的身高和体重的分布情况时,上海市的全体男大学生组成了母体,上海市的每个男大学生是个体。但是在统计里,由于我们关心的不是每个个体的种种具体特性,而仅仅是它的某一项或某几项数量指标 X (可以是向量) 和该数量指标 X 在母体中的分布情况。在上述例子中 X 是表示灯泡的寿命或男大学生的身高和体重。就此数量指标 X 而言,每个个体所取的值是不同的。在试验中,抽取了若干个个体就观察到了 X 的这样或那样的数值,因而这个数量指标 X 是一个随机变量(或向量),而 X 的分布就完全描写了母体中我们所关心的那个数量指标的分布状况。由于我们关心的正是这个数量指标,因此我们以后就把母体和数量指标 X 可能取值的全体组成的集合等同起来。所谓母体的分布也就是指数量指标 X 的分布。

一般说来母体这个集合的大小与所研究的问题有关,如在灯泡的自动化生产中,为了保证正常生产需每隔一定时间抽取几个灯泡测试其寿命,并认为灯泡的寿命 X 服从正态分布。在这里母体被想象为所生产的灯泡寿命的可能取值的全体所组成的集合,这样母体被视为由无限个元素组成。并视 X 为连续型随机变量。

为了对母体的分布律进行各种研究,就必需对母体进行抽样观察。一般说来,我们还不止进行一次抽样观察,而要进行几次观察。通过观察就得到母体指标 X 的一组数值 $(x_1, x_2, \dots, x_n)$, 其中每个 x_i 是一次抽样观察的结果。即某一个被观察的个体的 X 指标值。 $(x_1, \dots, x_n)$ 称为容量为 n 的子样的观察值。由于我们是利用子样观察来对母体的分布进行推断,因而从母体中抽取子样进行观察时必须是随机的。直观地说,如果我们要研究上海市男大学生的身高的分布情况,那么在抽样时就希望上海市的每个男大学生具有同等的可能被抽到测量身高,因为只有这样才能经过多次观察比较全面地了解母体。所以对于随机抽样来说,对其某一

次观察结果而论, $(x_1, x_2, \dots, x_n)$ 是完全确定的一组值, 但它又是随每次抽样观察而改变的, 由于我们要依据这一观察结果进行分析推断, 并研究比较各种推断方法的好坏, 因而一般考虑问题时, 就不能把 $(x_1, x_2, \dots, x_n)$ 看为确定的数值, 而应该看作为随机向量 $X = (X_1, X_2, \dots, X_n)$ ①. 称它为容量是 n 的子样, 因而对子样也有分布可言. 在不同的抽样观察中, X 得到不同的现实. $X = (X_1, X_2, \dots, X_n)$ 所可能取值的全体(这儿是 n 维空间或其中的一个子集)称为子样空间. 一个子样观察值 $(x_1, x_2, \dots, x_n)$ 就是子样空间中的一个点.

我们抽取子样的目的是为了对母体的分布律进行各种分析推断, 因而要求抽取的子样能很好地反映母体的特性, 这就必须对随机抽样的方法提出一定的要求. 通常提出下面两点: (i) 代表性: 要求子样的每个分量 X_i 与所考察的母体 X 具有相同的分布 $F(x)$; (ii) 独立性: $X_1, X_2, \dots, X_n$ 为相互独立的随机变量, 也就是说, 每个观察结果既不影响其它观察结果, 也不受其它观察结果的影响.(在有限母体中, 子样的各个观察结果可以是不独立的.)

满足上述两点性质的子样称为简单随机子样, 获得简单随机子样的抽样方法称为简单随机抽样. 在今后, 如果不作特殊声明, 所说的子样将理解为简单随机子样, 对于简单随机子样 $X = (X_1, X_2, \dots, X_n)$, 其分布可以由母体 X 的分布函数 $F(x)$ 完全决定, X 的分布函数是 $\prod_{i=1}^n F(x_i)$.

[例 1] 某血吸虫防治所为了研究某地区血吸虫病的流行情况, 对该地区所属的各生产队编制了一组卡片, 每张卡片上标明一个生产队的血吸虫病的感染率. 并将感染率的高低分为 4 类(无患者, 轻度流行, 中度流行和重度流行). 分别用数字 0, 1, 2, 3 表

① 第二册中, 除协方差矩阵符号用黑体外, 其他矩阵、向量符号均不用黑体.

