

工程硕士研究生教材

工程数学

下册 (数理统计与随机过程)

同济大学应用数学系 编著

同济大学出版社

7B11
7566

工程硕士研究生教材

工 程 数 学

(下册)

(数理统计与随机过程)

同济大学应用数学系 编著

同济大学出版社

图书在版编目(CIP)数据

工程数学(下册)/同济大学应用数学系编著. —上海:同济大学出版社, 2002. 8

工程硕士研究生教材

ISBN 7-5608-2443-9

I. 工… II. 同… III. 工程数学—研究生—教材

IV. TB11

中国版本图书馆 CIP 数据核字(2002)第 036289 号

工程硕士研究生教材

工程数学(下册)(数理统计与随机过程)

作者 同济大学应用数学系 编著

责任编辑 李炳钊 责任校对 郁 峰 装帧设计 李志云

出 版 同济大学出版社
发 行

(上海四平路 1239 号 邮编 200092 电话 021-65985622)

经 销 全国各地新华书店

印 刷 常熟市华顺印刷有限公司

开 本 787mm×960mm 1/16

印 张 21.25

字 数 425000

印 数 1—5000

定 价 28.00 元

版 次 2002 年 8 月第 1 版 2002 年 8 月第 1 次印刷

书 号 ISBN 7-5608-2443-9/O · 212

本书若有印装质量问题, 请向本社发行部调换

前 言

培养工程硕士研究生是为适应我国经济建设中对应用型、复合型高层次工程技术和工程管理人才的需要所采取的一项重要举措。“工程数学”课程是工程硕士研究生培养中一门重要的基础课,它适应不同专业、不同学习内容的要求,以及在较少的学时内掌握其所学专业必须具备的数学基础这一实际情况,编写一本可以根据各专业实际情况进行教学的教材是十分必要的。

我们自 1998 年开始为工程硕士研究生讲授工程数学课程,针对实际情况编写了《工程数学》讲义,该讲义已经在工程硕士研究生工程数学课程的教学中多次使用。根据近四年的使用情况,我们对原《工程数学》讲义进行了修改,形成了本书。

本书分上、下两册,上册由数值分析和矩阵论两部分组成,下册由数理统计和随机过程两部分组成。

本书为下册。数理统计部分内容通过介绍实际例子引进统计概念,阐明统计思想、统计理论与统计方法。其内容包括统计的基本概念,参数的估计、检验方法,回归分析和方差分析。随机过程部分对随机过程理论与方法作了浅显的介绍——主要介绍了马尔可夫链,平稳随机过程与平稳时间序列分析的理论与方法。书中内容力求精简,循序渐进,易于教学。

数理统计部分由何迎晖、蒋凤瑛编写完成;随机过程部分由何迎晖、钱伟民编写完成。

本书在出版过程中,得到同济大学应用数学系领导和教师的大力支持,同济大学研究生院给予大力帮助,特此表示感谢。

由于编者水平有限,错误和不妥之处在所难免,恳请读者指正。

编 者

2002 年 1 月于同济大学

目 录

第Ⅲ部分 数理统计

第一章 数理统计的基本概念

§ 1.1 数理统计问题	(3)
§ 1.2 总体与样本	(5)
§ 1.3 经验分布函数与直方图	(8)
§ 1.4 统计量	(12)
§ 1.5 三个常用分布	(17)
§ 1.6 正态总体下的抽样分布	(23)

第二章 参数估计

§ 2.1 参数估计问题	(29)
§ 2.2 求点估计的两种常用方法	(31)
§ 2.3 估计量的评选标准	(39)
§ 2.4 置信区间	(46)
§ 2.5 正态总体中未知参数的置信区间	(49)
§ 2.6 0-1 分布中未知参数的置信区间	(57)

