

工科研究生教材

应用数理统计

YING YONG SHU LI TONG JI

朱勇华 邵淑彩 孙韞玉 编

武汉水利电力大学出版社

工科研究生教材

应用数理统计

朱勇华 邵淑彩 孙韞玉 编

武汉水利电力大学出版社

1999年·武汉

(鄂) 新登字 15 号

图书在版编目 (CIP) 数据

应用数理统计 / 朱勇华等编. — 武汉: 武汉水利电力大学出版社, 1999. 11

ISBN 7-81063-059-8

I. 应… II. 朱… III. 数理统计 IV. 0212

中国版本图书馆 CIP 数据核字 (1999) 第 64277 号

责任编辑: 李汉保 责任校对: 徐虹 封面设计: 涂驰

武汉水利电力大学出版社出版发行

(武汉市武昌东湖南路 8 号, 邮编 430072)

湖北省国营人民大垸印刷厂印刷

开本: 850×1168 1/32 印张: 14.5 字数: 361 千字

2000 年 5 月第 1 版 2000 年 5 月第 1 次印刷 印数: 1-4 500 册

ISBN7-81063-059-8/O-4 定价: 18.00 元

前 言

数理统计一直是应用数学中最重要、最活跃的学科之一，对它的应用越来越广泛深入，它在国民经济和科学技术中的作用越来越重要。作为工科研究生，理应具备数理统计的基础知识和数理统计的思想方法。为此，我们根据多年的教学实践经验，结合工科研究生的特点，编写了这本数理统计教材，希望对工科研究生学习数理统计的基本知识，树立正确的统计观点能有所助益。

编写这本书的主要指导思想有以下三点：

第一，注重思想方法的介绍。我们认为，数理统计不仅是一门数学理论，而且还是一门“看世界”的方法。具备正确的数理统计的思想方法是工科研究生观察现实世界所必备的文化修养。为此，本书特别注重介绍各种统计方法的统计思想、问题的背景、统计方法产生的历史和不同统计学派的特点，使学生能够对数理统计的思想方法有一个系统的、全面的了解。

第二，注重应用性。数理统计是一门应用性很强的应用数学学科，它的应用几乎遍及社会科学、自然科学和工程技术的各个领域，成为解决国民经济和科学技术中实际问题的重要工具。作为一本好的数理统计教材，所介绍的理论和方法应该具有较强的应用性。因此，本书除了突出各种统计方法的应用性外，还编入了应用面较广的多元统计分析和贝叶斯估计方法，充实了一些应用性较强的内容，如估计量的评选标准的补充、样本容量的确定、

非参数检验的若干方法等,并介绍了目前国际上流行的统计软件,以适应工科研究生解决实际问题的需要。

第三,突出数理统计这门学科的特点。我们知道,统计推断是由样本去推断总体,由部分去推断整体,这就不可避免犯错误,数理统计的任务就是要充分利用样本的信息,作出尽可能精确、尽可能可靠的推断。书中通过方法的介绍、实例的分析来突出这个特点,让学生知道,数理统计不是一门纯技术性的学科,不是学会使用现成的方法就行了。而是要学生学会如何正确地使用统计方法,如何理解统计分析的结果。

全书共分六章,第一章、第六章由孙韞玉编写,第四章、第五章由邵淑彩编写,其余由朱勇华编写,朱勇华对全书进行了统稿。彭祖赠教授逐字逐句地审阅了全部手稿,提出了许多修改意见,周学良教授也审阅了部分手稿,提出了宝贵的意见。编者根据两位教授的意见,又对全书进行了全面的修改和补充。但限于编者的水平,仍然会有不少缺点和错误,敬请专家和读者批评指正。