示这4类,那么在每张卡片上可以标出数字0到3中的一个,以表示该生产队血吸虫病的流行情况属于那一类.用 $\times$ 记这数字,这 $\times$ 就是我们要研究的指标.

现在,母体是全部卡片,设一共有 N 张.如果用返回抽取方法从这些卡片中任意抽取 $n=10$ 张,即每抽一张卡片记录下该卡片上的感染率的类别数字后立即放回去,然后再抽下一张.在这里各次抽得的卡片其指标 X 取值 k ($k=0, 1, 2, 3$)的概率不变,从抽得的10张卡片上观察到的随机变量 $\times$ 的值 $(x_1, x_2, \dots, x_{10})$ 是一个10维随机向量的现实,这个10维随机向量的观察值构成一个子样;子样的容量 $n=10$;子样空间是由所有可能的各组 $(x_1, x_2, \dots, x_{10})$ 组成.由于 x_k ($k=1, 2, \dots, 10$)只可能取值0, 1, 2, 3,因此现在的子样空间由 4^{10} 个点组成.

二、统计量和子样矩

子样是母体的代表和反映,但在我们抽取子样之后,并不直接利用子样进行推断.而需要对子样进行一番“加工”和“提炼”,把子样中所包含的关于我们所关心的事物的信息集中起来.这便是针对不同的问题构造出子样的某种函数.这种函数在统计学中称为统计量.严格地说,一个统计量就是观察的 n 维随机向量即子样 $X=(X_1, X_2, \dots, X_n)$ 的一个(波雷尔可测)函数,且要求它不包含有任何未知参数.因此它也是一个随机变量(或向量).

若设 X_1, X_2 是从具有分布密度为 $N(\mu, \sigma^2)$ 的正态母体中抽取的一个2维子样,其中 μ, σ^2 是未知参数.则 $\frac{1}{2}(X_1+X_2)-\mu, \frac{X_1}{\sigma}$ 都不是统计量,因为它们含有未知参数.而 $X_1, X_1+3, X_1^2+X_2^2$ 却都是统计量.

下面先介绍一些常用的统计量——子样矩.

设 $(X_1, X_2, \dots, X_n)$ 是从母体 $\times$ 中抽取的一个子样,称统计量

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

为子样均值; 统计量

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

为子样方差; 统计量

$$A_\nu \triangleq \frac{1}{n} \sum_{i=1}^n X_i^\nu \quad (3)$$

为子样的 ν 阶矩(或 ν 阶原点矩); 统计量

$$B_\nu \triangleq \frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^\nu \quad (4)$$

为子样的 ν 阶中心矩.

现在我们研究子样矩的期望与方差, 在下面的讨论中, 我们用 μ 表示母体的均值, σ^2 表示母体的方差, α_k 表示母体的 k 阶原点矩, μ_k 表示母体的 k 阶中心矩, 即记

$$\begin{aligned} EX &\triangleq \mu, D(X) = E(X - \mu)^2 \triangleq \sigma^2 \\ EX^k &\triangleq \alpha_k, E(X - \mu)^k \triangleq \mu_k \end{aligned} \quad (5)$$

并且约定, 在我们用到 α_k (或 μ_k) 时, 假定它是存在的.

定理 1 设母体 X 服从分布 $F(x)$, $X = (X_1, \dots, X_n)$ 是从该母体中抽得的一个简单随机子样, 如果 $F(x)$ 的二阶矩存在, 则对子样均值 $\bar{X}$, 有

$$E\bar{X} = \mu$$

$$D(\bar{X}) = \frac{\sigma^2}{n}$$

$$[\text{证明}] \quad E(\bar{X}) = E\left(\frac{1}{n} \sum_{i=1}^n X_i\right) = \frac{1}{n} \sum_{i=1}^n E(X_i) = \frac{1}{n} \sum_{i=1}^n \mu = \mu$$