第三章 假设检验

§ 3.1 假设检验问题	(60)
§ 3.2 正态总体中未知参数的假设检验	(64)
§ 3.3 0-1 总体中未知参数的假设检验	(76)
§ 3.4 χ^2 拟合优度检验	(78)
§ 3.5 独立性检验	(83)
§ 3.6 秩和检验与符号检验	(86)

第四章 回归分析与方差分析

§ 4.1 相关关系问题	(96)
§ 4.2 一元线性回归分析	(99)

§ 4.3 多元线性回归分析简介	(114)
§ 4.4 单因子方差分析	(124)
§ 4.5 双因子方差分析	(129)
§ 4.6 正交试验设计方法简介	(139)
习题 III	(148)
习题 III 答案	(155)
参考书目 III	(158)

第 IV 部分 随机过程

第一章 随机过程的基本概念

§ 1.1 随机过程	(161)
§ 1.2 随机过程的分布	(164)
§ 1.3 随机过程的数字特征	(167)
§ 1.4 正态随机过程	(170)
§ 1.5 二维随机过程与复值随机过程	(172)
§ 1.6 常用随机过程	(177)

第二章 马尔可夫链

§ 2.1 马尔可夫过程	(187)
§ 2.2 马尔可夫链及其转移概率	(190)
§ 2.3 切普曼-柯尔莫哥洛夫方程	(196)
§ 2.4 马尔可夫链的绝对分布	(200)
§ 2.5 遍历性与平稳分布	(204)
§ 2.6 转移概率矩阵的估计与检验	(209)

第三章 平稳随机过程

§ 3.1 平稳过程	(214)
§ 3.2 平稳过程的相关函数	(218)
§ 3.3 各态历经性	(224)
§ 3.4 平稳过程的谱密度	(228)
§ 3.5 线性系统中的平稳过程	(239)

§ 3.6 平稳性检验	(244)
第四章 平稳时间序列分析	
§ 4.1 时间序列及其应用	(248)
§ 4.2 平稳时间序列的线性模型	(250)
§ 4.3 线性模型的性质	(259)
§ 4.4 线性模型的参数估计	(264)
§ 4.5 线性模型的识别方法	(270)
§ 4.6 线性模型的预报方法	(278)
习题IV	(286)
习题IV 答案	(292)
参考书目IV	(295)
附 录 概率论内容基本要点	(296)
附 表	
一、标准正态分布函数值表	(317)
二、 χ^2 分布的分位数 $\chi_p^2(n)$ 值表	(318)
三、 t 分布的分位数 $t_p(n)$ 值表	(322)
四、 F 分布的分位数 $F_p(m, n)$ 值表	(323)
五、相关系数检验的临界值表	(326)
六、秩和检验的临界值表	(327)
七、游程检验的临界值表	(328)
八、常用分布及其数字特征	(329)

第Ⅲ部分 数理统计

数理统计是工程类型硕士生的一门应用性很强的重要基础课程。通过本课程的学习,学生要掌握数理统计的基本概念和基本原理,并侧重于数理统计方法的应用。通过本课程的学习,学生应初步具有运用数理统计的思想、方法处理随机性数据和解决工程实际问题的能力。

数理统计是数学的一个分支,它主要研究如何以有效的方法去收集、整理与分析带有随机性影响的数据,从而对所考察的问题作出推断和预测,直至为采取某种决策提供依据和建议。

数理统计方法的应用极其广泛,几乎在人类活动的一切领域中都能不同程度地找到它的应用,其原因在于实验是科学研究的根本方法,而随机性因素对实验结果的影响是无所不在、无时不有的。另一方面,应用上层出不穷的需要又推动了数理统计的发展。

第一章 数理统计的基本概念

数理统计是一门应用性很强的数学课程. 数理统计应用的广泛性决定了这门学科的重要性. 最近几十年以来, 建立在概率论基础上的数理统计, 在理论上、方法上与应用上都得到了迅速的发展. 特别是计算机技术的长足进步, 为数理统计的普遍应用提供了广阔的前景.

