



第 1 章

预 备 知 识

本章将简要地介绍排列组合、概率统计、调查等方面的有关知识. 这些知识属于抽样技术的基础知识或预备知识, 在以后各章的叙述中会用到, 但又不属于本书中任何其他章节的介绍范围. 补充这些知识供读者在学习时参考, 能够使读者更加顺利地学好以后各章节的内容. 建议读者特别是那些没有系统学习过调查和数理统计的读者, 对本章所介绍的内容能够认真阅读, 因为这些内容不仅是进一步学习抽样技术的准备性知识, 同时这些知识本身也是十分重要和有用的.

1.1 排 列 组 合

1.1.1 基本原理

首先叙述两条基本原理, 这两条原理在排列组合分析公式的推导中起着重要作用.

(1) 加法原理: 假如完成一件事有 m 种方式, 第一种方式有 n_1 种方法, 第二种方式有 n_2 种方法, $\dots$, 第 m 种方式有 n_m 种方法, 而无论通过哪种方式方法都可以完成这件事, 则完成这件事总共有 $n_1 + n_2 + \dots + n_m$ 种方法.

(2) 乘法原理: 假如完成一件事有 m 个步骤, 第一个步骤有 n_1 种方法, 第二个步骤有 n_2 种方法, $\dots$, 第 m 个步骤有 n_m 种方法, 而每一步骤都不可缺少也不可重复, 则完成这件事共有 $n_1 n_2 \dots n_m$ 种方法.

1.1.2 排列

1. 选排列与全排列

定义 1.1 从 n 个不同的元素里, 每次任意取出 $r (1 \leq r \leq n)$ 个不同元素, 按照一定的顺序排成一列, 称为从 n 个不同元素中每次取 r 个不同元素的排列. 如果 $r < n$, 这样的排列称为选排列, 排列种数用 A_n^r 表示; 如果 $r = n$, 则称为全排列, 排列种数用 P_n 表示. 因此

$$A_n^r = n(n-1) \cdots (n-r+2)(n-r+1) = \frac{n!}{(n-r)!}, \quad (1.1)$$

$$P_n = n!.$$

为使以上两式在 $r = n$ 时统一, 规定 $0! = 1$.

2. 允许重复的排列

有 n 种不同的元素, 允许反复取出同种元素, 这样取出 r 个, 按照一定的顺序排成一列, 称为从 n 个不同元素中取出 r 个元素的允许重复的排列. 从 n 个不同元素中取出 r 个元素的允许重复的排列的种数为 n^r .

1.1.3 组合

由排列的定义可知, 在两个排列中元素不同固然是不同的排列, 元素相同但顺序不同也是不同的排列. 例如 $a_1 a_2 a_3$ 与 $a_1 a_3 a_2$ 是不同的排列. 如果我们考虑的问题, 只与元素有关, 而与顺序无关, 这样的问题称为组合.

定义 1.2 从 n 个不同的元素里, 每次任意取出 $r (1 \leq r \leq n)$ 个不同元素, 不考虑顺序的并成一组, 称为从 n 个不同元素中每次取 r 个不同元素的组合. 组合的种数用 C_n^r 或 $\binom{n}{r}$ 表示.

一般地, 从 n 个不同的元素中每次取 r 个不同元素的排列数为 A_n^r , 从而这时每个含某 r 个元素的组都被计算了 $r!$ 次, 故从 n 个不同的元素中每次任意取 r 个不同元素的组合数为

$$C_n^r = \frac{A_n^r}{r!} = \frac{n(n-1) \cdots (n-r+1)}{r(r-1) \cdots 2 \cdot 1} = \frac{n!}{r! (n-r)!}. \quad (1.2)$$

为使式(1.2)在 $n = r$ 时也成立, 规定 $C_n^0 = 1$.