本书的出版,得到了武汉水利电力大学研究生部、数学物理系的大力支持和鼓励,得到了武汉水利电力大学出版社的通力协作。在此向所有协助本书出版的同志致以衷心的感谢。

编者

一九九九年四月

目 录

前 言	i
第一章 数理统计基本概念与抽样分布	1
§1.1 引言	1
§1.2 数理统计中的几个基本概念	3
一、总体与样本 (3)	
二、统计量 (6)	
三、理论分布与经验分布 (8)	
§1.3 常用统计分布	11
一、 χ^2 分布 (12)	
二、 t 分布 (15)	
三、 F 分布 (16)	
四、概率分布的分位数 (17)	
§1.4 正态总体的抽样分布	21
§1.5 非正态总体的一些抽样分布	28
习题一	32
第二章 参数估计	35
§2.1 求点估计的方法	35
一、矩估计法 (36)	
二、极大似然估计法 (39)	
§2.2 估计量的评选标准	46
一、无偏性 (48)	
二、有效性 (51)	
三、相合性 (59)	
四、其它准则 (60)	
§2.3 区间估计	62

一、区间估计的基本思想 (62)	
二、单个正态总体参数的区间估计 (65)	
三、两个正态总体参数的区间估计 (68)	
四、非正态总体参数的区间估计 (76)	
§2.4 贝叶斯 (Bayes) 估计.....	79
一、先验分布与后验分布 (80)	二、贝叶斯估计 (81)
三、先验分布的选取 (85)	
习题二.....	86
第三章 假设检验	91
§3.1 假设检验的基本概念.....	91
一、假设检验问题 (91)	二、假设检验的基本原理 (92)
三、两类错误 (97)	四、假设检验的一般步骤 (100)
§3.2 单个正态总体参数的假设检验.....	101
一、当 σ^2 已知时, 关于 μ 的检验 (101)	
二、当 σ^2 未知时, 关于 μ 的检验 (103)	
三、当 μ 未知时, 关于 σ^2 的检验 (105)	
§3.3 两个正态总体参数的假设检验.....	107
一、当 σ_1^2, σ_2^2 已知时, 关于 μ_1, μ_2 的检验 (107)	
二、当 $\sigma_1^2 = \sigma_2^2 = \sigma^2$ 未知时, 关于 μ_1, μ_2 的检验 (108)	
三、当 σ_1^2, σ_2^2 未知且不相等时, 关于 μ_1, μ_2 的检验 (111)	
四、当 μ_1, μ_2 未知时, 关于方差 σ_1^2, σ_2^2 的检验 (114)	
§3.4 非正态总体参数的假设检验.....	116
一、大样本方法 (116)	二、特例 (119)
§3.5 样本容量 n 的确定.....	121
一、方差已知时, 正态总体均值的右边检验 (122)	
二、方差未知时, 正态总体均值的右边检验 (123)	
三、 μ 未知时, 正态总体方差的右边检验 (124)	

§3.6 非参数假设检验·····	126
一、皮尔逊 (K. Pearson) χ^2 拟合检验法 (126)	
二、 K 检验法 (135) 三、正态性检验 (138) 四、独立性检验 (142)	
习题三·····	146
第四章 回归分析·····	151
§4.1 一元线性回归·····	152
一、一元线性回归 (152) 二、未知参数的估计及统计性质 (154)	
三、回归效果的显著性检验 (161) 四、回归系数的置信区间 (170)	
五、预测与控制 (172)	
§4.2 多元线性回归·····	178
一、多元线性回归模型 (178) 二、未知参数的估计 (180)	
三、回归效果的显著性检验 (188) 四、单个回归系数的显著性检验，剔除变量的计算 (192) 五、预测 (197) 六、最优回归方程的选择 (199)	
§4.3 可化为线性回归的曲线回归·····	201
一、作变量替换化曲线回归为一元线性回归 (201)	
二、多项式回归 (207)	
习题四·····	210
第五章 方差分析及正交试验设计·····	215
§5.1 单因素试验方差分析·····	215
一、数学模型 (217) 二、统计分析 (219)	
§5.2 双因素试验方差分析·····	232
一、双因素有交互作用的方差分析的数学模型 (233)	
二、统计分析 (236) 三、双因素无交互作用 ($t=1$) 的方差分析 (246)	
§5.3 正交设计的基本方法·····	251
一、正交表 (252) 二、正交试验方案及其合理性解释 (253)	
三、正交设计的直观分析 (256)	