本章介绍数理统计中的一些基本知识, 内容包括总体、样本与统计量. 在此基础上, 将用概率论的方法来讨论一些常用的抽样分布.

§ 1.1 数理统计问题

数理统计是处理带有随机性影响的数据的一门学科. 为了使读者对这门学科的概貌有一个了解, 下面先举一个例子.

例 1.1 某食品厂为了加强质量管理, 在某天生产的一大批罐头中抽查了 100 个, 测得内装食品的净重数据(单位: g)如下:

342	341	348	346	343	342	346	341	344	348
346	346	341	344	342	344	345	340	344	344
343	344	342	342	343	345	339	350	337	345
349	336	348	344	345	332	342	341	350	343
347	340	344	353	341	340	353	346	345	346
341	339	342	352	342	350	348	344	350	335
340	338	345	345	349	336	342	338	343	343
341	347	341	347	344	339	347	348	343	347
346	344	345	350	341	338	343	339	343	346
342	339	343	350	341	346	341	345	344	342

试问该天生产的罐头中食品净重不足 340g 的概率有多大?

按照概率论的方法, 设随机地抽取一个罐头内装食品的净重为 X g, 所求的概率为 $P(X < 340)$. 然而, 在实际问题中, 随机变量 X 的分布是不清楚的, 因此无法算出这个概率的具体数值.

数理统计方法要利用已有的 100 个数据所提供的信息,因为这 100 个数据是对随机变量 X 作 100 次观测(即做了 100 次随机试验)得到的,它们的值反映了随机变量 X 的分布. 这 100 个数据中最小值为 332,最大值为 358. 我们把这 100 个数据所属的区间 $(331.5, 358.5]$ 等分成长度为 3 的 9 个小区间(称为组),并分别算出频数(即每组中数据的个数)与频率(即每组中数据的个数与数据总个数之比),列出下表:

序号 j	组 $(a_{j-1}, a_j]$	频数 n_j	频率 f_j
1	$(331.5, 334.5]$	1	0.01
2	$(334.5, 337.5]$	4	0.04
3	$(337.5, 340.5]$	12	0.12
4	$(340.5, 343.5]$	32	0.32
5	$(343.5, 346.5]$	30	0.30
6	$(346.5, 349.5]$	12	0.12
7	$(349.5, 352.5]$	7	0.07
8	$(352.5, 355.5]$	1	0.01
9	$(355.5, 358.5]$	1	0.01

根据上表中 9 个组 $(a_{j-1}, a_j]$ 及其相应的频数 n_j (或频率 f_j) ($j=1, \dots, 9$) 作出如图 1-1 的图形,这个图形称为直方图.

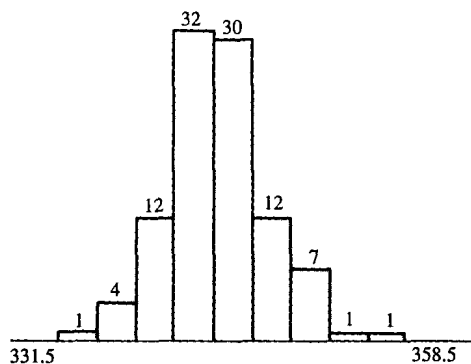


图 1-1 例 1.1 中的直方图

从直方图看,可以认为 X 大致上服从正态分布 $N(\mu, \sigma^2)$. 于是,所求概率为 $P(X < 340) = \Phi\left(\frac{340 - \mu}{\sigma}\right)^{\textcircled{1}}$. 尽管如此,我们还是无法得到这个概率的具体数值,因为 μ 与 σ 的值未知. 数理统计提供了一些方法,使我们可以从这 100 个数据出发对 μ 与 σ 的值作出推测. 例如,以后我们将会看到,可以推测 μ 为 343.83, σ 为 4.04. 这样,所求概率约为

$$P(X < 340) = \Phi\left(\frac{340 - 343.83}{4.04}\right) = \Phi(-0.948) = 0.1716$$

例 1.2 某厂要检查一大批产品的质量,希望知道这批产品的不合格率 p ,常用的方法是:从这批产品中随机地抽取若干样品进行测试,以样品中的不合格率来推测整批产品的不合格率 p .