组合的性质:

$$(1) C_n^r = C_n^{n-r}.$$

这个等式说明,从 n 个不同元素中取出 r 个元素的组合数等于从 n 个不同元素中取出 $n-r$ 个元素的组合数. 所以当 $r > \frac{n}{2}$ 时,通常不直接计算 C_n^r ,而是改为计算 C_n^r .

$$(2) C_n^r = C_{n-1}^{r-1} + C_{n-1}^r, r < n.$$

以上两个性质都可以根据组合的定义得出,这两个性质大大方便了计算.

1.2 概率统计中的一些基本原理

本节将介绍抽样调查中涉及的概率统计中的两个主要定理——大数定律和中心极限定理,以及参数估计的原理.

1.2.1 大数定律

大数定律和中心极限定理是概率论和数理统计中两类重要的定理,抽样调查的原理主要涉及这两个定理. 其中大数定律奠定了用样本来估计总体的理论基础,其直观含义是随机事件的规律性总是在大量观察中才能显现出来,随着观察次数的增加,随机影响将相互抵消而使规律性有稳定的性质. 大数定律总结了这种规律性.

▮ 定理 1.1(切比雪夫(Chebyshev)大数定律) 设随机变量 $X_1, X_2, \dots, X_n, \dots$ 相互独立,且具有相同的数学期望和方差: $E(X_k) = \mu, V(X_k) = \sigma^2 (k = 1, 2, \dots)$, 前 n 个随机变量的算术平均为

$$Y_n = \frac{1}{n} \sum_{k=1}^n X_k,$$

则对于任给正数 ϵ , 有

$$\lim_{n \rightarrow \infty} P\{|Y_n - \mu| < \epsilon\} = \lim_{n \rightarrow \infty} P\left\{\left|\frac{1}{n} \sum_{k=1}^n X_k - \mu\right| < \epsilon\right\} = 1. \quad (1.3)$$

切比雪夫大数定律表明,当 n 很大时,在概率意义下,随机变量 $X_1, X_2, \dots, X_n$ 的算术平均 $\frac{1}{n} \sum_{k=1}^n X_k$ 接近于数学期望 $E(X_k) = \mu$. 也就是说,在定理的条件下, n 个随机变量的算术平均,当 n 无限增加时将几乎变成一个常数.

▮ 定理 1.2(伯努利(Bernoulli)大数定律) 设 n_A 是 n 次独立重复试验

中事件 A 发生的次数, p 是事件 A 在每次试验中发生的概率, 则对于任给的正数 $\varepsilon > 0$, 有

$$\lim_{n \rightarrow \infty} P \left\{ \left| \frac{n_A}{n} - p \right| < \varepsilon \right\} = 1, \quad (1.4)$$

或

$$\lim_{n \rightarrow \infty} P \left\{ \left| \frac{n_A}{n} - p \right| \geq \varepsilon \right\} = 0. \quad (1.5)$$

可见, 伯努利大数定律证明了试验次数 n 无限增大时, 其频率可以无限接近概率.

定理 1.3 (辛钦 (Khinchine) 大数定律) 设随机变量 $X_1, X_2, \dots, X_n, \dots$ 相互独立, 服从同一分布, 且具有数学期望 $E(X_k) = \mu (k=1, 2, \dots)$, 则对于任给的正数 ε , 有

$$\lim_{n \rightarrow \infty} P \left\{ \left| \frac{1}{n} \sum_{k=1}^n X_k - \mu \right| < \varepsilon \right\} = 1. \quad (1.6)$$

显然, 伯努利大数定律是辛钦大数定律的特殊情况. 辛钦大数定律证明了在 n 增大时, 样本均值会无限接近总体的数学期望. 这就给抽样估计提供了理论基础, 即样本比例和样本均值在 n 比较大时可以作为总体比例和总体均值的一个近似估计. 同样, 可以用样本方差来估计总体方差, 具体证明过程将在第 3 章详细介绍.

1.2.2 中心极限定理

中心极限定理奠定了用样本估计量对总体参数进行区间估计的理论基础. 它证明了不论总体服从何种分布, 只要方差有限, 在观察值足够多时, 许多估计量的抽样分布就趋向正态分布. 利用中心极限定理可以对总体参数做区间估计及确定其相应的置信概率, 是抽样结论可靠性的理论依据.