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

定理 2 如果 $F(x)$ 存在 4 阶矩, 则对子样方差, 有

$$E S_n^2 = \frac{n-1}{n} \sigma^2$$

$$D(S_n^2) = \frac{\mu_4 - \mu_2^2}{n} - \frac{2(\mu_4 - 2\mu_2^2)}{n^2} + \frac{\mu_4 - 3\mu_2^2}{n^3} \quad (\mu_2 = \sigma^2)$$

$$\begin{aligned} \text{[证明]} \quad E S_n^2 &= E \left\{ \frac{1}{n} \sum_{i=1}^n X_i^2 - \frac{1}{n^2} \left(\sum_{i=1}^n X_i \right)^2 \right\} \\ &= \alpha_2 - \frac{1}{n^2} E \left\{ \sum_{i=1}^n X_i^2 + \sum_{i \neq j} X_i X_j \right\} \\ &= \alpha_2 - \frac{1}{n^2} (n\alpha_2 + n(n-1)\mu^2) \\ &= \frac{n-1}{n} (\alpha_2 - \mu^2) = \frac{n-1}{n} \sigma^2 \end{aligned}$$

现记 $Y_i \triangleq X_i - \mu$, 则 $EY_i = 0$, $D(Y_i) = \sigma^2$, $EY_i^4 = \mu_4$, 且有

$$\begin{aligned} n^2 (S_n^2)^2 &= \left[\sum_{i=1}^n (X_i - \bar{X})^2 \right]^2 = \left[\sum_{i=1}^n (X_i - \mu)^2 - n(\bar{X} - \mu)^2 \right]^2 \\ &= \left(\sum_{i=1}^n Y_i^2 - n\bar{Y}^2 \right)^2 = \left(\sum_{i=1}^n Y_i^2 \right)^2 - \frac{2}{n} \left(\sum_{i=1}^n Y_i^2 \right) \left(\sum_{j=1}^n Y_j \right)^2 \\ &\quad + \frac{1}{n^2} \left(\sum_{j=1}^n Y_j \right)^4 \\ &= \left(\sum_{i=1}^n Y_i^4 + \sum_{i \neq j} Y_i^2 Y_j^2 \right) - \frac{2}{n} \left(\sum_{i \neq j} Y_i^2 Y_j^2 + \sum_{i=1}^n Y_i^4 \right) \\ &\quad + \sum_k \sum_{i \neq j} Y_k^2 Y_i Y_j + \frac{1}{n^2} \left(\sum_{i=1}^n Y_i^4 + 3 \sum_{i \neq j} Y_i^2 Y_j^2 \right. \\ &\quad \left. + 2 \sum_k \sum_{i \neq j} Y_k^2 Y_i Y_j + \sum Y_i Y_j Y_k Y_l \right) \textcircled{1} \end{aligned}$$

对上式两边取期望得:

$$\begin{aligned} n^2 E(S_n^2)^2 &= n\mu_4 + n(n-1)\mu_2^2 - \frac{2}{n} [n(n-1)\mu_2^2 \\ &\quad + n\mu_4] + \frac{1}{n^2} [n\mu_4 + 3n(n-1)\mu_2^2] \\ &= \left(n-2 + \frac{1}{n} \right) \mu_4 + \left(n-2 + \frac{3}{n} \right) (n-1) \mu_2^2 \end{aligned}$$

① Σ' 表示对 $i \neq k, j \neq l, (i, j) \neq (l, k)$ 的 i, j, k, l 求和,

所以

$$\begin{aligned}
 D(S_n^2) &= E(S_n^2)^2 - (ES_n^2)^2 \\
 &= \frac{1}{n^3} \left[\left(n-2 + \frac{1}{n} \right) \mu_4 + \left(n-2 + \frac{3}{n} \right) (n-1) \mu_2^2 \right] \\
 &\quad - \left(\frac{n-1}{n} \right)^2 \mu_2^2 \\
 &= \frac{1}{n^3} \left(n-2 + \frac{1}{n} \right) \mu_4 + (n-1) (3-n) \frac{\mu_2^2}{n^3} \\
 &= \frac{\mu_4 - \mu_2^2}{n} - \frac{2(\mu_4 - 2\mu_2^2)}{n^2} + \frac{\mu_4 - 3\mu_2^2}{n^3}
 \end{aligned}$$