从例 1.1 与例 1.2 可以看出,数理统计的内容有以下两个特点:

(1) 数理统计的出发点是数据,因此必须先收集一批数据. 如何收集数据? 这是数理统计的两个重要分支——抽样方法与试验设计的基本内容.

(2) 有了数据之后,我们需要根据数据对所关心的问题(例如,例 1.1 中的概率 $P(X < 340)$ 与例 1.2 中的 p)进行推测. 当然,这种推测是不可能绝对准确的,它总含有一定程度的不确定性,而概率正是不确定性大小的数量表示. 这样,任何一种推测都必须用一定的概率来表明推测的可靠程度. 这种在一定的概率意义下的推测称为统计推断. 统计推断的基本形式是估计(包括点估计与区间估计)和检验.

本书主要介绍统计推断的基本原理与方法,除第四章 § 4.6 外,一般不涉及抽样方法与试验设计的内容.

§ 1.2 总体与样本

在数理统计中,我们把研究对象的全体称为总体(或母体),把组成总体的每个成员称为个体. 例如,在例 1.1 中,该天生产的所有罐头构成了一个总体,而每一个罐头是个体. 在这个问题中,实际关心的仅仅是罐头内装食品的净重. 因此,为了方便,我们也把该天生产的所有各个罐头的净重看成一个总体,每个罐头的净重视为个体. 在例 1.2 中,这一大批产品构成了一个总体,而每一件产品是个体. 在这个问题中,实际关心的仅仅是产品的质量等级(合格或不合格). 因此,为了方便,我们也把

① $\Phi(x)$ 是标准正态分布 $N(0, 1)$ 的分布函数,即 $\Phi(x) = \int_{-\infty}^x \frac{1}{\sqrt{2\pi}} e^{-\frac{t^2}{2}} dt, -\infty < x < \infty$.

这一大批产品中所有各个产品的质量等级看成一个总体(等级允许重复),每一件产品的质量等级视为个体. 这里,产品的质量等级是定性描写的,我们不难把它用一个数值来表示. 通常用数字“0”表示合格品,用数字“1”表示不合格品.

通过上面例子可以看到,在实际问题中,我们所关心的往往是总体中的个体的某个数值指标. 不妨把个体的这个数值指标看作是随机变量,记作 X ,因为在实际测得一个罐头的净重(见例 1.1)或一件产品的质量(见例 1.2)之前,试验结果是不确定的. X 的分布称为总体分布, X 的分布函数称为总体分布函数,记作 $F(x)$. 在例 1.2 中,随机变量 X 服从 0-1 分布 $B(1, p)$,因为随机地抽取一件产品作质量检查,试验结果或者是“ $X=0$ ”(表示被检查的那件产品是合格品),或者是“ $X=1$ ”(表示被检查的那件产品是不合格品).

随机变量 X 的分布(或分布函数)反映了总体数值指标取值的统计规律性,因此,以后我们常说“总体 X 的分布函数为 $F(x)$ ”. 这意味着我们把总体用与其相应的随机变量 X 来表示.

当 X 是离散型随机变量时,相应地有总体概率函数;当 X 是连续型随机变量时,相应地有总体密度函数. 我们把两者都记作 $f(x)$. 这样将会给今后的讨论带来叙述上的方便. 这里要注意,当 X 为离散型随机变量时,概率函数 $f(x)$ 的定义域不是 $(-\infty, \infty)$,而是仅含有限个或可列个元素的数值,且 $f(x) \triangleq P(X=x)$ ^①. 在例 1.2 中,总体 $X \sim B(1, p)$,因此, X 的概率函数

$$f(x) = \begin{cases} 1-p, & x=0 \\ p, & x=1 \end{cases}$$

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

它就是概率函数.