定理 1.4 (列维-林德伯格 (Levy-Lindberg) 中心极限定理) 设随机变量 $X_1, X_2, \dots, X_n, \dots$ 相互独立, 服从同一分布, 且具有数学期望和方差: $E(X_k) = \mu, V(X_k) = \sigma^2 \neq 0 (k=1, 2, \dots)$, 则随机变量

$$\bar{Y}_n = \frac{\sum_{k=1}^n X_k - E\left(\sum_{k=1}^n X_k\right)}{\sqrt{V\left(\sum_{k=1}^n X_k\right)}} = \frac{\sum_{k=1}^n X_k - n\mu}{\sqrt{n}\sigma}$$

的分布函数 $F_n(x)$ 对于任意 x 满足

$$\lim_{n \rightarrow \infty} F_n(x) = \lim_{n \rightarrow \infty} P \left\{ \frac{\sum_{k=1}^n X_k - n\mu}{\sqrt{n}\sigma} \leq x \right\} = \int_{-\infty}^x \frac{1}{\sqrt{2\pi}} e^{-\frac{t^2}{2}} dt. \quad (1.7)$$

中心极限定理对随机变量 $X_1, X_2, \dots, X_n, \dots$ 的分布不作要求, 这为样本容量的确定和在大样本情况下的统计推断提供了理论依据. 定理 1.4 还可以放宽条件, 仅要求 $X_1, X_2, \dots, X_n, \dots$ 是相互独立但不服从同一分布的随机变量序列, 此时, 有一个应用更为广泛的中心极限定理如下:

定理 1.5 (李雅普诺夫 (Liapunov) 定理) 设随机变量 $X_1, X_2, \dots, X_n, \dots$ 相互独立, 它们具有数学期望和方差: $E(X_k) = \mu, V(X_k) = \sigma^2 \neq 0 (k = 1, 2, \dots)$, 记 $B_n^2 = \sum_{k=1}^n \sigma_k^2$. 若存在正数 δ , 使得当 $n \rightarrow \infty$ 时, $\frac{1}{B_n^{2+\delta}} \sum_{k=1}^n E(|X_k - \mu_k|^{2+\delta}) \rightarrow 0$, 则随机变量

$$Z_n = \frac{\sum_{k=1}^n X_k - E\left(\sum_{k=1}^n X_k\right)}{\sqrt{V\left(\sum_{k=1}^n X_k\right)}} = \frac{\sum_{k=1}^n X_k - \sum_{k=1}^n \mu_k}{B_n}$$

的分布函数 $F_n(x)$ 对于任意 x , 满足

$$\lim_{n \rightarrow \infty} F_n(x) = \lim_{n \rightarrow \infty} P \left\{ \frac{\sum_{k=1}^n X_k - \sum_{k=1}^n \mu_k}{B_n} \leq x \right\} = \int_{-\infty}^x \frac{1}{\sqrt{2\pi}} e^{-\frac{t^2}{2}} dt. \quad (1.8)$$

这个定理表明, 在定理的条件下, 随机变量 $Z_n = \frac{\sum_{k=1}^n X_k - \sum_{k=1}^n \mu_k}{B_n}$, 当 n 很大时, 近似地服从正态分布 $N(0, 1)$. 由此, 当 n 很大时, $\sum_{k=1}^n X_k = B_n Z_n + \sum_{k=1}^n \mu_k$ 近似地服从正态分布 $N\left(\sum_{k=1}^n \mu_k, B_n^2\right)$. 这就是说, 无论各个随机变量 $X_k (k=1, 2, \dots)$ 服从什么分布, 只要满足定理的条件, 那么它们的和 $\sum_{k=1}^n X_k$ 当 n 很大时就近似地服从正态分布, 从而均值 $\frac{1}{n} \sum_{k=1}^n X_k$ 当 n 很大时也近似服从正态分布. 中心极限定理是大样本统计推断的理论基础.