定理 3 如果 $F(x)$ 存在 2ν 阶矩, 则对子样的 ν 阶原点矩 A_ν , 有

$$EA_\nu = \alpha_\nu$$

$$D(A_\nu) = \frac{\alpha_{2\nu} - \alpha_\nu^2}{n}$$

$$[\text{证明}] \quad EA_\nu = E\left(\frac{1}{n} \sum_{i=1}^n X_i^\nu\right) = \frac{1}{n} \sum_{i=1}^n E(X_i^\nu) = \alpha_\nu$$

$$\begin{aligned}
 D(A_\nu) &= EA_\nu^2 - (EA_\nu)^2 = E\left(\frac{1}{n} \sum_{i=1}^n X_i^\nu\right)^2 - \alpha_\nu^2 \\
 &= E\left(\frac{1}{n^2} \sum_{i=1}^n X_i^{2\nu} + \frac{1}{n^2} \sum_{i \neq j} X_i^\nu X_j^\nu\right) - \alpha_\nu^2 \\
 &= \frac{1}{n} \alpha_{2\nu} + \frac{1}{n^2} \cdot n(n-1) \alpha_\nu^2 - \alpha_\nu^2 = \frac{\alpha_{2\nu} - \alpha_\nu^2}{n}
 \end{aligned}$$

三、顺序统计量和经验分布函数

设 $(X_1, \dots, X_n)$ 是从母体 X 中抽取的一个子样. 记 $(x_1, x_2, \dots, x_n)$ 是子样的一个观察值, 将观察值的各分量按大小递增次序排列, 得到

$$x_1^* \leq x_2^* \leq \dots \leq x_n^*$$

当 $(X_1, \dots, X_n)$ 取值为 $(x_1, \dots, x_n)$ 时, 我们定义 $X_k^{(n)}$ 取值为 x_k^* . 称由此得到的 $X_1^{(n)}, \dots, X_n^{(n)}$ 为 $(X_1, \dots, X_n)$ 的一组顺序统计量. 显然 $X_1^{(n)} \leq X_2^{(n)} \leq \dots \leq X_n^{(n)}$, $X_1^{(n)} = \min_{1 \leq i \leq n} X_i$, 即 $X_1^{(n)}$ 的观察

值是子样观察值中最小的一个, 而 $X_n^{(n)} = \max_{1 \leq i \leq n} X_i$, 即 $X_n^{(n)}$ 的观察值是子样观察值中最大的一个. 关于顺序统计量的进一步讨论可参阅第 7 章 § 2.

记

$$F_n^*(x) \triangleq \begin{cases} 0, & \text{当 } x \leq x_1^* \\ \frac{k}{n}, & \text{当 } x_k^* < x \leq x_{k+1}^*, \quad k=1, 2, \dots, n-1 \\ 1, & \text{当 } x > x_n^* \end{cases}$$

显然 $0 \leq F_n^*(x) \leq 1$, 且作为 x 的函数是一非减左连续函数, 把 $F_n^*(x)$ 看作为 x 的函数, 它具备分布函数所要求的性质, 故称为经验分布函数(或子样分布函数).

经验分布函数也是子样的函数, 它与子样矩之间具有下列关系: 设 $(x_1, x_2, \dots, x_n)$ 是子样观察值, $F_n^*(x)$ 是对应的经验分布函数, 则有:

$$\begin{aligned} \bar{x} &= \frac{1}{n} \sum_{i=1}^n x_i = \int x dF_n^*(x) \\ s_n^2 &= \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2 = \int (x - \bar{x})^2 dF_n^*(x) \\ a_\nu &= \frac{1}{n} \sum_{i=1}^n x_i^\nu = \int x^\nu dF_n^*(x), \quad \nu=2, 3, \dots \\ b_\nu &= \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^\nu = \int (x - \bar{x})^\nu dF_n^*(x), \quad \nu=3, 4, \dots \end{aligned}$$

§ 3. 求估计量的两种常用方法

一、问题的提出

参数估计是统计推断的基本问题之一, 这类问题对我们并不生疏. 如在第一册例 1.3.7 中提出要求估计鱼池中鱼的条数; 在 § 2 中提出要求估计上海市男大学生的平均身高等, 这些都是参数估计问题. 下面再举几个例子.