X	0	1
P_r	$1-p$	p

在数理统计中,总体分布永远是未知的. 在例 1.1 中,罐头的净重究竟服从什么分布是不清楚的;即使有足够的理由可以认为罐头的净重服从正态分布 $N(\mu, \sigma^2)$,但 μ 与 σ^2 的值还是未知的. 在例 1.2 中,虽然我们可以确定总体 $X \sim B(1, p)$,但 p 的值有多大是不清楚的. 诚然,我们把这一大批产品逐个作检查是可以得到 p 值的. 但是,一则产品很多,逐个检查很不经济;二则在破坏性试验的情形(例如产品质量是

① “ $\triangleq$ ”表示“记作”,这里表示规定 $f(x) = P(X=x)$.

灯泡寿命)下这是不现实的. 我们希望从客观存在的总体中按一定原则选取一些个体(即抽样), 通过对这些个体作观察或测试来推断关于总体分布的某些量(例如总体 X 的均值、方差、中位数等). 观测到的这些个体的值便是实际问题中常见的数据.

在数理统计中, 称数据为样本观测值, 记作 $(x_1, \dots, x_n)$, 称 n 为样本大小^①. 在作观测前, 样本观测值是不确定的, 为了体现随机性, 记作 $(X_1, \dots, X_n)$, 称为样本(或子样). 因而样本 $(X_1, \dots, X_n)$ 是一个 n 维随机向量. 这个随机向量 $(X_1, \dots, X_n)$ 可能取值的全体(即值域)称为样本空间. 样本空间可以是 n 维欧氏空间, 也可以是 n 维欧氏空间的一个子集. 样本观测值 $(x_1, \dots, x_n)$ 是样本空间中的一个点.

从理论统计工作者的立场看, 样本是一个 n 维随机向量; 从应用统计工作者的立场看, 样本就是一批已知的数据. 这是因为前者站在抽样前的立场上, 而后者站在抽样后的立场上. 样本的这种二重性虽然是一件平凡的事情, 但却很重要. 对理论统计工作者而言, 为了使统计方法不仅仅适用于某些具体样本观测值而带有一定的普遍性, 必须更多地注意到样本是随机向量; 对应用统计工作者而言, 虽然他们已经习惯于把样本看成数据, 但仍要注意“样本是随机向量”这一概率背景, 不然的话, 样本便成了一堆杂乱无章、毫无规律性可言的数字, 无法进行任何统计处理.

在实际问题中, 可以有各种不同的方法从总体中抽取样本. 本书主要涉及一种简单随机抽样方法, 用这种抽样方法得到的样本称为简单随机样本. 由于本书仅仅涉及简单随机样本, 因此, 今后提到“样本”均指简单随机样本.

简单随机抽样相当于概率论中的有放回抽样. 当我们从某个总体中抽取大小为 n 的样本时, 每抽取一个个体作观测后立即放回并搅匀, 然后再抽取下一个个体. 当然, 实际情形中往往采用的是无放回抽样. 但是, 当样本大小 n 与总体所含个体的个数 N 之比很小(一个经验法则是 $\frac{n}{N} \leq 0.1$) 时, 可以近似地把无放回抽样看作是有放回抽样.

在作简单随机抽样时, 由于在抽取每个样本分量 X_i 前所面临的总体的状况都相同, 因此, $X_1, \dots, X_n$ 的分布都相同, 且都与总体分布相同; 另外, 由于这种抽样是有放回抽样, 因此, 样本分量 $X_1, \dots, X_n$ 相互独立. 以后我们常常用到术语“ $(X_1, \dots, X_n)$ 是取自总体 X 的一个样本”, 这里蕴含了上述“同分布”与“独立性”的概率背景.