四、有交互作用的正交设计及其结果的直观分析 (264)	
§5.4 正交设计的方差分析.....	267
一、正交设计的方差分析 (268)	二、最优生产条件 (274)
三、带重复试验的方差分析 (284)	
习题五.....	287
第六章 多元统计分析.....	293
§6.1 多元正态分布及其参数的估计与检验.....	294
一、多元正态分布 (294)	二、均值向量与协方差阵的估计 (298)
三、均值向量与协方差阵的检验 (303)	
§6.2 相关分析.....	310
一、主成份分析 (311)	二、典型相关分析 (323)
§6.3 判别分析.....	337
一、距离判别 (338)	二、贝叶斯判别 (347)
三、费歇尔判别 (356)	
§6.4 聚类分析.....	361
一、聚类统计量 (362)	二、系统聚类法 (366)
三、动态聚类法 (374)	四、模糊聚类法 (378)
习题六.....	384
习题答案.....	395
附录 1: 常用统计分析软件简介.....	403
附录 2: 常用数理统计表.....	407
参考文献.....	453

第一章 数理统计基本概念 与抽样分布

§1.1 引言

数理统计学是数学的一个分支。它的任务是研究怎样用有效的方法去收集和使用带随机性影响的数据。

数理统计学研究的对象是带随机性影响的数据。是否假定数据有随机性，这是区别数理统计方法和其它数据处理方法的根本点。数据的随机性来源有二：一是抽样的随机性，由于经济原因的考虑或时间限制或问题性质决定，不可能或者没有必要得到研究对象的全部资料，而只能用“一定的方式”抽取其一部分进行考察，这样所得数据的随机性就是来自抽样的随机性，二是试验过程中的随机误差，即在试验过程中未加控制，或无法控制，或不便控制，甚至是不了解的因素所引起的误差。在实际问题中这两类随机性常常交织在一起。例如某工厂生产出大量的电视机显象管，为了检测显象管的寿命，推断寿命的分布类型、有关参数的具体数值以及是否达到生产要求等等，必须对显象管的寿命进行测试，由于寿命试验具有破坏性，所以只能抽取少量显象管以一定方式进行加速老化试验而得到部分数据，这里，抽样的随机性对数据便有影响。另外产品即使是在同一条件下生产出来的，

但各台显象管的寿命仍会有差异,这就是随机误差对数据有影响。

数理统计学研究的内容随着科学技术和生产实际的不断发展而逐步扩大,概括起来可分为两大类:(1)用有效的方法去收集数据。这里“有效”一词有两方面的含义,一是可以建立一个在数学上便于处理的模型来描述所得数据,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。对此问题的研究构成了数理统计学中的两个分支,即抽样理论和试验设计,这些不是本书的主要内容。(2)有效地使用数据。获取数据以后,必须使用有效的方法去集中和提取数据中的有关信息,以对所研究的问题作出尽可能精确和可靠的结论,这种“结论”在统计上叫做“推断”。有效地使用数据是比有效地收集数据更为复杂的问题,这一问题的研究构成了数理统计学的中心内容——统计推断。上面提到的推断显象管寿命的分布类型、有关参数的具体数值以及是否达到生产要求等都是统计推断所要解决的问题。本书将主要讨论统计推断。数理统计学中除了上面所提到的一些分支外,还有不少内容,如质量控制,可靠性理论,统计决策理论,时间序列分析等等,限于篇幅,本书未作介绍。

需要强调的是由于收集和使用的仅仅是部分数据,且带有随机性,要利用部分去推断总体时,其结论只能做到尽可能而非绝对的精确和可靠,而结论的正确性程度显然可以用概率来度量,因此概率论是数理统计的基础,这是毫无疑义的,不过统计方法的具体使用并不需要很高深的数学知识,但这些方法的理论依据,不具备较多较深的数学知识就说不清楚,本书着重介绍数理统计方法,也给出一些必要的数学推导,但不追求其严密性及完整性。

由于随机性影响无所不在,所以数理统计方法的应用十分广泛,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。例如,在工农业生产中最佳生产工艺的安排,产品质量的控制管理,技术革新前后产品质量的鉴定,元器件寿命的计算,工程设

计中安全系数的统计分析等等；在医学卫生领域中药物疗效的检验，某种疾病与某特定因素的关系大小的确定，大气污染的诸有害成分的主成分分析等等；在社会、经济领域中的抽样调查、民意测验、市场预测等等都要用到数理统计方法。此外，在测量、通信、气象、水文、地震预报、地质探矿、考古研究、刑事鉴别等各方面数理统计方法都得到愈来愈多的应用。