1.2.3 参数估计

抽样调查通过对样本的观察来推断总体的某些特征,所依据的是数理统计学中参数估计的原理.这里主要讨论参数估计的两种方式——点估计和区间估计,以及评价估计量的标准.

1. 点估计(point estimation)

设总体 X 的分布函数 $F(x; \theta)$ 的形式已知, θ 是待估参数, $X_1, \dots, X_n$ 是 X 的一个样本, $x_1, x_2, \dots, x_n$ 是相应的一个样本值. 点估计问题就是要构造一个适当的统计量 $\theta(X_1, \dots, X_n)$, 用它的观察值 $\hat{\theta}(x_1, x_2, \dots, x_n)$ 来估计未知参数 θ . 这里称 $\theta(X_1, \dots, X_n)$ 为 θ 的估计量, 称 $\hat{\theta}(x_1, x_2, \dots, x_n)$ 为 θ 的估计值. 在不致混淆的情况下统称估计量和估计值为估计, 并都简记为 $\hat{\theta}$. 由于估计量是样本的函数, 因此对于不同的样本值, θ 的估计值往往是不相同的. 下面介绍两种常用的构造估计量的方法: 矩估计法和极大似然估计法.

(1) 矩估计法

设 X 为连续型随机变量, 其概率密度为 $f(x; \theta_1, \theta_2, \dots, \theta_k)$, 或 X 为离散型随机变量, 其分布率为 $P\{X=x\} = p(x; \theta_1, \theta_2, \dots, \theta_k)$, 其中 $\theta_1, \theta_2, \dots, \theta_k$ 为待估参数, $X_1, \dots, X_n$ 是来自 X 的样本. 假设总体 X 的前 k 阶矩存在, 且为

$$\mu_l = E(X^l) = \int_{-\infty}^{\infty} x^l f(x; \theta_1, \theta_2, \dots, \theta_k) dx \quad (X \text{ 为连续型}),$$

或

$$\mu_l = E(X^l) = \sum_{x \in R_X} x^l p(x; \theta_1, \theta_2, \dots, \theta_k) \quad (X \text{ 为离散型}),$$

$l=1, 2, \dots, k$; R_X 是 x 可能取值的范围. 一般来讲, 这 k 个矩都是 $\theta_1, \theta_2, \dots, \theta_k$ 的函数. 基于样本矩

$$A_l = \frac{1}{n} \sum_{i=1}^n X_i^l$$

依概率收敛于相应的总体矩 μ_l ($l=1, 2, \dots, k$). 若样本矩的连续函数依概率收敛于相应的总体矩的连续函数, 就用样本矩作为相应的总体矩的估计量, 而以样本矩的连续函数作为相应的总体矩的连续函数的估计量, 这种估计方法就称为矩估计法. 矩估计法的具体做法是:

令

$$\mu_l = A_l, \quad l=1, 2, \dots, k,$$

这是一个包含 k 个未知参数 $\theta_1, \theta_2, \dots, \theta_k$ 的联立方程组. 一般来讲, 可以从中解出 $\theta_1, \theta_2, \dots, \theta_k$, 此时, 就用这个方程组的解 $\hat{\theta}_1, \hat{\theta}_2, \dots, \hat{\theta}_k$ 分别作为 $\theta_1, \theta_2, \dots, \theta_k$ 的估计量, 这种估计量称为矩估计量. 矩估计量的观察值称为矩估计值.

(2) 极大似然估计法

若总体 X 属离散型, 其分布率 $P\{X=x\}=p(x;\theta), \theta \in \Theta$ 的形式为已知, θ 为待估参数, Θ 是 θ 可能取值的范围. 设 $X_1, \dots, X_n$ 是来自 X 的样本, 则 $X_1, \dots, X_n$ 的联合分布率为

$$\prod_{i=1}^n p(x_i; \theta).$$

又设 $x_1, x_2, \dots, x_n$ 是相应于样本 $X_1, \dots, X_n$ 的一个样本值, 易知样本 $X_1, \dots, X_n$ 取到观察值 $x_1, x_2, \dots, x_n$ 的概率, 亦即事件 $\{X_1=x_1, X_2=x_2, \dots, X_n=x_n\}$ 发生的概率为

$$L(\theta) = L(x_1, x_2, \dots, x_n; \theta) = \prod_{i=1}^n p(x_i; \theta), \quad \theta \in \Theta. \quad (1.9)$$

这一概率随 θ 的取值而变化, 它是 θ 的函数. $L(\theta)$ 称为样本的似然函数.