设 $(X_1, \dots, X_n)$ 是取自总体 X 的一个样本. 如果 X 的分布函数为 $F(x)$, 那么, 由于 $X_1, \dots, X_n$ 是独立同分布的随机变量, 且每一个 X_i 的分布函数都是 $F(x)$, $i=1, \dots, n$, 因此, 样本 $(X_1, \dots, X_n)$ 的分布函数为

① 有些书上称 n 为样本容量.

$$F^*(x_1, \dots, x_n) = \prod_{i=1}^n F(x_i), \quad -\infty < x_1, \dots, x_n < \infty.$$

类似地, 当总体 X 的密度函数(或概率函数)为 $f(x)$ 时, 样本 $(X_1, \dots, X_n)$ 的密度函数(或概率函数)为

$$f^*(x_1, \dots, x_n) = \prod_{i=1}^n f(x_i).$$

在例 1.2 中, 总体 $X \sim B(1, p)$, 它的概率函数为

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

因此样本 $(X_1, \dots, X_n)$ 的概率函数为

$$f^*(x_1, \dots, x_n) = p^{\sum_{i=1}^n x_i} (1-p)^{n - \sum_{i=1}^n x_i}, \\ x_i = 0, 1, i=1, \dots, n.$$

在例 1.1 中, 假定知道罐头净重 $X \sim N(\mu, \sigma^2)$, 那么, 样本 $(X_1, \dots, X_n)$ 的密度函数为

$$f^*(x_1, \dots, x_n) = \prod_{i=1}^n \left(\frac{1}{\sqrt{2\pi}\sigma} \exp \left\{ -\frac{(x_i - \mu)^2}{2\sigma^2} \right\} \right) \\ = (2\pi\sigma^2)^{-\frac{n}{2}} \exp \left\{ -\frac{1}{2\sigma^2} \sum_{i=1}^n (x_i - \mu)^2 \right\}, \\ -\infty < x_1, \dots, x_n < \infty.$$

上面介绍的一切概念都可以推广到需要考察的总体数值指标有 k 个的情形, $k \geq 2$. 例如, 我们要同时考察某种圆柱形零件的长度与外径, 这时, 相应地有二维总体, 二维随机向量 (X, Y) , 以及二维总体分布, 二维总体分布函数 $F(x, y)$, 等等. 大小为 n 的样本 $(X_1, Y_1; \dots, X_n, Y_n)$ 构成一个 $2n$ 维随机向量, 它的观测值 $(x_1, y_1; \dots; x_n, y_n)$ 仍是样本空间中的一个点, 但样本空间是 $2n$ 维欧氏空间或 $2n$ 维欧氏空间的一个子集. 对 k 维总体的情形是类似的.

§ 1.3 经验分布函数与直方图

设 $(X_1, \dots, X_n)$ 是取自总体 X 的一个大小为 n 的样本, $(x_1, \dots, x_n)$ 是样本观测值. 定义一个函数

$$F_n(x) \triangleq \frac{1}{n} \cdot \{X_1, \dots, X_n \text{ 中} \leq x \text{ 的个数}\}, \\ -\infty < x < \infty,$$

称这个函数 $F_n(x)$ 为经验分布函数;相应地,称函数

$$\tilde{F}_n(x) \triangleq \frac{1}{n} \cdot \{x_1, \dots, x_n \text{ 中小于或等于 } x \text{ 的个数}\},$$

$$-\infty < x < \infty$$

为经验分布函数的观测值.

这里要注意,经验分布函数的观测值 $\tilde{F}_n(x)$ 与通常意义下的分布函数相同,它具有分布函数的一切性质. 不难由 $\tilde{F}_n(x)$ 的定义式看出, $\tilde{F}_n(x)$ 是一个跳跃度为 $\frac{1}{n}$ (或 $\frac{1}{n}$ 的整数倍) 的非降阶梯函数. 经验分布函数则不然,因为给定自变量 x 时, $F_n(x)$ 的取值不是一个数,而是一个随机变量.