数理统计学是一门较年轻的学科。它正式诞生于 19 世纪后期，至 20 世纪 20 年代这门学科已稳稳地站住了脚跟，到 40 年代已形成一个成熟的数学分支。二次大战后工农业和科技等方面迅速发展，对数理统计不断提出新的课题，使其不断地向纵深发展。特别是由于电子计算机具有大批量、高速度处理数据的能力，使得数理统计方法得到广泛的应用，各种使用方便的统计程序包的出现使得数理统计发挥着越来越大的作用。

§1.2 数理统计中的几个基本概念

一、总体与样本

直观地说，把研究对象的全体称为总体（又称母体），而组成总体的每个元素称为个体。总体包含的个体数可以是有限的，也可以是无限制的。对每个个体来说，它有各个方面的特性，而人们关心的往往不是个体的一切方面，只是它的某个（或某几个）数量指标以及该指标在总体中的概率分布情况。例如，在研究一批电视机组成的总体时，可能关心的是电视机显象管寿命的概率分布情况，由于任何一台显象管的寿命事先是不能确定的，而每一台显象管都确实对应有一个寿命值，所以可认为显象管寿命是一个随机变量，也就是说，把总体与一个随机变量（显象管寿命）

联系起来，对总体的研究也就转化为对表示总体的随机变量的统计规律的研究。这样，就可以用精确的语言来描述总体与个体。总体就是一个具有确定概率分布的随机变量（一维或多维），而一个个体则是随机变量的一次观测值。以后常用大写字母 X 、 Y 等表示总体。

总体 X 的概率分布是确定的，但又是未知的，或至少分布的某些参数是未知的。为了研究总体的情况，必须在总体中抽取一定数量的个体进行观测，这个过程称为抽样（也称取样、采样）。从一个总体 X 中抽取 n 个个体为观测值 $(x_1, x_2, \dots, x_n)$ ，这样取得的 $(x_1, x_2, \dots, x_n)$ 称为取自总体 X 的一个样本（又称子样）。样本中个体的数目 n 称为样本容量。注意到数理统计学的对象是受到随机性影响的数据，对于某一次抽样来说，所得观测数据 $(x_1, x_2, \dots, x_n)$ 是完全确定的一组数，但由于抽样的随机性，观测值 $(x_1, x_2, \dots, x_n)$ 也是随机变化的，为了强调随机性，用 $(X_1, X_2, \dots, X_n)$ 表示样本，也就是说，应从两个角度来看待样本，在抽样之前，样本是一个（ n 维）随机变量，在进行一次具体的抽样之后它就是一个数组，这就是样本的二重性。认识样本的二重性是十分重要的，一般在理论推导中总把样本视为随机变量，而在用理论推导所得出的结论进行具体推断时，样本就成了具体数字了。把样本 $(X_1, X_2, \dots, X_n)$ 可能取值的全体称为样本空间，记为 $\mathcal{X}$ ，它通常为 n 维空间或其中的某一个子集，一个具体样本值 $(x_1, x_2, \dots, x_n)$ 则是样本空间中的一个点。

为了使样本能很好地反映总体，除了要求抽样随机性以外，较为自然也最有实用价值的要求是：（1）独立性。因为独立观测是一种最简单的观测方法，对所得数据进行处理也较方便，所以要求在抽取 n 个个体时，每一次抽样都是独立进行的，即各次抽取的结果彼此互不影响，用概率论语言说，即要求 $X_1, X_2, \dots, X_n$ 是相互独立的随机变量。（2）代表性。要求抽取 n 个个体时每一

次抽样都是在完全相同条件下进行，这样就能保证每一个 X_i ($i = 1, 2, \dots, n$)都具有总体的特征，即要求每一个 X_i ($i = 1, 2, \dots, n$)都与总体有相同的分布。上述两点要求在多数情况下是容易得到满足的。凡满足这两个要求所得样本称为简单随机样本。为明确起见下面用定义形式表述出来。