由费希尔 (Fisher) 引进的极大似然估计法, 就是固定样本观察值 $x_1, x_2, \dots, x_n$, 在 θ 取值的可能范围 Θ 内挑选使概率 $L(x_1, x_2, \dots, x_n; \theta)$ 达到最大的参数值 $\hat{\theta}$, 作为参数 θ 的估计值, 即取 $\hat{\theta}$ 使

$$L(x_1, x_2, \dots, x_n; \hat{\theta}) = \max_{\theta \in \Theta} L(x_1, x_2, \dots, x_n; \theta). \quad (1.10)$$

这样得到的 $\hat{\theta}$ 与样本值 $x_1, x_2, \dots, x_n$ 有关, 常记为 $\hat{\theta}(x_1, x_2, \dots, x_n)$, 称为参数 θ 的极大似然估计值, 而相应的统计量 $\hat{\theta}(X_1, \dots, X_n)$ 称为参数 θ 的极大似然估计量.

若总体 X 属连续型, 其概率密度 $f(x; \theta), \theta \in \Theta$ 的形式已知, θ 为待估参数, Θ 是 θ 可能取值的范围. 设 $X_1, \dots, X_n$ 是来自 X 的样本, 则 $X_1, \dots, X_n$ 的联合密度为

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

又设 $x_1, x_2, \dots, x_n$ 是相应于样本 $X_1, \dots, X_n$ 的一个样本值, 则随机点 $(X_1, \dots, X_n)$ 落在点 $(x_1, x_2, \dots, x_n)$ 的邻域 (边长分别为 $dx_1, dx_2, \dots, dx_n$ 的 n 维立方体) 内的概率近似为

$$\prod_{i=1}^n f(x_i; \theta) dx_i, \quad (1.11)$$

其值随 θ 的取值而变化. 与离散型的情况一样, 取 θ 的估计值 $\hat{\theta}$ 使概率 (1.11)

取到最大值, 但因子 $\prod_{i=1}^n dx_i$ 不随 θ 而变, 故只需考虑函数

$$L(\theta) = L(x_1, x_2, \dots, x_n; \theta) = \prod_{i=1}^n f(x_i; \theta) \quad (1.12)$$

的最大值. 这里 $L(\theta)$ 称为样本的似然函数. 若

$$L(x_1, x_2, \dots, x_n; \hat{\theta}) = \max_{\theta \in \Theta} L(x_1, x_2, \dots, x_n; \theta),$$

则称 $\hat{\theta}(x_1, x_2, \dots, x_n)$ 为 θ 的极大似然估计值, 称 $\hat{\theta}(X_1, \dots, X_n)$ 为 θ 的极大似然估计量.

在很多情形, $p(x; \theta)$ 和 $f(x; \theta)$ 关于 θ 可微, 这时 θ 常可从方程

$$\frac{d}{d\theta} L(\theta) = 0 \quad (1.13)$$

解得. 又因 $L(\theta)$ 与 $\ln L(\theta)$ 在同一 θ 处取到极值, 因此, θ 的极大似然估计 θ 可以从方程

$$\frac{d}{d\theta} \ln L(\theta) = 0 \quad (1.14)$$

求得, 这一求解方法往往比较方便.

2. 估计量的评选标准

参数估计中估计量的选择标准也适用于抽样调查, 只是由于在抽样调查中有各种不同的抽选方法, 不同的抽选方法有不同的估计量, 使得在抽样调查中具有更多可供选择的估计量. 一个好的估计量的标准主要有无偏性、有效性和一致性.