在实际计算 $\tilde{F}_n(x)$ 时,用下述方法比较方便. 假定样本观测值 $(x_1, \dots, x_n)$ 中有 n_1 个 x_1^* , $\dots$, n_m 个 x_m^* , 其中 $x_1^* < \dots < x_m^*$. 记频率 $f_j = \frac{n_j}{n}$, $j=1, \dots, m$, 即样本观测值的频率分布为

样本观测值	x_1^*	x_2^*	$\dots$	x_m^*
出现的频率	f_1	f_2	$\dots$	f_m

不难验证经验分布的观测值为

$$\tilde{F}_n(x) = \begin{cases} 0, & \text{当 } x < x_1^* \\ \sum_{j=1}^k f_j, & \text{当 } x_k^* \leq x < x_{k+1}^*, k=1, \dots, m-1, \\ 1, & \text{当 } x \geq x_m^*, \end{cases}$$

其中, $\sum_{j=1}^k f_j$ 称为相应于 x_k^* 的累积频率, $k=1, \dots, m$.

例 1.3 对某厂生产的电子仪器作寿命测试,得到样本观测值(单位:100 小时)为

5 3 7 5 4

试求经验分布函数的观测值 $\tilde{F}_5(x)$ ①.

解 先写出样本观测值(按从小到大次序排列)的频率分布:

① 在实际统计问题中,样本大小 n 往往很大. 为了便于计算,对应用统计方法的说明性例子与习题,本书有时取较小的 n 值.

样本观测值	3	4	5	7
出现的频率	$\frac{1}{5}$	$\frac{1}{5}$	$\frac{2}{5}$	$\frac{1}{5}$

把上述表格看成是某个离散型随机变量的概率函数, 与其相应的分布函数便是经验分布函数的观测值:

$$\tilde{F}_5(x) = \begin{cases} 0, & x < 3 \\ \frac{1}{5}, & 3 \leq x < 4 \\ \frac{2}{5}, & 4 \leq x < 5 \\ \frac{4}{5}, & 5 \leq x < 7 \\ 1, & x \geq 7 \end{cases}$$

图 1-2 给出了 $\tilde{F}_5(x)$ 的图像.

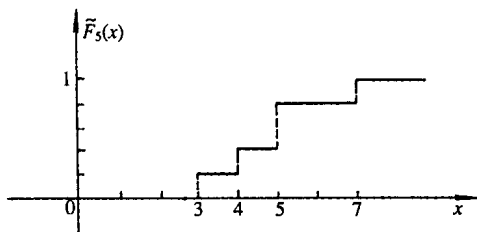


图 1-2 例 1.3 中经验分布函数观测值 $\tilde{F}_5(x)$

由经验分布函数 $F_n(x)$ 的定义式知道, $F_n(x)$ 是样本 $(X_1, \dots, X_n)$ 中不大于 x 的个体个数所占的比例. 另一方面, 对每一个固定的 x , 总体分布函数 $F(x)$ 是总体中不大于 x 的个体个数所占的比例 (即概率). 自然会想到, 未知的总体分布函数是否可以用能观测到 (即能根据数据具体算出) 的经验分布函数去近似? 这当然取决于 $|F_n(x) - F(x)|$ 这个量的大小. 下面的理论结果表明这种做法是合理的.

定理 1.1 设 $(X_1, \dots, X_n)$ 是取自总体 X 的一个样本, 总体分布函数为 $F(x)$. 对任意一个 $x \in (-\infty, \infty)$ 与任意一个 $\epsilon > 0$,

$$\lim_{n \rightarrow \infty} P(|F_n(x) - F(x)| \geq \epsilon) = 0.$$

证 对任意一个固定的 $x \in (-\infty, \infty)$, 定义随机变量

$$Y_i = \begin{cases} 1, & \text{当 } X_i \leq x, \\ 0, & \text{当 } X_i > x, \end{cases} \quad i = 1, \dots, n.$$