〔例1〕 灯泡厂生产的灯泡，由于种种随机因素的影响，每批生产出来的灯泡，其中每个灯泡的使用寿命是不一致的，也就是说，灯泡的使用寿命 X 是一个随机变量。由中心极限定理和实际经验知道，灯泡使用寿命 X 服从正态 $N(\mu, \sigma^2)$ 分布。但一般我们事先并不能确切知道 (μ, σ^2) 的具体数值，而只知道它们是落在某个范围中，例如 $(0, \infty) \times (0, \infty)$ 或某个更小的区域里。为了断定所生产的这批灯泡的质量，自然提出要求估计这批灯泡的平均寿命以及寿命长短的差异程度。即要求估计 μ 和 σ^2 。有时还希望以一定的可靠性来估计其平均寿命 μ 是在某个范围内或者不低于某个数。

〔例2〕 大家知道，纺织厂细纱机上的断头次数可以用普阿松分布 $P(\lambda)$ 来描述，现在希望知道每只纱锭在某一时间间隔内断头数为 k 次的概率 $(k=0, 1, 2, 3, \dots)$ 也就是要找出母体的分布律。而对于普阿松分布，只要知道数学期望 λ ——平均断头次数，就可以确定其分布律 $P(\lambda)$ 。同样我们常常不知道 λ 的确切值，而只知道它落在某个范围里，例如 $(0, \infty)$ 或某个区间中。在这里，我们也提出了要求估计母体分布的均值的问题。

由上述可见，在参数估计问题中，我们总是首先假设母体 X 具有一族可能的分布 F ，且 F 的函数形式是已知的，仅包含有几个未知参数。记 θ 是支配这分布的未知参数(可以是向量)，在统计学上，我们把分布 F 的未知参数 θ 的全部可容许值组成的集合称为参数空间，记为 Θ 。

今后，我们用 $F(\cdot; \theta)$ 表示 X 的分布，又称集合 $\{F(\cdot; \theta), \theta \in \Theta\}$ 为 X 的分布函数族。类似地，如果 X 是连续型随机变量，我们有概率密度函数族；如果 X 是离散型随机变量，我们有概率分布族。

在例1中，灯泡寿命 X 服从 $N(\mu, \sigma^2)$ 分布，参数空间是 $\Theta = \{(\mu, \sigma^2); 0 < \mu < \infty, \sigma^2 > 0\}$ 。 $\{N(\mu, \sigma^2), (\mu, \sigma^2) \in \Theta\}$ 就是

X 的分布函数族。

在例 2 中记断头次数为 X , X 服从 $P(\lambda)$, 此处参数空间是 $\Theta = \{\lambda: 0 < \lambda < \infty\}$, $\{P(\lambda); \lambda \in \Theta\}$ 是 X 的概率分布族。

一个参数估计问题就是要求通过子样估计母体分布所包含的未知参数 θ 或 θ 的某函数。

由于母体分布的函数形式是已知的, 未知的仅是一个或几个未知参数, 而母体的真分布也完全由这些参数所完全决定。因此通过估计参数可以估计母体的真分布。但我们有时也直接感兴趣于估计这些参数的值。

那末如何估计未知参数呢? 仍以例 1 为例, 为了估计灯泡的平均寿命, 当然要抽取若干个灯泡作试验(即抽取子样), 若它们的寿命分别是 $X_1, X_2, \dots, X_n$ 。由大数定律知道, 当 n 很大时, $\bar{X} = \frac{1}{n}(X_1 + \dots + X_n)$ 以很大的概率与 $\mu = EX$ 任意接近, 因而自然地吧子样寿命的平均值 $\bar{X}$ 作为母体平均寿命 μ 的估计量。

一般地, 设母体具有分布族 $\{F(x; \theta), \theta \in \Theta\}$, $X_1, X_2, \dots, X_n$ 是它的一个子样。点估计问题就是要求构造一个统计量 $T(X_1, \dots, X_n)$ 作为参数 θ 的估计(T 的维数与 θ 的维数相同)。在统计学上, 我们称 T 为 θ 的估计量。如果 $x_1, \dots, x_n$ 是子样的一组观察值, 代入统计量 T 就得到 T 的具体数值, 这个数值常称为 θ 的估计值, 今后, 估计量和估计值这两个名词将不强调它们的区别通称为估计。

除了点估计问题以外, 还有另一类估计问题, 它要求用区间 (T_1, T_2) 作为 θ 可能取值范围的一种估计。这个区间称为置信区间, 这类估计问题称为区间估计问题。由于参数的区间估计方法与假设检验具有密切的联系, 所以放在 § 3.3 中介绍。

在这里, 估计量 T 是子样 $X_1, \dots, X_n$ 的函数, 它不可以包含

有未知参数,也就是说 T 是一个作估计用的统计量. 是随机变量(向量). 当我们获得了一组子样观察值 $(x_1, \dots, x_n)$ 之后, 就用 $T(x_1, \dots, x_n)$ 作为未知参数 θ 的估计值. 对于不同的子样观察值, 所得的估计值是不同的. 寻找一个估计量就是要找一个依据观察结果求得未知参数 θ 的估计值的方法. 而不是具体去寻找一个估计值. 衡量一个估计的好坏当然不能按一次试验结果得到的参数估计值与参数真值的偏差大小来确定, 而必需从总体出发, 也就是要以多次试验结果所得的估计值与参数真值的偏差大小来确定. 这方面的详细讨论请阅第五章, 这里仅介绍两种常用的求估计量的方法.