(1) 无偏性

若估计量 $\hat{\theta}(X_1, \dots, X_n)$ 的数学期望 $E(\hat{\theta})$ 存在, 且对于任意 $\theta \in \Theta$ 有 $E(\hat{\theta}) = \theta$, 则称 $\hat{\theta}$ 是 θ 的无偏估计量. 从直观意义上讲, 用个别样本来对总体参数进行估计时, 由于随机的原因, 可能会偏离真值, 然而一个好的估计量从平均的意义上应该等于所要估计的总体参数值.

(2) 有效性

有效性是用来比较参数 θ 的两个无偏估计量 $\hat{\theta}_1$ 和 $\hat{\theta}_2$ 的, 如果在样本容量 n 相同的情况下, $\hat{\theta}_1$ 的观察值较 $\hat{\theta}_2$ 更密集在真值 θ 的附近, 我们就认为 $\hat{\theta}_1$ 较 $\hat{\theta}_2$ 为理想. 由于方差是随机变量取值与其数学期望的偏离程度的度量, 所以无偏估计以方差小者为好. 因此, 有效性的定义为, 设 $\hat{\theta}_1$ 和 $\hat{\theta}_2$ 都是 θ 的无偏估计

量,若有 $V(\hat{\theta}_1) < V(\hat{\theta}_2)$, 则称 $\hat{\theta}_1$ 较 $\hat{\theta}_2$ 有效.

(3) 一致性

一致性又称相合性. 前面讲的无偏性和有效性都是在样本容量 n 固定的前提下提出的, 我们自然希望随着样本容量的增大, 一个估计量的值稳定于待估参数的真值, 因此一致性定义为, 设 $\hat{\theta}(X_1, \dots, X_n)$ 为参数 θ 的估计量, 若对于任意 $\theta \in \Theta$, 当 $n \rightarrow \infty$ 时, $\hat{\theta}(X_1, \dots, X_n)$ 依概率收敛于 θ , 则称 $\hat{\theta}$ 为 θ 的一致估计量. 因此一致性是在无限总体的情况下进行抽样, 是从极限意义上对统计量与总体参数的关系来说的. 在抽样调查中, 总体 N 通常是有限总体, 采用不重复抽样, 显然当 $n \rightarrow N$ 时, 样本的估计值就等于总体的参数, 在这个意义上一致性也是成立的.

3. 区间估计(interval estimation)

对于一个未知量, 人们在测量或计算时, 常不以得到近似值为满足, 还需估计误差, 即要求确切地知道近似值的精确程度 (亦即所求真值所在的范围). 类似地, 对于未知参数 θ , 除了求出它的点估计 $\hat{\theta}$ 外, 还希望估计出一个范围, 并希望知道这个范围包含参数 θ 真值的可信程度. 这样的范围通常以区间的形式给出, 同时还给出此区间包含参数 θ 真值的可信程度. 这种形式的估计称为区间估计, 这样的区间即所谓置信区间, 下面引入置信区间的定义.

置信区间 设总体 X 的分布函数 $F(X; \theta)$ 含有一个未知参数 θ , 对于给定值 $\alpha (0 < \alpha < 1)$, 若由样本 $X_1, \dots, X_n$ 确定的两个统计量 $\underline{\theta} = \underline{\theta}(X_1, \dots, X_n)$ 和 $\bar{\theta} = \bar{\theta}(X_1, \dots, X_n)$ 满足

$$P\{\underline{\theta}(X_1, \dots, X_n) < \theta < \bar{\theta}(X_1, \dots, X_n)\} = 1 - \alpha, \quad (1.15)$$

则称随机区间 $(\underline{\theta}, \bar{\theta})$ 是 θ 的置信度为 $1 - \alpha$ 的置信区间, $\underline{\theta}$ 和 $\bar{\theta}$ 分别称为置信度为 $1 - \alpha$ 的双侧置信区间的置信下限和置信上限, $1 - \alpha$ 称为置信度.