定义 1.2.1 设随机变量 $X_1, X_2, \dots, X_n$ 相互独立，且每一个 X_i ($i = 1, 2, \dots, n$)与总体 X 有相同的概率分布，则称 $(X_1, X_2, \dots, X_n)$ 为来自总体 X 的容量为 n 的简单随机样本。

今后如无特别声明，凡提到样本，都是简单随机样本。

引入简单随机样本，就可利用概率论中对独立同分布的随机变量序列所建立的许多重要定理，这些重要定理的结论为数理统计提供了必要的理论基础。下面的定理以后会多次用到，它由多维随机变量的分布性质推出。

定理 1.2.1 若 $(X_1, X_2, \dots, X_n)$ 是来自总体 X 的样本，设 X 的分布函数为 $F(x)$ ，则样本 $(X_1, X_2, \dots, X_n)$ 的联合分布函数为 $\prod_{i=1}^n F(x_i)$ 。

例 1.2.1 设 $(X_1, X_2, \dots, X_n)$ 是取自总体 X 的样本， X 服从参数为 λ 的指数分布，则 $(X_1, X_2, \dots, X_n)$ 的联合分布函数为

$$F(x_1, x_2, \dots, x_n) = \begin{cases} \prod_{i=1}^n (1 - e^{-\lambda x_i}), & \forall x_i > 0 \\ 0 & \text{其他} \end{cases}$$

联合概率密度函数为

$$f(x_1, x_2, \dots, x_n) = \begin{cases} \prod_{i=1}^n \lambda e^{-\lambda x_i}, & \forall x_i > 0 \\ 0 & \text{其他} \end{cases}$$

二、统计量

样本是对总体进行统计分析和推断的依据, 虽然样本含有总体的信息, 但比较分散, 必须经过一定的加工、提炼, 把分散在样本中有用的信息集中起来。具体地说, 就是针对不同问题构造样本的各种函数, 再利用这些函数去推断总体的性质, 在数理统计学中称这种函数为统计量。

定义 1.2.2 设 $(X_1, X_2, \dots, X_n)$ 为取自总体 X 的一个样本, $T(X_1, X_2, \dots, X_n)$ 为 $(X_1, X_2, \dots, X_n)$ 的一个实值函数, 且 T 中不包含任何未知参数, 则称 $T = T(X_1, X_2, \dots, X_n)$ 为一个统计量。

作为统计量必须不含任何未知参数, 这一点是非常重要的, 因为在有些情形, 统计量 T 是作为未知参数 θ 的估计量而构造的, 若 T 中含有未知参数 θ , 就无法作为 θ 的估计了。注意到样本的二重性, 作为样本的函数的统计量也就具有二重性, 即统计量 $T(X_1, X_2, \dots, X_n)$ 为随机变量, 它应有确定的概率分布, 称之为抽样分布, 而对于样本的一个观测值 $(x_1, x_2, \dots, x_n)$, 统计量 $T(X_1, X_2, \dots, X_n)$ 也有一个相应的值 $T(x_1, x_2, \dots, x_n)$ 。

下面介绍几个常用的重要统计量。

定义 1.2.3 设 $(X_1, X_2, \dots, X_n)$ 是从总体 X 中抽取的一个样本, 下列统计量分别被称为

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

$$\text{样本方差} \quad S^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2 = \frac{1}{n-1} [\sum_{i=1}^n X_i^2 - n\bar{X}^2]$$

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

$$\text{样本 } k \text{ 阶原点矩} \quad M_k = \frac{1}{n} \sum_{i=1}^n X_i^k, \quad k = 1, 2, \dots$$

$$\text{样本 } k \text{ 阶中心矩} \quad M'_k = \frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^k, \quad k = 2, 3, \dots$$

这些统计量统称为总体的样本矩。

显然 $M_1 = \bar{X}$, 它是样本的算术平均, $M'_2 = \frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^2$, 本书中常将 M'_2 用 $\bar{S}^2$ 表示, $\bar{S}^2$ 与 S^2 略有不同, 但它们都是样本平均偏差平方和。 $\bar{X}$ 和 S^2 是以后用得最多的统计量, 由下面定理可以看出, $\bar{X}$ 集中反映了总体均值的信息, S^2 集中反映了总体方差的信息。