二、矩方法

矩方法是一种古老的估计方法. 大家知道, 矩是描写随机变量的最简单的数字特征. 子样来自于母体, 从定理 2.1~2.3 看到子样矩在一定程度上也反映了母体矩的特征, 因而自然想到用子样矩作为母体矩的估计.

设 $\{F(x; \theta), \theta \in \Theta\}$ 是母体 X 的可能分布族, $\theta = (\theta_1, \dots, \theta_k)$ 是待估计的未知参数. 假定母体分布的 k 阶矩存在, 则母体分布的 ν 阶矩

$$\alpha_\nu(\theta_1, \dots, \theta_k) = \int_{-\infty}^{\infty} x^\nu dF(x; \theta), \quad 1 \leq \nu \leq k,$$

是 $\theta = (\theta_1, \dots, \theta_k)$ 的函数.

对于子样 $X = (X_1, \dots, X_n)$, 其 ν 阶子样矩是

$$A_\nu = \frac{1}{n} \sum_{i=1}^n X_i^\nu, \quad 1 \leq \nu \leq k$$

现在用子样矩作为母体矩的估计, 即令

$$\alpha_\nu(\theta_1, \dots, \theta_k) = A_\nu = \frac{1}{n} \sum_{i=1}^n X_i^\nu, \quad \nu = 1, 2, \dots, k \quad (1)$$

这样, (1) 式确定了包含 k 个未知参数 $\theta = (\theta_1, \dots, \theta_k)$ 的 k 个

方程式。解此方程组 (1) 就可以得到 $\theta = (\theta_1, \dots, \theta_k)$ 的一组解 $\hat{\theta} = (\hat{\theta}_1, \dots, \hat{\theta}_k)$ 。因为 A_ν 是随机变量, 故解得的 $\hat{\theta}$ 也是随机变量。现在将 $\hat{\theta}_1, \dots, \hat{\theta}_k$ 分别作为 $\theta_1, \dots, \theta_k$ 的估计, 称为矩方法的估计, 这种求估计量的方法称为矩方法。

[例 3] 母体均值和方差的矩估计。

设 $X_1, \dots, X_n$ 是一子样, 设母体的二阶矩存在, 则有 $\alpha_2 = \sigma^2 + \mu^2$ 。用矩方法得方程组

$$\begin{cases} \hat{\mu} = \frac{1}{n} \sum_{i=1}^n X_i = \bar{X} \\ \hat{\mu}^2 + \hat{\sigma}^2 = \hat{\alpha}_2 = \frac{1}{n} \sum_{i=1}^n X_i^2 \end{cases}$$

解之得

$$\hat{\mu} = \bar{X}$$

$$\hat{\sigma}^2 = \frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^2 = S_n^2$$

所以母体均值 μ 和方差 σ^2 的矩估计分别是子样均值 $\bar{X}$ 和子样方差 S_n^2 。

运用定理 2.1 有

$$E\hat{\mu} = E\bar{X} = \mu$$

$$D(\hat{\mu}) = \frac{1}{n} \sigma^2$$

和

$$E(\hat{\sigma}^2) = \frac{n-1}{n} \sigma^2$$

由此可见, $\hat{\mu}$ 作为 μ 的估计它是在 μ 的真值的周围波动, 且其平均值恰好是真值 μ 。这一性质在统计学上称为无偏性。

一般地, 如果 $T(X)$ 是未知参数 θ 的一个估计量, 且满足下面的关系式,

$$E_\theta T(X) = \theta, \text{ 对一切 } \theta \in \Theta \quad (2)$$

则称 $T(X)$ 是 θ 的无偏估计。

由上可见, 子样均值 $\bar{X}$ 是母体均值 μ 的无偏估计, 同样如令