这就是说, 若反复抽样多次 (各次得到的样本的容量相等, 都是 n), 每个样本值确定一个区间 $(\underline{\theta}, \bar{\theta})$, 每个这样的区间要么包含 θ 的真值, 要么不包含 θ 的真值. 按照伯努利大数定律, 在这样多的区间中, 包含 θ 真值的约占 $100(1 - \alpha)\%$, 不包含 θ 真值的约仅占 $100\alpha\%$. 例如, 若 $\alpha = 0.01$, 反复抽样 1000 次, 则得到的 1000 个区间中不包含 θ 真值的约为 10 个.

【例 1.1】 总体 $X \sim N(\mu, \sigma^2)$, σ^2 为已知, μ 为未知, 设 $X_1, \dots, X_n$ 是来自 X 的样本, 求 μ 的置信度为 $1 - \alpha$ 的置信区间.

解 因为 $\bar{X}$ 是 μ 的无偏估计, 且有

$$\frac{\bar{X}-\mu}{\sigma/\sqrt{n}} \sim N(0,1).$$

$\frac{\bar{X}-\mu}{\sigma/\sqrt{n}}$ 所服从的分布 $N(0,1)$ 是不依赖于任何未知参数的, 按标准正态分

布的上 α 分位点的定义, 有

$$P\left\{\left|\frac{\bar{X}-\mu}{\sigma/\sqrt{n}}\right| < u_{\alpha/2}\right\} = 1-\alpha,$$

即

$$P\left\{\bar{X}-\frac{\sigma}{\sqrt{n}} u_{\alpha/2} < \mu < \bar{X}+\frac{\sigma}{\sqrt{n}} u_{\alpha/2}\right\} = 1-\alpha.$$

这样, 就得到了 μ 的一个置信度为 $1-\alpha$ 的置信区间

$$\left(\bar{X}-\frac{\sigma}{\sqrt{n}} u_{\alpha/2}, \bar{X}+\frac{\sigma}{\sqrt{n}} u_{\alpha/2}\right).$$

如果取 $\alpha=0.05$, 查表知 $u_{\alpha/2}=u_{0.025}=1.96$, 又若 $\sigma=1, n=16$, 于是得到一个置信度为 0.95 的置信区间

$$\left(\bar{X} \pm \frac{1}{\sqrt{16}} \times 1.96\right), \quad \text{即} \quad (\bar{X} \pm 0.49).$$

再者, 若由一个样本值算得样本均值的观察值 $\bar{x}=5.20$, 则得到一个区间

$$(5.20 \pm 0.49), \quad \text{即} \quad (4.71, 5.69).$$

通过这个例子, 可以看到求未知参数 θ 的置信区间的具体做法如下:

(1) 寻求一个样本 $X_1, \dots, X_n$ 的函数

$$Z=Z(X_1, \dots, X_n; \theta),$$

它包含待估参数 θ , 而不含其他未知参数, 并且 Z 的分布已知且不依赖于任何未知参数(当然不依赖于待估参数 θ);

(2) 对于给定的置信度 $1-\alpha$, 定出两个常数 a, b , 使

$$P\{a < Z(X_1, \dots, X_n; \theta) < b\} = 1-\alpha;$$

(3) 若能从 $a < Z(X_1, \dots, X_n; \theta) < b$ 得到等价的不等式 $\underline{\theta} < \theta < \bar{\theta}$, 其中 $\underline{\theta} = \underline{\theta}(X_1, \dots, X_n)$, $\bar{\theta} = \bar{\theta}(X_1, \dots, X_n)$ 都是统计量, 那么 $(\underline{\theta}, \bar{\theta})$ 就是 θ 的一个置信度为 $1-\alpha$ 的置信区间.

函数 $Z(X_1, \dots, X_n; \theta)$ 的构造, 通常可以从 θ 的点估计着手考虑, 许多常用的正态总体参数的置信区间可以用上述步骤推得. 读者若想详细了解具体方法, 可以参看概率统计书籍的相关章节.

抽样估计是参数估计的应用和进一步推广, 但它与传统数理统计学中的