定理 1.2.2 设 $(X_1, X_2, \dots, X_n)$ 是取自总体 X 的一个样本, 若 X 的二阶矩存在, 并记 $E(X) = \mu$, $D(X) = \sigma^2$, 则有

$$E(\bar{X}) = \mu, \quad D(\bar{X}) = \frac{\sigma^2}{n}, \quad E(S^2) = \sigma^2$$

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

$$D(\bar{X}) = D\left(\frac{1}{n} \sum_{i=1}^n X_i\right) = \frac{1}{n^2} \sum_{i=1}^n D(X_i) = \frac{1}{n^2} \sum_{i=1}^n D(X) = \frac{\sigma^2}{n}$$

$$E(S^2) = E\left[\frac{1}{n-1} (\sum_{i=1}^n X_i^2 - n\bar{X}^2)\right] = \frac{1}{n-1} \left[\sum_{i=1}^n E(X_i^2) - nE(\bar{X}^2)\right]$$

$$= \frac{1}{n-1} \left[\sum_{i=1}^n (D(X_i) + (E(X_i))^2) - n(D(\bar{X}) + (E\bar{X})^2)\right]$$

$$= \frac{1}{n-1} \left[n(\sigma^2 + \mu^2) - n\left(\frac{\sigma^2}{n} + \mu^2\right)\right] = \sigma^2$$

还可以推出 $D(S^2)$ 与总体矩的关系式, 因推导较烦琐, 故而略去。

需要指出, 若总体 X 的 k 阶矩存在, 则样本的 k 阶矩必依概率收敛于总体的 k 阶矩。例如, $M_k = \frac{1}{n} \sum_{i=1}^n X_i^k$ 为样本 k 阶原点矩, $\mu_k = E(X^k)$ 为总体 k 阶原点矩, 因为 $X_1, X_2, \dots, X_n$ 相互独立且与 X 同分布, 所以 $X_1^k, X_2^k, \dots, X_n^k$ 相互独立且与 X^k 同分布, 再注意到

$$E(M_k) = E\left(\frac{1}{n} \sum_{i=1}^n X_i^k\right) = \frac{1}{n} \sum_{i=1}^n E(X_i^k) = \frac{1}{n} \sum_{i=1}^n E(X^k) = \mu_k$$

故由独立同分布的辛钦 (Хинчин) 大数定律可知, 当 $n \rightarrow \infty$ 时, M_k 依概率收敛于 μ_k 。

定义 1.2.4 设 $(X_1, X_2, \dots, X_n)$ 是取自总体 X 的一个样本, 将样本观测值 $(x_1, x_2, \dots, x_n)$ 的各分量按大小递增顺序排序。

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

当 $(X_1, X_2, \dots, X_n)$ 取值为 $(x_1, x_2, \dots, x_n)$ 时, 定义 $(X_{(1)}, X_{(2)}, \dots, X_{(n)})$ 取值为 $(x_{(1)}, x_{(2)}, \dots, x_{(n)})$, 则称 $(X_{(1)}, X_{(2)}, \dots, X_{(n)})$ 为由 $(X_1, X_2, \dots, X_n)$ 导出的一组顺序统计量, 称 $X_{(k)}$ 为第 k 个顺序统计量, 特别地分别称

$$X_{(1)} = \min_{1 \leq i \leq n} \{X_i\}, \quad X_{(n)} = \max_{1 \leq i \leq n} \{X_i\}$$

为最小顺序统计量和最大顺序统计量, 即 $X_{(1)}$ 的取值为样本观测值中最小的一个, $X_{(n)}$ 的取值为样本观测值中最大的一个。 $X_{(1)}, X_{(n)}$ 统称为极值统计量。

后面将会看到上述统计量在统计推断中的作用。

三、理论分布与经验分布

总体 X 的分布函数称为理论分布函数, 样本的分布函数则称为经验分布函数, 其具体定义如下。

定义 1.2.5 设 $(X_1, X_2, \dots, X_n)$ 为取自总体 X 的一个样本, $(x_1, x_2, \dots, x_n)$ 是样本的一个观测值, 将这些值按大小递增顺序排列成 $x_{(1)} \leq x_{(2)} \leq \dots \leq x_{(n)}$, 并作函数